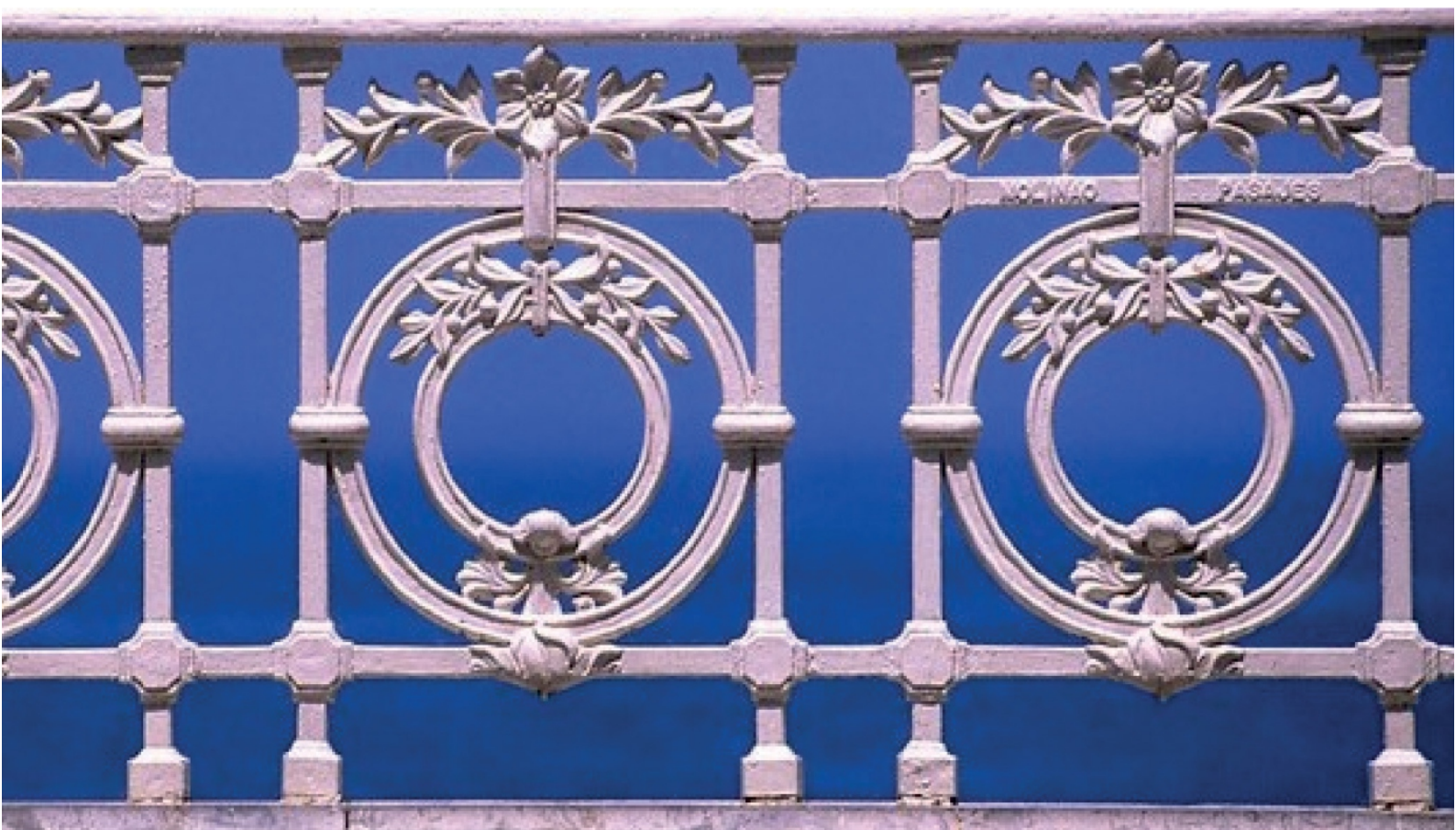


ESTYLF 2016



XVIII Congreso Español Sobre
Tecnologías y Lógica Fuzzy



Libro de Resúmenes

XVIII Congreso Español sobre Tecnologías y Lógica Fuzzy



Libro de Resúmenes

Donostia-San Sebastián, 25-27 de mayo de 2016

ORGANIZADORES:

Departamento de Matemática Aplicada.
Universidad del País Vasco/Euskal Herriko Unibertsitatea

EDITORES:

Cristina Alcalde
Humberto Bustince
Francisco Javier Fernández

ENTIDADES PATROCINADORAS:

Universidad del País Vasco/Euskal Herriko Unibertsitatea
Escuela de Ingeniería de Gipuzkoa
Grupo de Investigación de la UPV/EHU «*Sistemas Inteligentes y Energía (SI+E)*»
European Society for Fuzzy Logic and Technology
Red temática de excelencia «*Lógica Difusa y Soft Computing (LODISCO)*»



GIPUZKOAKO
INGENIARITZA
ESKOLA
ESCUELA
DE INGENIERÍA
DE GIPUZKOA



Red temática de Excelencia:
**Lógica Difusa y Soft Computing
(LODISCO)**

ISBN: 978-84-608-7381-5

Presentación

En 2016, coincidiendo con su designación como *Capital Europea de la Cultura*, Donostia-San Sebastián acogerá del 25 al 27 de mayo la celebración del *XVIII Congreso Español sobre Tecnologías y Lógica Fuzzy (ESTYLF 2016)*.

El ESTYLF celebra en esta especial ocasión sus 25 años de existencia desde que tuvo lugar su primera edición en Granada en 1991. Desde ese momento, el congreso ha constituido un foro de discusión, abierto a todas las personas interesadas en la Lógica Fuzzy y sus aplicaciones, en el que presentar resultados, debatir ideas y exponer proyectos relacionados con el área.

Este volumen recopila los resúmenes de las 5 conferencias plenarias y los 117 trabajos aceptados para su presentación en el congreso. Se han organizado 13 sesiones especiales cuyas diferentes temáticas constituyen una muestra de la extensa variedad de aplicaciones en las que la lógica y las tecnologías fuzzy tienen cabida.

Deseamos agradecer a todos los que han hecho posible el buen desarrollo que este congreso, incluyendo a autores, miembros de los comités, organizadores de sesiones especiales, evaluadores, patrocinadores y, en general, a todas aquellas personas que de una u otra manera han contribuido a su celebración.

Donostia-SanSebastián, Mayo de 2016

Cristina Alcalde
Humberto Bustince
Francisco Javier Fernández

Organización

Presidente Honorífico

Enric Trillas (European Centre for Soft Computing)

Presidenta del Comité Organizador

Cristina Alcalde (Universidad del País Vasco/Euskal Herriko Unibertsitatea)

Comité Organizador

Mariano Jiménez (Universidad del País Vasco/Euskal Herriko Unibertsitatea)

Itziar Zubia (Universidad del País Vasco/Euskal Herriko Unibertsitatea)

Xabier Ostolaza (Universidad del País Vasco/Euskal Herriko Unibertsitatea)

Ana Burusco (Universidad Pública de Navarra)

Aránzazu Jurío (Universidad Pública de Navarra)

Daniel Paternain (Universidad Pública de Navarra)

Luka Eciolaza (IK4-Lortek)

Comité del Programa

Javier Montero (Universidad Complutense de Madrid)

Alberto Bugarín (Universidad de Santiago de compostela)

Enrique Herrera-Viedma (Universidad de Granada)

Eduarne Barrenechea (Universidad Pública de Navarra)

Francesc Esteva (IIIA-CSIC)

Sesiones Especiales

Humberto Bustince (Universidad Pública de Navarra)

Francisco Javier Fernández (Universidad Pública de Navarra)

Contenidos

Conferencias Plenarias

<i>The age of aggregation</i>	1
Bernard de Baets	
<i>El reto del uso de los Sistemas Difusos en Big Data y Ciencia de Datos</i>	2
Francisco Herrera	
<i>De las funciones de agregación a las multidistancias: toda una experiencia</i>	4
Gaspar Mayor	
<i>Vínculos en entornos difusos</i>	5
Manuel Ojeda-Aciego	
<i>Cosas que he aprendido en mis cuarenta años de lógica borrosa</i>	6
Enric Trillas	

Sesión regular

<i>Computing the whole set of adjoint negations</i>	8
M.Eugenia Cornejo, Francesc Esteva, Jesús Medina, Eloísa Ramírez Poussa	
<i>Computing the minimal solutions of fuzzy relation equations</i>	10
Juan Carlos Díaz Moreno, Jesús Medina, Esko Turunen	
<i>Algoritmos de Lógica Difusa para el Análisis del Deterioro de Electrodo en Soldadura por Resistencia</i> Luka Eciolaza, Ander Muniategui, Mikel Ayuso, Marijo Garmendia, Pedro Alvarez	12
<i>Sobre la relación entre operadores de consecuencias y preórdenes borrosos usando conjuntos</i> ...	14
Jorge Elorza, S. Montes, C. Bejines, J. Bragard	
<i>Sobre los conceptos de dual y anulador de un grupo topológico abeliano fuzzy</i>	16
Sergio Ardanza-Trevijano, María Jesús Chasco, Jorge Elorza	
<i>Utilización de Prototipos Deformables Borrosos en la Generación de Recomendaciones sobre Gestión del Tiempo</i>	18
Juan García-Aguilar, Francisco Romero, Jose A. Olivas, Jesus Serrano-Guerrero	
<i>Control de suspensión activa basado en lógica difusa</i>	20
Javier López de la Nieta, Matilde Santos Peñas	
<i>Extended Set of Hesitant Fuzzy Linguistic Term</i>	22
Jordi Montserrat-Adell, Núria Agell, Monica Sanchez, Francisco J. Ruiz	
<i>Storing heterogeneous information in fuzzy ontologies using multi-granular fuzzy linguistic methods</i> .	24
Juan Antonio Morente-Molinera, Ignacio Javier Pérez, Francisco Javier Cabrerizo, Sergio Alonso, Enrique Herrera-Viedma	
<i>Convergence and parameter learning of Fuzzy Cognitive Maps</i>	26
Gonzalo Nápoles-Ruiz, Elpiniki Papageorgiou, Rafael Bello, Koen Vanhoof	
<i>A Fuzzy System to Assign Reviewers for Conference Paper Evaluation</i>	28
Jennifer Nguyen, Germán Sánchez, Núria Agell, Cecilio Angulo	
<i>El cálculo modal de la borrosidad en el <i>Pharum Scientiarum</i> (1659) de Sebastián Izquierdo (1601-1681). Una posible prolongación de la lógica fuzzy de Lofti Zadeh.</i>	30
Carlos Ortiz de Landázuri	

<i>Una propuesta para el modelado del conocimiento del sentido común en la Programación Lógica</i> ...	32
Clemente Rubio-Manzano, Martín Pereira-Fariña	
<i>Uniform Interpolants in Nilpotent Minimum Logic</i>	34
Diego Valota	

Conditional, implication and indistinguishability in ordinary reasoning

Organizadores: Isabel Aguiló, Gaspar Mayor, Enric Trillas

<i>Condicionalidad y Continua Lingüísticos</i>	36
Gaspar Mayor, Enric Trillas	
<i>Modelos de conjeturas y conexiones de Galois</i>	38
Adolfo Rodríguez de Soto	
<i>Razonamiento difuso mediante modelos CHC y representación por niveles</i>	40
Daniel Sánchez	
<i>Enunciados condicionales y enunciados causales, precisos e imprecisos</i>	42
Alejandro Sobrino	

Ecuaciones funcionales relacionadas con operadores de agregación y subconjuntos fuzzy

Organizadores: María Jesús Campión, Esteban Induráin

<i>Ecuaciones funcionales relacionadas con conjuntos fuzzy, ordenaciones representables y quasi-métricas ponderadas</i>	44
María Jesús Campión, Raquel G. Catalán, Esteban Indurain, Gustavo Ochoa, Oscar Valero	
<i>Conjuntos Fuzzy y la Ecuación de Separabilidad</i>	46
María Jesús Campión, Raquel G. Catalán, José Garay, Esteban Induráin	
<i>Characterizations of lattice OWA operators</i>	48
Laura De Miguel, Daniel Paternain, Gustavo Ochoa, Inmaculada Lizasoain	
<i>Convolución de funciones entre retículos</i>	50
Laura De Miguel, Humberto Bustince, Bernard De Baets	
<i>Relaciones binarias definidas a través de ecuaciones funcionales</i>	52
María Jesús Campión, Raquel G. Catalán, Laura De Miguel, Esteban Induráin	
<i>Point-sensitive aggregation operators: functional equations</i>	54
Juan Carlos Candeal, Esteban Induráin	

Funciones de agregación y conectivos lógicos

Organizadores: Sebastià Massanet, Juan Vicente Riera, Daniel Ruiz-Aguilera, Joan Torrens

<i>Una introducción a las FNI-implicaciones</i>	56
Isabel Aguiló, Jaume Suñer, Joan Torrens	
<i>Negaciones para un orden admisible en conjuntos fuzzy valorados en intervalos</i>	58
María José Asiain Olló, Humberto Bustince, Javier Fernández, Razko Mesiar, Anna Kolesárová	
<i>La noción de función de pre-agregación</i>	60
Humberto Bustince, Aranzazu Jurio, Daniel Paternain, Mikel Galar, José Sanz, Mikel Elkano, Giancarlo Lucca, Radko Mesiar, Anna Kolesarova, Graçaliz Dimuro	
<i>Funciones monótonas direccionalmente ordenadas</i>	62
Humberto Bustince, Edurne Barrenechea, Miguel Pagola, Javier Fernandez, Laura de Miguel, Carlos Lopez-Molina, Julio Lafuente, Mikel Sesma, Radko Mesiar, Anna Kolesárová	
<i>La clase $(\leq^\omega)_{\omega \in L}$ de órdenes de actividad en retículos distributivos $(L, \leq)$ y su familia de nulnormas asociada $(\wedge^\omega)_{\omega \in L}$</i>	64
Ramón Fuentes-González	

<i>Hesitant fuzzy relations and their transitive closure</i>	66
Luis Garmendia, Ramón González del Campo, Jordi Recasens	
<i>Agregaciones computables</i>	68
Ramón González del Campo ¹ , Luis Garmendia, Javier Montero, J. Tinguaro Rodríguez, Daniel Gómez	
<i>On some properties of SUOWA operators</i>	70
Bonifacio Llamazares	
<i>OWAs defining multidistances and measures of dispersion on R^k</i>	72
Javier Martín, Gaspar Mayor	
<i>Relating absolute dispersion measures, multidistances and remainders of aggregation functions</i>	74
Miguel Martínez-Panero, José Luis García Lapresta, Luis Carlos Meneses	
<i>Estudio del Modus Ponens respecto de una uninorma</i>	76
Margarita Mas, Miquel Monserrat, Daniel Ruiz-Aguilera, Joan Torrens	
<i>Sobre Implicaciones Racionales</i>	78
Sebastià Massanet Massanet, Juan Vicente Riera, Daniel Ruiz-Aguilera	
<i>Sumas ordinales de funciones de implicación borrosas</i>	80
Sebastià Massanet Massanet, Juan Vicente Riera, Joan Torrens	
<i>Qualitative orness for OWA operators defined on lattice intervals</i>	82
Gustavo Ochoa, Daniel Paternain, Inmaculada Lizasoain, Humberto Bustince	
<i>Órdenes Parciales para Hesitant Fuzzy Sets</i>	84
Luis Garmendia, Ramón González del Campo, Jordi Recasens	

Herramientas de soft computing para la toma de decisiones en ambientes de incertidumbre

Organizadores: Rocío de Andrés, Teresa González Arteaga, José Carlos Rodríguez Alcantud

<i>A fuzzy approach to the problem of group identification</i>	86
José Carlos R. Alcantud, Rocio de Andres Calle	
<i>Ajuste de dosis de insulina mediante redes neuronales y control borroso</i>	88
Ana García Alvarez, Matilde Santos Peñas	
<i>A new consensus measure based on Pearson correlation coefficient</i>	90
Teresa González-Arteaga, Rocio de Andres Calle, Francisco Chiclana	
<i>Una nueva visión de los Z-números de Zadeh</i>	92
Sebastià Massanet Massanet, Juan Vicente Riera, Joan Torrens	
<i>The weighted Hamming distance as a divergence measure</i>	94
Susana Montes, Jorge Elorza, Carlos Bejines*, Jean Bragard	
<i>Aplicación del uso de valoraciones hesitant lingüísticas en una red social de economía colaborativa</i> .	96
Rosana Montes, Ana María Sánchez, Pedro Villar, Francisco Herrera	
<i>A Novel Consensus Model for Group Decision Making with Hesitant Fuzzy Linguistic Information</i>	98
Rosa M. Rodríguez, Luis Martinez	
<i>Bell-type inequalities and fuzzy probability</i>	100
Agnieszka Siwec, Susana Díaz, Bernard De Baets, Susana Montes	

Lógica borrosa y educación

Organizadora: Itziar García Honrado

<i>Presentación de una red social docente que integra un sistema de recomendaciones lingüístico difuso para personalizar el acceso a recursos educativos</i>	102
Alberto Ching-López, Carlos Porcel, Enrique Herrera-Viedma	

<i>Estudio del aprovechamiento de oportunidades de aprendizaje matemático en un aula de 3º de la ESO</i> Miquel Ferrer, Josep María Fortuny, Itziar García-Honrado	104
<i>Propuesta interdisciplinar en Educación Secundaria: creación de un modelo con lógica fuzzy para la hormonal del ciclo ovárico.</i>	106
Edith Herrera, Itziar García-Honrado	
<i>Estructura formal de los sistemas cognitivos desde el Diagrama de Marlo</i>	108
Marcos López Aznar	
<i>Extensión borrosa del método de aprendizaje de Leitner</i>	110
José Otero, Rosario Suárez, Luís Junco	
<i>EvalGRAPHS: herramienta para implementar evaluadores inteligentes. Metodología de estudio del comportamiento de los evaluadores</i>	112
Gloria Sanchez-Torruboa, Carmen Torres-Blanc, Silvia García Guerrero, Ruben Gonzalez-Martin	
<i>Un software de modelado de color basado en lógica difusa como herramienta docente para aprender el concepto de imprecisión del color de manera visual</i>	114
Jose Manuel Soto Hidalgo, Pedro Manuel Martínez Jiménez, Jesús Chamorro-Martínez, Daniel Sánchez	
<i>¿Qué debe saberse para emplear la lógica fuzzy?</i>	116
Enric Trillas	

Lógica difusa, análisis de conceptos formales, conjuntos rugosos y temas relacionados

Organizadores: Ana Burusco, Jesús Medina, Manuel Ojeda, Cristina Alcalde

<i>Hipercontextos L-fuzzy</i>	118
Cristina Alcalde, Ana Burusco	
<i>Fuzzy adjunction revisited</i>	120
Inma P. Cabrera, Pablo Cordero, Bernard De Baets, Francisca García-Pardo, Manuel Ojeda-Aciego	
<i>Towards attribute implications in multi-adjoint context</i>	122
M.Eugenia Cornejo, Jesús Medina, Eloísa Ramírez Poussa, , Rafael Rodríguez Galván,	
<i>On the size of reducts in multi-adjoint contexts</i>	124
M.Eugenia Cornejo, Jesús Medina, Eloísa Ramírez Poussa	
<i>Similarity-based reasoning using prototypes and counterexamples</i>	126
Soma Dutta, Francesc Esteve, Lluís Godó Lacasa	
<i>Extensión a la Lógica Borrosa de algunas nociones asociadas a Relaciones y a Grafos ordinarios utilizando Retículos de Conceptos Borrosos</i>	128
Ramón Fuentes-González	
<i>Nuevos órdenes sobre los números y cantidades borrosas</i>	130
Pablo Hernández Varela , Carmen Torres-Blanc, Susana Cubillo, José Atilio Guerrero	
<i>Application of F-transforms to Reduction of Relational Databases</i>	132
Petra Hodakova, Nicolas Madrid, Pavel Rusnok	
<i>Removing redundancy for attribute implications in data with grades</i>	134
Estrella Rodríguez-Lorenzo, Pablo Cordero, Manuel Enciso, Angel Mora Bonilla	
<i>Un método para obtener una representación difusa de series de tiempo</i>	136
Antonio Moreno Garcia, Juan Moreno Garcia	
<i>Aprendizaje de Reglas de Asociación Difusas en Flujo de Datos para el Análisis del Electroencefalograma en Psicofisiología</i>	138
Elena Ruiz, Pandelis Perakakis, Jorge Casillas	
<i>An Approach to Multilabel Classification using K-Formal Concept Analysis</i>	140
Francisco Valverde-Albacete, Carmen Peláez-Moreno	

Minería de datos, textos y web

Organizadores: M. Dolores Ruiz, Juan Gómez Romero, María José Martín Bautista, Daniel Sánchez

<i>On the Application of Data Mining Techniques to Graded Ontology Building</i>	142
Fernando Bobillo, M. Dolores Ruiz, Juan Gómez-Romero, Daniel Sánchez	
<i>Reglas de asociacion difusas en Big Data</i>	144
Carlos Fernandez-Basso, M. Dolores Ruiz, Maria Jose Martin-Bautista	
<i>Open Linked Data Mining for Environmental Scanning</i>	146
Juan Gómez-Romero, Fernando Bobillo, M. Dolores Ruiz, Miguel Delgado	
<i>Análisis de Tópicos en Redes Sociales mediante Técnicas de Clustering Difuso</i>	148
Karel Gutierrez-Batista, Jesús Campaña, Amparo Vila, Maria Jose Martin-Bautista	
<i>Fuzzy approximate dependencies in terms of representation by levels</i>	150
Carlos Molina, M. Dolores Ruiz, Daniel Sánchez, Jose Serrano	

Modelos de optimización y decisión 5D

Organizadores: Dagoberto Castellanos, Carlos Cruz Corona

<i>Sistema Híbrido AHP-Fuzzy Logic, para Toma de Decisiones en Entornos de Incertidumbre. Aplicación a FMA</i>	152
Luis Santacreu Rios, Alejandro Talavera Ortiz, Ricardo Aguasca Colomo, Maximo Mendez Babey, Blas Galvan Gonzalez, Dagoberto Castellano Nieves	
<i>The Team Orienteering Problem with Fuzzy Times</i>	154
Julio Antonio Brito Santana, Airam Expósito Márquez, José Andrés Moreno Pérez	
<i>FLMICMAC aplicado a las moratorias turísticas en las Islas Canarias</i>	156
Dagoberto Castellanos Nieves, Carlos Cruz, Julio Brito Santana	
<i>Solving a Fuzzy Multi-Criteria Location Problem by Using Georeferenced Data</i>	158
Christopher Expósito-Izquierdo, Airam Expósito-Márquez, Belén Melián-Batista, José Marcos Moreno-Vega	
<i>Método multicriterio basado en Soft Computing para evaluar la calidad de software</i>	160
Yamilis Fernández Pérez, Carlos Cruz, José Luis Verdegay Galdeano	
<i>Análisis de patrones en predicción de series temporales fuzzy</i>	162
A. Rubio, E. Vercher, J.D. Bermúdez	
<i>Banzhaf values for cooperative games with a proximity relation among the agents</i>	164
Inés Gallego, Julio R. Fernández, Andrés Jiménez-Losada, Manuel Ordoñez	
<i>Un enfoque multicriterio difuso para la selección de carteras socialmente responsables para inversores con aversión a las pérdidas</i>	166
Amelia Bilbao Terol, Mariano Jiménez López, Mar Arenas Parra, Verónica Cañal Fernández, Pablo Nguema Obama	
<i>Ordenación de fondos de inversión mediante valoración difusa de su responsabilidad social</i>	168
Maria Teresa Lamata, Vicente Liern, Blanca Pérez-Gladish	
<i>Programación Lineal Difusa y Modelos de Localización con Cobertura para Diseminación de Información en VANETs</i>	170
Antonio D. Masegosa, Idoia de la Iglesia, Unai Hernandez-Jayo	

Problemas de big data con técnicas de soft computing

Organizadores: José Antonio Sanz, Mikel Galar

<i>Optimización de la fase de aprendizaje del Sistema de Clasificación Basado en Reglas Difusas para Big Data Chi-FRBCS-BigDataCS</i>	172
Mikel Elkano, Mikel Galar, José Sanz, Humberto Bustince	

<i>Sistemas de Clasificación Basados en Reglas Difusas para Big Data con MapReduce: Análisis de la Granularidad</i>	174
Alberto Fernandez, Sara del Río, Francisco Herrera	
<i>Un estudio preliminar sobre el uso de técnicas de aprendizaje incremental para problemas de big data</i>	
Juan Carlos Gámez, David García, Antonio González, Raúl Pérez	176
<i>A MapReduce Implementation of a Genetic Fuzzy System for Regression</i>	178
Ismael Rodríguez, Manuel Mucientes, Alberto Bugarin	

Procesamiento de imágenes

Organizadores: Aránzazu Jurio, Daniel Paternáin

<i>Detección de estructuras curvilíneas usando la transformación morfológica borrosa todo-nada</i>	180
Pedro Bibiloni, Manuel González-Hidalgo, Sebastià Massanet Massanet	
<i>FPGA Implementation of the Two-Dimensional Fuzzy-ELA Algorithm for Image Enlargement</i>	182
Maria Brox Jimenez, Santiago Sánchez-Solano, Piedad Brox, Andrés Gersnoviez, Iluminada Baturone	
<i>On Fuzzy Modelling of Shape Vertices</i>	184
Jesús Chamorro-Martínez, Luis Suárez-Llorens	
<i>Face Recognition Using Color Local Binary Patterns: A Comparative Study</i>	186
Fadi Dornaika, Ignacio Gonzales, Yassine Ruichek	
<i>Agregación de imágenes de contornos usando OWAs</i>	188
Manuel González-Hidalgo, Daniel Ruiz-Aguilera, Arnau Mir, Sebastià Massanet Massanet	
<i>Una metodología para evaluar algoritmos de segmentación jerárquica</i>	190
Carely Guada, J. Tinguaro Rodríguez, Daniel Gómez, Javier Yáñez, Javier Montero	
<i>Reducción del efecto vignetting usando medidas de borrosidad</i>	192
Laura Lopez-Fuentes, Sebastià Massanet Massanet, Manuel González-Hidalgo	
<i>Membrane proteins Detection Using Unsupervised Fuzzy Competitive Learning</i>	194
Mohammed Madiati ¹ , Abdelaziz Bouroumi	
<i>Modelado lingüístico de relaciones espacio-temporales de los vectores de movimiento H264/AVC para la detección de eventos en vídeos de tráfico</i>	196
Luis Rodríguez Benitez, Juan Giralt, Juan Moreno Garcia, Luis Jimenez Linares	
<i>Hardware implementation of fuzzy inference systems for real-time video processing applications</i> ...	198
Santiago Sánchez-Solano, María Brox, Elisa Calvo-Gallego, Andrés Gersnoviez, Piedad Brox	

Relaciones multivaluadas

Organizadores: Gaspar Mayor, Tomasa Calvo, Pilar Fuster

<i>Acciones de Grupos Borrosas</i>	200
Dionis Boixader, Jordi Recasens	
<i>Medidas de proximidad ordinal: aplicaciones</i>	202
José Luis García Lapresta, Raquel González del Pozo, David Pérez Román	
<i>From S-implications to similarities</i>	204
Erick González-Caballero, Susana Díaz, Susana Montes, Rafael Espín	
<i>Functions preserving a class of fuzzy binary relations</i>	206
Tomas Sánchez, Gaspar Mayor, Isabel Aguilo, Pilar Fuster, Jaume Suñer	
<i>A study on dissimilarities and negative transitivity</i>	208
Agnieszka Siwiec, Susana Díaz, Bernard De Baets, Susana Montes	
<i>Partial Metric Spaces and Indistinguishability Operators: A Duality Relationship</i>	210
Oscar Valero, Javier Martín, Pilar Fuster	

Soft computing en aprendizaje

Organizadores: José Antonio Sanz, Mikel Galar

<i>Clasificación difusa aplicada a la estimación de la intensidad del esfuerzo físico a partir de rasgos de expresión facial</i>	212
Andrés Cencerrado, Jordi Moreno, Lluís Capdevila	
<i>Análisis exploratorio gráfico borroso con Radviz para la representación de la evolución del aprendizaje</i>	214
Luis Junco, José Otero, Rosario Suárez	
<i>Generalizing the Choquet Integral Using Copulas: An Application to Fuzzy Rule-Based Classification Systems</i>	216
Giancarlo Lucca, José Antonio Sanz, Humberto Bustince, Graçaliz Dimuro, Benjamin Bedregal	
<i>Tackling Brain Computing Interfaz problems using Fuzzy Rule-Based Classification Systems</i>	218
Maria Perez-Zabalza, José Sanz, Aranzazu Jurio, Marisol Gomez, Humberto Bustince	
<i>Graph Construction Using Adaptive Local Hybrid Coding Scheme: Application to Face Recognition</i>	220
Mahdi Tavassoli Kejani, Fadi Dornaika, Alireza Bosaghzadeh, Hussein Mahvash	
<i>Information Dynamics of Machine Learning</i>	222
Francisco Valverde-Albacete, Carmen Peláez-Moreno	
<i>Una aproximación difusa del vecino más cercano para clasificación multi-instancia</i>	224
Pedro Villar, Rosana Montes, Ana María Sánchez, Francisco Herrera	

Soft computing y generación de lenguaje natural

Organizadores: Alberto Bugarín, Nicolás Marín, Daniel Sánchez, Gracián Triviño

<i>Automatic Simplification of Numerical Expressions using Fuzzy Logic</i>	226
Susana Bautista, Raquel Hervas, Pablo Gervas	
<i>Computing protoforms from high-rate sensor streams</i>	228
J. Medina ¹ , M. Espinilla, S. Campana, L. Martinez	
<i>Una aproximación a la generación de expresiones de referencia con propiedades difusas</i>	230
Nicolás Marín, Gustavo Rivas-Gervilla, Daniel Sánchez	
<i>Perspective-based Character Description in Interactive 3D Environments</i>	232
Gonzalo Mendez, Raquel Hervas, Pablo Gervas, Ricardo de la Rosa, Daniel Ruiz	
<i>Using Fuzzy Logic to Model Character Affinities for Story Plot Generation</i>	234
Gonzalo Mendez, Pablo Gervas, Carlos Leon	
<i>Generación de lenguaje natural a partir de secuencias de datos utilizando modelado espacio-temporal</i>	236
Juan Moreno Garcia, Luis Jimenez Linares, Luis Rodriguez Benitez	
<i>A Contribution to the Study of Paraphrasing using Fuzzy Logic and WordNet</i>	238
Martín Pereira-Fariña, Chris Reed	
<i>Fuzzy sets and natural language generation. What to make out of it?</i>	240
Alejandro Ramos Soto, Alberto Bugarin, Senén Barro Ameneiro	

The age of aggregation

Bernard De Baets

*KERMIT, Department of Mathematical Modelling, Statistics and Bioinformatics
Ghent University, Coupure links 653, 9000 Gent, Belgium
Bernard.DeBaets@ugent.be*

The study of aggregation functions is undeniably one of the most important spin-offs of fuzzy set theory. Although the major part of the literature considers the aggregation of numerical values, also qualitative linearly or partially ordered scales are witnessing increased attention. One could say that there is a shift towards the aggregation on structures.

Given the data-generating era we are living in, materialized in the buzzword big data, a dedicated study of aggregation processes has become even more indispensable. This lecture is an appeal to the fuzzy set community to play a key role in this endeavour. This can only be achieved with an open mind. Firstly, although the unifying conditions of aggregation functions (boundary conditions and monotonicity) are extremely basic, we claim that their generality might be questioned. Secondly, the aggregation of structured information, such as rankings, graphs, trees, posets is likely to become a hot topic.

In short, we expect to see an increasing interest in the aggregation of relational data, just as has been the case already in the closely related field of machine learning. We will focus on the aggregation of rankings, point to largely ignored monotonicity issues and downgrade the role of distance metrics in favour of the recently introduced monometrics, which should be appealing to the fuzzy set community.

In this way we might all work together to achieve aggregation 2.0.

El reto del uso de los Sistemas Difusos en Big Data y Ciencia de Datos

Francisco Herrera

Dpto. de Ciencias de la Computación e Inteligencia Artificial. Universidad de Granada

<http://decsai.ugr.es/herrera>

herrera@decsai.ugr.es

La combinación de la gran cantidad de datos disponible en la actualidad y de las herramientas necesarias para su procesamiento conforman lo que se conoce como *big data*. Aunque el término *big data* se refiere al tamaño de los conjuntos de datos a procesar, también implica la necesidad de su análisis en tiempo real cuando recibimos flujos continuos e ingentes de datos, y la posibilidad de manejar una gran variedad de estructuras, datos numéricos, textuales, enlaces. Hay que tratar con datos cuyo volumen, diversidad y complejidad requieren el uso de nuevas arquitecturas, técnicas, algoritmos y análisis para gestionar y extraer el valor y conocimiento oculto en ellos... por ello se habla de las 3 V's (volumen, velocidad y variedad).

Los datos y la capacidad para procesarlos y extraer conocimiento son el "nuevo oro" en la economía digital en la que nos adentramos. Como consecuencia, ha surgido un área denominada ciencia de datos que recoge todos los escenarios en los que los datos tienen un papel estelar con el objetivo de convertirlos en conocimiento. La ciencia de datos engloba a las áreas conocidas como aprendizaje automático, minería de datos, big data,

Los sistemas difusos comenzaron utilizándose para la representación de conocimiento experto, conocimiento expresado lingüísticamente que permitió el desarrollo de los controladores difusos. Posteriormente se emplearon para representar conocimiento extraído de bases de datos y comenzó el auge del aprendizaje automático de sistemas difusos, desarrollándose múltiples propuestas con diferentes herramientas tales como redes neuronales, algoritmos genéticos, clustering, modelos ad-hoc, ... Después asistimos al largo y rico debate sobre el equilibrio entre interpretabilidad y precisión en el modelado difuso, y el papel que la interpretabilidad jugaba en el diseño y posterior uso de los sistemas difusos, debate como ejercicio dialéctico que perdura y que en la mayoría de los casos no termina con conclusiones que tengan una utilidad real a nivel práctico. Abierta queda la discusión sobre el papel que juegan los sistemas difusos frente a otros modelos igualmente interpretables como los árboles de decisión o sistemas basados en reglas en general, así como su papel dentro de la teoría del aprendizaje automático y la ciencia de datos. Todos estos estudios han ido en paralelo al desarrollo de aplicaciones, muy reconocidas en el uso de controladores difusos, actualmente una tecnología más a disposición de los ingenieros. Igualmente podemos encontrar múltiples aplicaciones de sistemas difusos en el modelado de datos.

Hoy en día está aceptado que los sistemas difusos son una herramienta más en el modelado de datos, aunque todavía no se han desarrollado estudios que analicen el papel que juegan:

"cuándo y dónde deben ser utilizados y por qué". No basta hacer referencia a la interpretabilidad, hay que mostrar el porqué de esa interpretabilidad, o por qué apostar por el uso de los modelos de reglas difusos, modelos con solapamiento frente a los modelos tipo divide y vencerás como los árboles de decisión. Se pueden plantear varias preguntas para reflexionar sobre la interpretabilidad: ¿Cuándo es útil la descripción lingüística o las reglas fuzzy? ¿Son útiles los modelos interpretables en problemas complejos y con gran cantidad de datos? ¿Se deben seguir diseñando modelos fuzzy más complejos y no interpretables vía fusión de modelos, para ser competitivos en precisión? Por otra parte, existen otros modelos más allá de los sistemas basados en reglas, como los prototipos difusos representados por técnicas de vecino más cercano, donde los algoritmos "fuzzy kNN" son otra excelente alternativa a los modelos clásicos, o el uso de la teoría de los conjuntos difusos fusionada con otros modelos de aprendizaje (SVM, redes neuronales fuzzy, árboles de decisión fuzzy,). Todos ellos son alternativas estudiadas en la literatura especializada, cuyo uso no está analizado a nivel teórico ni caracterizado a nivel práctico.

En el escenario actual de big data, se nos abren importantes posibilidades para el empleo de los sistemas difusos. El uso del modelo de programación "MapReduce" nos sitúa en un nuevo escenario donde nos encontramos los problemas ya conocidos de "falta de densidad de datos", "pequeños conjuntos aislados", consecuencia del particionamiento de los datos en la etapa "Map". Por otra parte obtenemos modelos parciales que hay que combinar para mostrar modelos globales en un entorno de big data. Estudiar bajo qué condiciones los sistemas difusos pueden ser una herramienta importante para estos problemas es un reto actual en la teoría de sistemas difusos que puede hacer de éstos una herramienta importante en este nuevo escenario.

En este contexto, la conferencia se planteará desde una doble perspectiva: por una parte hacer una breve introducción al uso de los sistemas difusos en big data, mostrando los resultados existentes y los retos abiertos; y por otra discutir su aplicación en la ciencia de datos, analizando el papel que juegan los modelos de aprendizaje basados en conjuntos difusos en la actualidad, y los retos abiertos en los múltiples escenarios de la ciencia de datos. Una pregunta queda en el aire: ¿cuál es el sitio de los modelos de aprendizaje basados en conjuntos difusos y/o sistemas difusos en la ciencia de datos? Es labor de la comunidad de soft computing dar argumentos teóricos y prácticos para que jueguen un papel importante en los nuevos restos de la ciencia de datos.

De las funciones de agregación a las multidistancias: toda una experiencia

Gaspar Mayor

Universitat de les Illes Balears, Palma (Mallorca)

Cuando uno llega a una cierta edad y la vida le ha permitido dedicarse durante mucho tiempo a la investigación, no estoy seguro que sea una buena idea hacer balance de lo que ha significado para los demás y para uno mismo esa dedicación. No sé si es bueno recordar los logros alcanzados, los fracasos acumulados, el tiempo perdido en búsquedas inútiles, la emoción de una idea fructífera, etc., pero la amable invitación de la que he sido objeto, y que agradezco de todo corazón, me libra de la duda, y me aboca a hacer un repaso del trabajo de investigación llevado a cabo desde principios de los años ochenta (del siglo pasado) hasta hoy.

Lo primero que tengo que hacer es recordar a los principales responsables de que durante más de 35 años haya trabajado en esto de la creación científica y, más concretamente, en algunos de los temas que pertenecen al intervalo marcado por el título de esta contribución. Por otra parte, vienen a mi memoria aquellos libros y artículos que me aportaron ideas para empezar a trabajar en mis primeros temas de investigación. Sin duda es importante para un investigador haber tenido la suerte de contar con buenos maestros, y de haber topado con publicaciones sugerentes. Este es mi caso, y por ello agradezco a unos (los maestros) y a otros (los autores de las publicaciones) su decisiva contribución a mi carrera como investigador.

Entre las mayores satisfacciones que he tenido en todos estos años tengo que citar las derivadas de la dirección de tesis doctorales, que, además, han marcado hitos en mi propia actividad investigadora. No importa decir, en esta misma línea, que la constitución del grupo LOBFI (Lógica Borrosa y Fusión de la Información) a finales de los años ochenta ha representado seguramente uno de los mayores logros en los que he podido participar. Ha sido una satisfacción ver crecer a este grupo de investigadores que se ha convertido al cabo de los años en un grupo de referencia tanto en España como fuera de ella.

Para terminar esta introducción, sería obligado referirme a las líneas de investigación en las que he trabajado y que, en definitiva, constituyen la principal razón de ser de mi contribución, pero creo que será mejor que las exponga, junto con el desarrollo de lo anterior, hablando de viva voz.

Dedicado a mis maestros y a mis discípulos.

Vínculos en entornos difusos

Manuel Ojeda-Aciego

*Universidad de Málaga, Depto. Matemática Aplicada
Blv. Louis Pasteur 35, 29071 Málaga, Spain
aciego@uma.es*

El análisis de conceptos formales (FCA) se ha convertido en una línea de investigación muy activa, tanto desde el punto de vista teórico como del práctico. Su gran aplicabilidad es motivo suficiente para estudiar de un modo más profundo sus mecanismos subyacentes, y una forma importante de este conocimiento adicional se obtiene a través de la generalización.

Se han desarrollado diversas variantes de FCA, la mayoría de las cuales se centran en incluir modos de razonamiento alternativos que dan lugar al uso de adjetivos adicionales (difuso, posibilista, rough, etc.). Sin embargo, no se ha publicado mucho en ciertas nociones específicas, tales como los vínculos entre contextos formales.

Una de las motivaciones para introducir la noción de vínculo fue la de proporcionar una herramienta para el estudio de correlaciones entre contextos formales, de alguna manera imitando el comportamiento de las conexiones de Galois entre sus retículos de conceptos correspondientes. En esta charla trataremos con generalizaciones de la noción de vínculo en un entorno L-difuso.

Cosas que he aprendido en mis cuarenta años de lógica borrosa

Enric Trillas

Oviedo

etrillasetrillas@gmail.com

Trataré, únicamente, acerca de dos temas que me parecen relevantes para quienes investiguen, o quieran aplicar, la lógica borrosa. Ambos tienen a ver con una corrección que me gusta hacerle a la conocida afirmación de Zadeh, ‘fuzzy logic is a matter of degree’, para transformarla en ‘fuzzy logic is a matter of degree and design’ remarcando la diferencia con metodologías de tipo más formal ^[1]. Hablaré sobre los conceptos de ‘fuzzy set’ y de enunciado condicional, que son parte de cuanto mis más de cuarenta años dedicados a la lógica borrosa, me han enseñado que presentan sombras que pueden llevar a diseñar mal un sistema. Ninguno de los dos puede darse por cerrado y, como se mostrará, en ambos existen cuestiones aún abiertas que les otorgan un interés que va más allá de una simple descripción.

1. De forma abusiva, se confunde un conjunto borroso de etiqueta lingüística P en un universo del discurso X , con una de sus funciones de pertenencia $\mu_P: X \rightarrow [0, 1]$ cuando, de los conjuntos borrosos, no se sabe bien qué tipo de entidades son, al no conocerse ningún criterio de identidad que los individúe. Por el contrario, al identificarlos con una función de pertenencia y con la identidad funcional punto a punto, resulta que cada predicado P origina en X una multitud de conjuntos borrosos y sin que exista evidencia de ello en el lenguaje natural que está en la base de la llamada lógica borrosa. Es una visión formal que deja lagunas tales que parece razonable repensar aspectos como es, por ejemplo, qué representa una función μ_P , qué nos muestra realmente del conjunto borroso. Se intentará explicarlo partiendo de que P genera una pseudo-entidad (ciertamente nebulosa) llamada el ‘colectivo lingüístico de los P ’^[2], como es, por ejemplo, si $P = \text{joven}$ y $X = \text{habitantes del País Vasco}$, ‘los jóvenes vascos’, o si $P = \text{pequeño}$ en $X = [0, 10]$, ‘los números pequeños entre 0 y 10’; tales colectivos lingüísticos están, por lo menos, bien anclados en el lenguaje, significan algo que sabemos reconocer. Por más que sea fácil contra-argumentar que así no se hace sino cambiar ‘conjunto borroso’ por ‘colectivo’, conviene añadir que si, en un ámbito que no es puramente formal ^[3], algo observable debe tomarse siempre por inicial, lo significativo es no sólo explicar qué representa μ_P más allá de una función ^[2], sino ligarlo con su diseño, con su especificación y, así, poder ver al colectivo bajo otra luz; algo en lo que aún hay algún cabo suelto que se mostrará en la charla y que hace que el tema aún deba considerarse abierto.

2. Nada en la lógica borrosa es universal, sino situacional, dependiente del significado contextual y, con frecuencia, guiado por un propósito como son, por ejemplo, las funciones de pertenencia, las que representan a los conectivos, los modificadores y los cuantificadores ^[1]. También es así con los enunciados condicionales o reglas, ‘Si p , entonces q ’ ($p \rightarrow q$) que, en general, no muestran sino una relación entre el antecedente p y el consecuente q y presentan, para su diseño y por lo menos, dos caras distintas; la de su significado y la forma que este le impone, así como la de su representación acorde con el

uso que se pretenda del mismo. En suma, su forma simbólica y su especificación ^[4]. Por ejemplo y en el caso clásico, una forma ‘material’ que quiera usarse para deducir-adelante, deberá cumplir la desigualdad del Modus Ponens ($p \cdot (p \rightarrow q) \leq q$) y para deducir-hacia-atrás la del Modus Tollens ($q' \cdot (p \rightarrow q) \leq p'$). Sin embargo y en el razonamiento de sentido común aparecen otros usos en los que el consecuente sólo es conjeturable desde el antecedente y la regla, debiéndose cumplir por tanto, $p \cdot (p \rightarrow q) \leq q'$. Un problema relacionado es el de la representación de los condicionales en las formas conjuntivas de Mamdani o Larsen ^[1], imponiendo la rara conmutatividad de la conjunción en el lenguaje. La importancia práctica de los sistemas descritos por reglas lingüísticas imprecisas, muestra el interés del ya viejo problema de los enunciados condicionales que, no obstante, tampoco puede considerarse resuelto por completo. Las más de cuarenta funciones que han sido empleadas para representar las reglas borrosas y pertenecientes a varias formas distintas, son una muestra de ello.

La lógica borrosa mantiene sus puertas abiertas a nuevas, no rutinarias y creativas ideas, especialmente si se la ve, en lugar de otro tipo de lógica deductiva, como un estudio experimental y teórico en los ámbitos reales del lenguaje natural y el razonamiento ordinario no necesariamente de tipo deductivo formal. Ese es, a mi modo de ver, el camino a seguir hacia el ‘Computing with Words’ de Zadeh ^[5] que no hace, en realidad, sino volver a sus primeras ideas ^[6] y que, a la vez, podría permitir desanudar el llamado ‘nudo gordiano’ de la Inteligencia Artificial.

References

- [1] E. Trillas, L. Eciolaza, 2015, Fuzzy Logic, Springer.
- [2] E. Trillas, C. Moraga, S. Termini, A Naïve way of looking at fuzzy sets, forthcoming in Fuzzy Sets and Systems.
- [3] E.H. Mamdani, E. Trillas, 2012, Correspondence between an experimentalist and a theoretician; Combining Experimentation and Theory (Eds. E. Trillas *et al*): 1-18, Springer.
- [4] E. Trillas, I. García-Hontado, 2012, A Reflection on Fuzzy Conditionals; Combining Experimentation and Theory (Eds. E. Trillas *et al*): 391-406.
- [5] L.A. Zadeh, 2012, Computing with Words: Principal Concepts and Ideas; Springer.
- [6] L.A. Zadeh, 1965, Fuzzy Sets; Information and Control, 8: 338-353.

Computing the whole set of adjoint negations

M. Eugenia Cornejo¹, Francesc Esteva², Jesús Medina³, Eloísa Ramírez-Poussa³

¹ *Department of Statistic and O.R., University of Cádiz. Spain*
mariaeugenia.cornejo@uca.es

² *Artificial Intelligence Research Institute (IIIA - CSIC), Bellaterra, Spain*
esteva@iiia.csic.es

³ *Department of Mathematics, University of Cádiz. Spain*
{jesus.medina,eloisa.ramirez}@uca.es

From a theoretical point of view, negation operators have been extensively studied in different works [4, 5]. Two of the environments in which these operators have a remarkable importance are fuzzy logic and fuzzy logic programming. In this abstract, we will remind the notion of adjoint negations which arise as a generalization of residuated negations. Adjoint negations are defined from the implications of an adjoint triple [1] and several properties of them have been recently introduced [2]. We have also introduced an algorithm which computes the whole set of negation operators from two finite posets considering the strong negations only.

Definition 1. Let $(P_1, \leq_1)$ and $(P_2, \leq_2)$ be two posets, $(P_3, \leq_3, \perp_3)$ be a lower bounded poset and $(\&, \swarrow, \searrow)$ an adjoint triple with respect to P_1 , P_2 and P_3 . The mappings $n_n: P_1 \rightarrow P_2$ and $n_s: P_2 \rightarrow P_1$ defined as: $n_n(x) = \perp_3 \searrow x$, $n_s(y) = \perp_3 \swarrow y$, for all $x \in P_1$, $y \in P_2$, are called adjoint negations with respect to P_1 and P_2 . The operators n_s and n_n satisfying that $x = n_s(n_n(x))$ and $y = n_n(n_s(y))$, for all $x \in P_1$ and $y \in P_2$, are called strong adjoint negations.

The following theorem shows that adjoint negations can be uniquely defined from strong adjoint negations by means of a bijection between adjoint negations defined on two posets and strong adjoint negations defined on two complete meet-semilattices [3]. This theorem extends an interesting result about the relationship among weak and strong negations defined on a complete lattice studied in [4].

Henceforth, two posets $(P, \leq_P)$, $(Q, \leq_Q)$ and two complete meet-semilattices $(P', \preceq_{P'})$, $(Q', \preceq_{Q'})$ with maximum elements $\top_{P'}$ and $\top_{Q'}$, respectively, such that $P' \subseteq P$ and $Q' \subseteq Q$ are considered. Hence, they are complete lattices although they are not complete sublattices since, for example, the join in P' is not the join in P .

Moreover, $N_{(P', Q')}(P, Q)$ will denote the set of pairs of adjoint negations (n_s, n_n) with respect to P and Q satisfying that $n_s(P) = Q'$ and $n_n(Q) = P'$, and $SN(P', Q')$ will denote the set of pairs of strong adjoint negations (n'_s, n'_n) with respect to P' and Q' .

Theorem 1. *There exists an one to one correspondence between $N_{(P', Q')}(P, Q)$ and $SN(P', Q')$.*

Corollary 1. *Given a pair of strong adjoint negations (n'_s, n'_n) with respect to P' and Q' , there exists a pair of adjoint negations (n_s, n_n) with respect to P and Q defined as:*

$$\begin{aligned} n_s(p) &= n'_s(z_p) \quad \text{with} \quad z_p = \bigwedge_{P'} \{y \in P' \mid p \leq y\} \\ n_n(q) &= n'_n(z_q) \quad \text{with} \quad z_q = \bigwedge_{Q'} \{x \in Q' \mid q \leq x\} \end{aligned}$$

such that $n_{s|_{P'}} = n'_s$, $n_{n|_{Q'}} = n'_n$, and $n_s(P) = Q'$, $n_n(Q) = P'$.

Now, from these results the following algorithm to compute the whole set of adjoint negations w.r.t. two finite posets is introduced.

Algorithm 1: Computing $N(P, Q)$

input : Two finite posets $(P, \leq_P)$ and $(Q, \leq_Q)$
output: $N(P, Q)$ - Adjoint negations which can be defined on $(P, \leq_P)$ and $(Q, \leq_Q)$

```

1 Meet-Subsemilattices of  $P$  with Maximum ( $MSM(P)$ ) :=  $\emptyset$ ;
2 Meet-Subsemilattices of  $Q$  with Maximum ( $MSM(Q)$ ) :=  $\emptyset$ ;
3  $Dual[MSM(P)] := \emptyset$ ;
4  $Isomorphisms(Dual[MSM(P)], MSM(Q)) := \emptyset$ ;
5 Computing the power set of  $P$ :  $\mathcal{P}(P) = \{P' \mid P' \subseteq P\}$ ;
6 for each  $A \in \mathcal{P}(P)$  do
7   if  $(A, \leq_P)$  is a complete meet-semilattice with maximum element then
8      $\lfloor$  add  $(A, \leq_P)$  to  $MSM(P)$ 
9    $Dual[MSM(P)] := \{(A, \leq_P) \mid (A, \leq_P) \in MSM(P)\}$ ;
10 Computing the power set of  $Q$ :  $\mathcal{P}(Q) = \{Q' \mid Q' \subseteq Q\}$ ;
11 for each  $B \in \mathcal{P}(Q)$  do
12   if  $(B, \leq_Q)$  is a complete meet-semilattice with maximum element then
13      $\lfloor$  add  $(B, \leq_Q)$  to  $MSM(Q)$ 
14 for each  $X \in Dual[MSM(P)]$  and  $Y \in MSM(Q)$  do
15   if  $Isomorphisms(X, Y) \neq \emptyset$  then
16      $\lfloor$  add  $Isomorphisms(X, Y)$  to  $Isomorphisms(Dual[MSM(P)], MSM(Q))$ 
17  $N(P, Q) := Isomorphisms(Dual[MSM(P)], MSM(Q))$ ;
18 return  $N(P, Q)$ 

```

where $Dual(MSM(P))$ will contain the dual poset of each meet-semilattice with maximum in P and $Isomorphisms(X, Y)$ is the set of isomorphisms between X and Y at the end.

References

- [1] M. E. Cornejo, F. Esteva, J. Medina, and E. Ramírez-Poussa. Relating adjoint negations with strong adjoint negations. *7th European Symposium on Computational Intelligence and Mathematics (ESCIM 2015)*, I:66–71, 2015.
- [2] M. E. Cornejo, J. Medina, and E. Ramírez-Poussa. A comparative study of adjoint triples. *Fuzzy Sets and Systems*, 211:1–14, 2013.
- [3] M. E. Cornejo, J. Medina, and E. Ramírez-Poussa. General negations for residuated fuzzy logics. *Lecture Notes in Computer Science*, 8536:13–22, 2014.
- [4] F. Esteva. Negaciones en retículos completos. *Stochastica*, I:49–66, 1975.
- [5] E. Trillas. Sobre negaciones en la teoría de conjuntos difusos. *Stochastica*, III:47–60, 1979.

Computing the minimal solutions of general fuzzy relation equations

J.C. Díaz-Moreno¹, J. Medina¹, E. Turunen²

¹ *Department of Mathematics, University of Cádiz. Spain*
{juancarlos.diaz,jesus.medina}@uca.es

² *Research Unit Computational Logic. Vienna University of Technology. Wien, Austria.*
esko.turunen@tut.fi

Fuzzy relation equations were introduced by E. Sanchez [5] and they have been applied to different frameworks, such as, approximate reasoning, decision making, fuzzy control, etc. The existence of the solutions of these equations have been studied in many papers [1, 2, 4]. Nevertheless, the characterization and computation of the minimal solutions has not been studied so much.

Based on a general framework, in which the minimal solutions of each solvable fuzzy relation equation exist, this paper introduces an efficient procedure in order to obtain the minimal solutions of a general fuzzy relation equation, in which the operators may neither be commutative nor associative. For that, an extra property in the operators is assumed which is not very restrictive since, for instance, it is satisfied by the usual t-norms.

The general residuated operator considered to define the fuzzy relation equation will be denoted as $\odot: L \times L \rightarrow L$ and will be order preserving operators with an operator $\rightarrow: L \times L \rightarrow L$ satisfying the following adjoint property with $\odot$, for each $x, y, z \in L$,

$$x \odot y \preceq z \quad \text{if and only if} \quad y \preceq x \rightarrow z \quad (1)$$

Definition 1. *Given the pair $(\odot, \rightarrow)$, a fuzzy relation equation is the equation:*

$$R \circ X = T, \quad (2)$$

where $R: U \times V \rightarrow L$, $T: U \times W \rightarrow L$ are given finite L -fuzzy relations and $X: V \times W \rightarrow L$ is unknown; and $R \circ X: U \times W \rightarrow L$ is defined, for each $u \in U$, $w \in W$, as

$$(R \circ X)\langle u, w \rangle = \bigvee \{R\langle u, v \rangle \odot X\langle v, w \rangle \mid v \in V\}.$$

It is well known that the fuzzy relation equation (2) has a solution if and only if

$$(R \Rightarrow T)\langle v, w \rangle = \bigwedge \{R\langle u, v \rangle \rightarrow T\langle u, w \rangle \mid u \in U\}$$

is a solution and, in that case, it is the greatest solution, see [5, 6].

The following property is fundamental in order to ensure the computation of the minimal solutions.

Definition 2. Given an operator $\odot: L \times L \rightarrow L$, we will say that it holds the IPNE-condition (making reference to that $\odot$ is Infimum Preserving of arbitrary Non-Empty sets), if it verify

$$a \odot \bigwedge \{b_i \mid i \in \Gamma\} = \bigwedge \{a \odot b_i \mid i \in \Gamma\} \quad (3)$$

for each element $a \in L$ and each non-empty subset $\{b_i \mid i \in \Gamma\} \subseteq L$.

In [3], the following procedure was proposed for a linear lattice $(L, \leq)$. Given the fuzzy relation equation (2), where R, X, T are finite, and $\odot$ satisfies the IPNE-Condition; if it is solvable, then, for each $\langle u, w \rangle \in U \times W$

$$V_{uw} = \{v_s \in V \mid R\langle u, v_s \rangle \odot (R \Rightarrow T)\langle v_s, w \rangle = T\langle u, w \rangle\} \neq \emptyset$$

For each $v_s \in V_{uw}$, let us consider

$$e_s = \bigwedge \{d \in L \mid R\langle u, v_s \rangle \odot d = T\langle u, w \rangle\},$$

and the fuzzy subsets $S_{uws}: V \rightarrow L$ of V , defined by $S_{uws}(v) = \begin{cases} e_s & \text{if } v = v_s \\ 0 & \text{otherwise} \end{cases}$

And, let us consider $S_{uw} = \{S_{uws} \mid v_s \in V_{uw}\}$. From these sets, the following characterization arises.

Theorem 1. A fuzzy relation $X: V \times W \rightarrow L$ is a minimal solution of Equation (2) if and only if, for each $w \in W$, the fuzzy set $X_w: V \rightarrow L$, defined by

$$X_w(v) = X\langle v, w \rangle$$

is a minimal cover of $\{S_{uw} \mid u \in U\}$.

The necessity of considering nonlinear lattices in the applications has increased in the last years. Hence, generalizing the previous characterization is a need. The procedure is more complex, since other level of subsets need to be considered and the conjunction among them is different. The details of this generalization will be presented in a next extension of this paper.

References

- [1] J. C. Díaz-Moreno and J. Medina. Solving systems of fuzzy relation equations by fuzzy property-oriented concepts. *Information Sciences*, 222:405–412, 2013.
- [2] J.C. Díaz-Moreno and J. Medina. Using concept lattice theory to obtain the set of solutions of multi-adjoint relation equations. *Information Sciences*, 266:218–225, 2014.
- [3] J.C. Díaz-Moreno, J. Medina, and E. Turunen. Minimal solutions of general fuzzy relation equations on linear carriers. an algebraic characterization. *Fuzzy Sets and Systems*, 2016. In press.
- [4] I. Perfilieva. Finitary solvability conditions for systems of fuzzy relation equations. *Information Sciences*, 234:29–43, 2013.
- [5] E. Sanchez. Resolution of composite fuzzy relation equations. *Information and Control*, 30(1):38–48, 1976.
- [6] E. Turunen. On generalized fuzzy relation equations: necessary and sufficient conditions for the existence of solutions. *Acta Universitatis Carolinae. Mathematica et Physica*, 28:33–37, 1987.

Algoritmos de Lógica Difusa para el Análisis del Deterioro de Electrodo en Soldadura por Resistencia

Luka Eciolaza, Ander Muniategui, Mikel Ayuso, Marijo Garmendia, Pedro Alvarez

IK4 LORTEK, Ordizia, Gipuzkoa, Spain

{leciolaza, amuniategui, mayuso, mjgarmendia, palvarez}@lortek.es

El uso de las tecnologías digitales a lo largo de la cadena de valor manufacturera juega un factor clave para el impulso del crecimiento en la productividad de los procesos industriales [1]. La fabricación avanzada implica una monitorización de las líneas de producción y el acceso rápido a la información generada. Por lo que el análisis inteligente de datos se está estableciendo en el sector como una herramienta esencial a la hora de extraer conocimiento útil en forma de modelos, y su posterior uso para la toma de decisiones.

Este artículo presenta el trabajo realizado para analizar un proceso de *soldadura por resistencia* (figura 1.a) que representa un paso crítico dentro de una línea de fabricación de alta cadencia. En la soldadura por resistencia, la unión entre diferentes metales se consigue mediante la aplicación combinada de calor y fuerza en el punto de unión. En el proceso analizado, ésta combinación es muy delicada tanto para la degradación de los electrodos, como para la calidad de la pieza producida. Los electrodos se van deteriorando progresivamente y afectan directamente a la calidad de las piezas producidas. Se ha realizado un análisis de datos de proceso basado en la lógica difusa [2], con el objetivo final de avanzar hacia un escenario de fabricación con cero defectos. El análisis se centra en evaluar la calidad de las piezas producidas y el estado de degradación de los electrodos, para lo que se busca: (i) determinar cuantitativamente las relaciones entre la aparición de defectos con diferentes patrones de funcionamiento, (ii) correlacionar las condiciones de fabricación con los atributos de calidad de las uniones obtenidas, (iii) identificar parámetros de proceso que afectan significativamente en la calidad de producto final. Esta evaluación permitirá determinar (a) cuándo deben cambiarse los electrodos y (b) la combinación óptima de parámetros de soldadura para extender la vida útil de los electrodos asegurando asimismo la calidad de las piezas producidas.

Actualmente, el conocimiento del proceso está basado en la experiencia de operarios cualificados. Este trabajo propone realizar un análisis y modelado “*off-line*”, en el que se ejecuta un agrupamiento de atributos de calidad de pieza (parámetros de salida) empleando el método FKM (*fuzzy k-means clustering*) y se obtienen cinco clases. Posteriormente, se utilizan los parámetros de soldadura (8 parámetros de entrada) para clasificar las clases identificadas. Para implementar el modelo de clasificación difuso se ha utilizado inicialmente el método de Isibuchi (FRBCS.W) del paquete *frbs* [3]. Para cada uno de los 8 parámetros de entrada se han definido 5 etiquetas con funciones de pertenencia triangulares. El tipo de *t-norma* y *t-conorma* seleccionados son MIN y MAX respectivamente y se ha utilizado la función de implicación de Zadeh. Los ratios de clasificación medios obtenidos son del 88.5%.

Proyecto realizado en colaboración con la empresa ORKLI, S.Coop.

Los resultados obtenidos (figura 1.c) muestran que se pueden determinar los atributos de calidad de las piezas producidas y la degradación de los electrodos en función de los parámetros de soldadura. El trabajo realizado se usará como base para desarrollar un sistema de inspección de calidad no destructivo en tiempo real que será utilizado para la toma de decisiones. Por último, cabe mencionar que este trabajo también representa un intento de ver la aplicabilidad, en forma de velocidad y facilidad de uso, de la lógica difusa para entornos de fabricación reales.

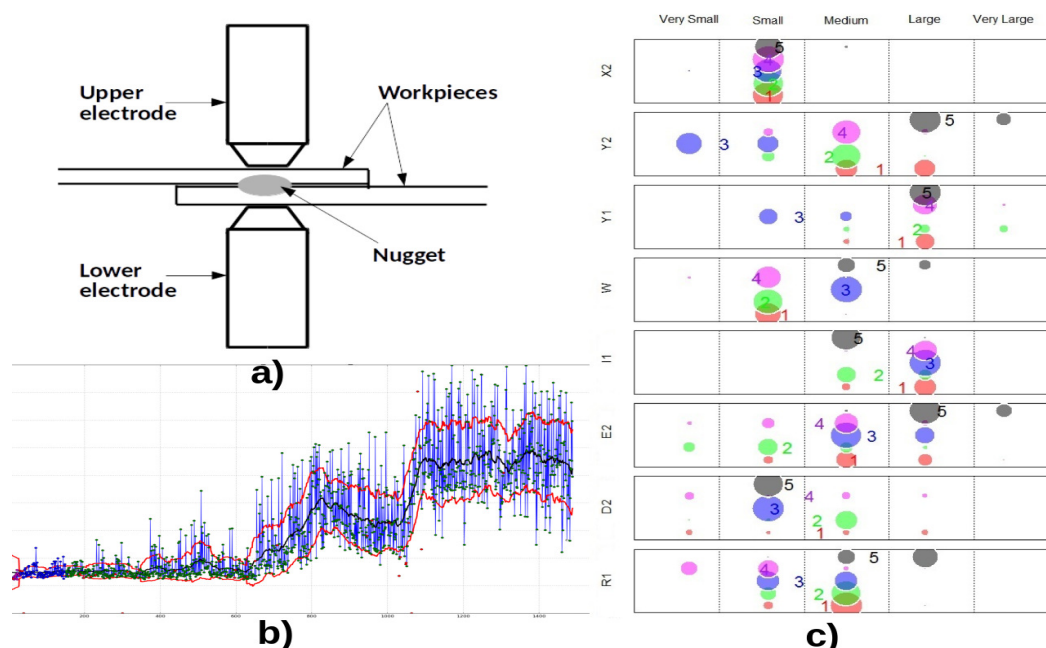


Figura 1: a) Esquema del proceso de unión por resistencia para unir chapas metálicas. b) Evolución de una variable durante el ciclo de vida de un electrodo. En verde los valores de cada pieza producida. Las líneas negra y rojas muestran el valor medio y las desviaciones estandar respectivamente. c) Reglas obtenidas en el clasificador difuso. El área de los círculos es proporcional al número de reglas de cada clase donde aparece la etiqueta correspondiente.

Referencias

- [1] N. Athanassopoulou and H. Thompson. *Development of a strategic research and innovation roadmap for future architectures and services for manufacturing in europe*. Road4FAME project, 2015.
- [2] Enric Trillas and Luka Eciolaza. *Fuzzy Logic: An Introductory Course for Engineering Students*, volume 320. Springer, studies in fuzziness and soft computing edition, 2015.
- [3] Lala Septem Riza, et. al. frbs: Fuzzy rule-based systems for classification and regression in R. *Journal of Statistical Software*, 65(6):1–30, 2015.

Sobre la relación entre operadores de consecuencias y preórdenes borrosos usando conjuntos

J. Elorza¹, S. Montes², C. Bejines³, J. Bragard¹

¹ Dept.de Física y Matemática Aplicada. Facultad de Ciencias. Universidad de Navarra
jelorza@unav.es , jbragard@unav.es

² Dept.de Estadística e Investigación Operativa. Facultad de Ciencias. Universidad de Oviedo
montes@uniovi.es

³ Facultad de Ciencias, Universidad de Málaga
cbejines@uma.es

La relación que existe entre preórdenes y operadores de consecuencias borrosos usando t-normas ha sido ampliamente estudiada ([1],[2]). En esta relación subyace, entre otras nociones, la transitividad borrosa.

El concepto de transitividad se define usando una t-norma. Sin embargo, en algunos casos, una t-norma puede ser demasiado restrictiva. En este contexto, se introduce un concepto de transitividad más general, usando un operador más general, denominado conjuntor, exigiendo sólo la monotonía y la preservación de las condiciones de contorno que extienden la noción de transitividad en conjuntos clásicos ([3]).

En el presente trabajo se generalizan la regla composicional de Zadeh, el concepto de preorden y la noción de coherencia usando conjuntos y se analizan bajo qué condiciones los resultados fundamentales que existen con el uso de t-normas son mantenidos usando el concepto de conjuntor.

Definición 1.-Una operación binaria $*$ en $[0,1]$ se dice que es un conjuntor si

(C1) $x_1 * y_1 \leq x_2 * y_2$ siempre que $x_1 \leq x_2$; $y_1 \leq y_2$

(C2) $0 * 0 = 0$; $0 * 1 = 0$; $1 * 0 = 0$; $1 * 1 = 1$.

Definición 2.-Un conjuntor $*$ en $[0,1]$ verificando $a * 1 = 1 * a = 1$ para todo $a \in [0,1]$, se dirá conjuntor con unidad o con identidad.

En lo que sigue, $*$ denotará un conjuntor cualquiera con unidad (es sencillo probar mediante ejemplos que los resultados presentados en este trabajo no se mantienen, en general, para conjuntos sin unidad).

Definición 3.-Una relación borrosa R definida en un universo X se dice $*$ -preorden borroso si

(P1) $R(x,x) = 1 \forall x \in X$ (reflexividad)

(P2) $R(x,z) \geq R(x,y) * R(y,z) \forall x,y,z \in X$ ($*$ -transitividad).

Recordemos la definición de operador de consecuencias borroso.

Agradecemos la financiación parcial del Gobierno de España bajo los proyectos FIS2011-28820 y TIN2014-59543-P

Definición 4.-Un operador entre subconjuntos borrosos $C : [0, 1]^X \longrightarrow [0, 1]^X$ se dice *Operador de Consecuencias Borroso (FCO, abreviadamente)* si:

(C1) $\mu \subset C(\mu)$ para todo subconjunto borroso $\mu \in [0, 1]^X$ (inclusión)

(C2) $\mu_1 \subset \mu_2 \implies C(\mu_1) \subset C(\mu_2)$ para todos $\mu_1, \mu_2 \in [0, 1]^X$ (monotonía)

(C3) $C(C(\mu)) = C(\mu)$ para todo $\mu \in [0, 1]^X$ (idempotencia).

La siguiente noción de $*$ -coherencia juega un papel esencial en la construcción de preórdenes a partir de operadores de consecuencias borrosos. De hecho, como veremos, es una condición suficiente para que la relación inducida por un operador C mediante $R_C(x, y) = C(\varphi_x)(y)$ resulte un $*$ -preorden.

Definición 5.-Un operador entre subconjuntos borrosos $C : [0, 1]^X \longrightarrow [0, 1]^X$ se dice que es $*$ -coherente si $\mu(a) * C(\varphi_a)(x) \leq C(\mu)(x)$ para todo $\mu \in [0, 1]^X$ para todos $(a, x) \in X \times X$, donde φ_a representa la función característica ordinaria del singleton $\{a\}$.

Para el proceso de generación de FCOs a partir de relaciones borrosas, se tiene:

Teorema 1.-Si R es un $*$ -preorden borroso entonces el operador C_R^* dado por la regla composicional de Zadeh, $C_R^*(\mu) = \mu \circ^* R$ es un FCO, donde $(\mu \circ^* R)(x) = \sup_{w \in X} \{\mu(w) * R(w, x)\}$.

Además, C_R^* es un FCO $*$ -coherente. Para el proceso inverso, se tiene:

Teorema 2.-Si C es un FCO $*$ -coherente, entonces $R_C(x, y) = C(\varphi_x)(y)$ es un $*$ -preorden borroso.

Teorema 3.- $C_R^*(\varphi_x)(y) = R(x, y)$ y $R_{C_R^*}(x, y) = R(x, y)$.

Con los resultados anteriores, tenemos la siguiente cadena de contenidos:

$$\Omega_p^* \subset \tilde{\Omega}^* \subset \Omega_{rt}^* \subset \Omega \subset \Omega' \quad \text{donde,}$$

$$\Omega' = \{C \mid C : [0, 1]^X \longrightarrow [0, 1]^X\},$$

$$\Omega = \{C \mid C \text{ es FCO}\},$$

$$\Omega_{rt}^* = \{C \in \Omega \mid R_C \text{ es } * \text{-preorden}\},$$

$$\tilde{\Omega}^* = \{C \in \Omega \mid C \text{ es } * \text{-coherente}\} \text{ y}$$

$$\Omega_p^* = \{C \in \Omega \mid \exists R \text{ } * \text{-preorden}, C = C_R^*\}.$$

Finalmente, se muestra, mediante amplias familias de operadores, que la cadena de contenidos anterior está formada por contenidos estrictos.

Referencias

- [1] J.L. Castro, E. Trillas. *Tarski's fuzzy consequences*. Proceedings of the International Fuzzy Engineering Symposium'91 (1991) 70–81.
- [2] J. Elorza, P. Burillo. *On the relation between fuzzy preorders and fuzzy consequence operators*. International Journal of Uncertainty, fuzziness and Knowledge-Based Systems **7** (1999) 219–234.
- [3] S. Díaz, S. Montes, B. De Baets. *Transitive decomposition of fuzzy preference relations. the case of nilpotent minimum*. Kybernetika **40** (2004) 71–88.

Sobre los conceptos de dual y anulador de un grupo topológico abeliano fuzzy

Sergio Ardanza-Trevijano¹, María Jesús Chasco¹, Jorge Elorza¹

¹ *Departamento de Física y Matemática Aplicada. Universidad de Navarra.*
sardanza@unav.es, mjchasco@unav.es, jelorza@unav.es

En una serie de artículos clásicos Chang ([1]), Lowen ([4]) y otros desarrollan la teoría de espacios topológicos fuzzy. Ligadas a estas nociones aparecen distintas definiciones de compacidad fuzzy ([5]). Por otra parte Rosenfeld en [6] formula los elementos de una teoría de grupos fuzzy. A partir de los enfoques de Lowen y Rosenfeld, Foster ([2]) desarrolla el concepto de grupo topológico fuzzy.

Recordemos que un grupo fuzzy G en X se define como un conjunto fuzzy en X cumpliendo $G(xy) \geq \min(G(x), G(y))$ y $G(x^{-1}) = G(x)$.

Si X tiene una topología fuzzy ν , podemos considerar en G la topología inducida τ_G . Entonces G es un grupo topológico fuzzy si las aplicaciones $(x, y) \mapsto xy$ y $x \mapsto x^{-1}$ son relativamente fuzzy continuas.

Partimos de un grupo topológico localmente compacto abeliano Hausdorff (X, τ) al que dotamos de la topología fuzzy inducida por τ , $\omega(\tau)$, dada por el conjunto $\mathcal{C}(X, \mathbf{I}_r)$ de las funciones continuas de (X, τ) en $\mathbf{I}_r$, donde $\mathbf{I}_r$ es el intervalo unidad, con la topología $\tau_r = \{(\alpha, \infty) \cap \mathbf{I} : \alpha \in \mathbf{R}\}$ ([4]).

Consideramos el espacio de funciones formado por los homomorfismos fuzzy continuos de $(X, \omega(\tau))$ en $(\mathbf{T}, \omega(\tau_u))$ donde $\mathbf{T}$ denota el círculo unidad $\{z \in \mathbf{C} : |z| = 1\}$, y τ_u es la topología usual en $\mathbf{T}$. Este espacio de funciones tiene estructura de grupo que denotaremos

$$X^\wedge := FChom((X, \omega(\tau)), (\mathbf{T}, \omega(\tau_u)))$$

En $X^\wedge$ se pueden considerar distintas topologías fuzzy. En este trabajo consideramos la topología compacto-abierta fuzzy tal y como la definen Kohli y Prasannan en [3]. Esta topología compacto-abierta se determina del modo siguiente: para cada $x \in X$ definimos $e_x : X^\wedge \mapsto \mathbf{T}$ mediante $e_x(f) = f(x)$. Utilizando la noción apropiada de compacidad, los fuzzy-abiertos subbásicos en la topología compacto-abierta fuzzy que consideramos en $X^\wedge$, son subconjuntos fuzzy de la forma

$$\kappa_\mu = \bigwedge_{x \in \text{sop } \kappa} e_x^{-1}(\mu)$$

donde κ es un fuzzy-compacto en X , μ es un conjunto fuzzy-abierto en $\mathbf{T}$, $\text{sop } \kappa$ indica el soporte de κ y $e_x^{-1}(\mu)(f) = \mu(e_x(f)) = \mu(f(x))$ como es usual.

Los autores agradecen la financiación recibida del Ministerio de Economía y Competitividad de España dada por los proyectos MTM 2013-42486-P y FIS 2011-2820.

Tanto en X como en $X^\wedge$ tienen estructura de grupo y están dotados de topologías fuzzy, por tanto podemos considerar grupos topológicos fuzzy tanto en X como en $X^\wedge$.

Exploramos en este trabajo una noción adecuada de anulador y de grupo dual, que nos permita discernir cuándo dos grupos topológicos fuzzy G en X y H en $X^\wedge$ son anuladores uno del otro o están en dualidad.

Referencias

- [1] C.L. Chang: *Fuzzy Topological Spaces*. Journal of Mathematical Analysis and Applications. **24** (1968) 182–190.
- [2] D.H. Foster: *Fuzzy topological groups*. Journal of Mathematical Analysis and Applications. **67** (1979) 549–564.
- [3] J.K.Kohli, A.R. Prasannan: *Fuzzy topologies on function spaces*. Fuzzy Sets and Systems **116** (2000) 415–420.
- [4] R. Lowen: *Fuzzy Topological Spaces and Fuzzy Compactness*. Journal of Mathematical Analysis and Applications. **56** (1976) 621–633.
- [5] R. Lowen: *A Comparison of Different Compactness Notions in Fuzzy Topological Spaces*. Journal of Mathematical Analysis and Applications. **64** (1978) 446–454.
- [6] A. Rosenfeld: *Fuzzy Groups*. Journal of Mathematical Analysis and Applications. **35** (1971) 512–517.

Utilización de Prototipos Deformables Borrosos en la Generación de Recomendaciones sobre Gestión del Tiempo

Juan García-Aguilar, Francisco P. Romero, Jose A. Olivas, Jesús Serrano-Guerrero

*Departamento de Tecnologías y Sistemas de Información,
Universidad de Castilla La Mancha*

Juan. Garcia8@alu.uclm.es

{FranciscoP.Romero, JoseAngel.Olivas, Jesus.Serrano}@uclm.es

1. Introducción

En la actualidad existe una proliferación de herramientas, dispositivos y aplicaciones que permiten capturar datos de la actividad de cada persona y que ofrecen meramente la posibilidad de su visualización. La integración de la información disponible y su posterior análisis mediante algoritmos inteligentes daría la posibilidad de ofrecer unos servicios personalizados al usuario que le podrían ayudar, por ejemplo, a planificar su semana de trabajo.

En este artículo se propone el análisis de la información que proporcionan herramientas relacionadas con la productividad y el uso del tiempo (*Time Tracking Tools*, Cuantificadores, etc.) con el fin de crear un modelo complejo de comportamiento que permita ofrecer diferentes funcionalidades como una evaluación de la situación del estrés, diferentes recomendaciones con respecto a la situación actual y predicciones sobre lo que va a pasar en un futuro cercano. El principal problema a abordar cuando se intenta analizar estas fuentes de información reside en la imprecisión inherente al registro de los datos. Con el fin de poder manejar datos de estas características la lógica borrosa ofrece mecanismos que permiten introducir conocimiento experto y formalizar la incertidumbre que existe sobre el tema.

Con este fin, se parte del análisis de los factores que pueden influir en la organización del tiempo para posteriormente utilizar técnicas de lógica borrosa como los Prototipos Deformables Borrosos [1] que permiten realizar inferencias sobre las situaciones presentadas. De esta manera, los usuarios disponen de una herramienta que les proporciona no sólo una forma de visualización de la información recogida, sino conocimiento inferido que les ayudará a organizar sus tareas diarias.

2. Planteamiento

En el momento en el que una persona va a planificar el desarrollo de una semana de trabajo maneja un abanico de prototipos determinados por una serie de factores, es decir, debe decidir qué tipo de semana es previsible, existiendo diferentes formas de evaluación según las características de la misma. Siguiendo la teoría de prototipos de Zadeh[2] el concepto de semana con respecto a su productividad engloba un conjunto de prototipos, los cuales representan la buena, baja o media compatibilidad de los ejemplares con el concepto. Desde este punto de vista se puede hablar de una semana “altamente productiva”, “medianamente productiva” o “escasamente productiva”.

Con el fin de detectar las relaciones entre las diferentes semanas y con ello obtener aquellas de escasa, mediana y alta productividad, se realiza un proceso de *clustering* jerárquico sobre los indicadores extraídos para cada semana de los registros de tiempo (tiempos totales desglosados por día de la semana, tiempo total, número de interrupciones,

etc.). Posteriormente, se aplican funciones de resumen sobre las semanas de cada tipo con el fin de calcular los valores las características que van a definir cada uno de los prototipos de evolución del trabajo semanal(días trabajados, tiempo total, interrupciones, distribución diaria del trabajo, distribución del trabajo dentro de los días, etc.). Finalmente, cada uno de estos prototipos se representa mediante el uso de números borrosos para lo cual se realiza un proceso de ejemplificación. Los factores que determinan a priori la progresión del trabajo de una semana se pueden concentrar en los siguientes: desplazamientos previstos, presencia de días festivos, tiempo cambiante/extremo/normal, etc.

A partir de estas definiciones se trata simular la capacidad de interpretación de la situación, es decir, encontrar el modelo de evolución de la productividad más adaptado a las circunstancias reales. Con este fin, en primer lugar se pregunta al usuario el valor real de los factores de que determinan a priori la progresión del trabajo semanal. Su combinación dará unos grados de compatibilidad con los prototipos, aquellos prototipos con un grado de compatibilidad no nulo se modificarán para dar lugar a la caracterización de la semana real mediante una combinación lineal cuyos coeficientes son los grados de compatibilidad anteriormente obtenidos.

La previsión establecida sobre el semana actual sería válida si no existieran factores diarios que influyan en su desarrollo, lo cual ocurre con frecuencia, siendo necesario modificar la previsión realizada si aparecen cambios inesperados en la situación diaria. Los factores diarios que se han extraído de la información disponible son los siguientes: tiempo de trabajo el día anterior, finalización del día anterior, previsión meteorológica especial, etc. Para modificar el desarrollo de la semana se realiza una modificación diaria de acuerdo a determinados valores de algunos de los factores. Por ejemplo, si aparece una previsión de nieve se debe modificar la previsión del día siguiente teniendo en cuenta dificultad en los desplazamientos, etc.

Como prueba de concepto se ha realizado una herramienta Web que incluye gran parte de las funcionalidades planteadas. En primer lugar permite la importación de la información de los registros de tiempo del usuario generados por diferentes aplicaciones (*Toggl*, etc.). Además extrae la información meteorológica (histórico y previsiones) de forma automática mediante técnicas de *WebScraping*. A partir de estos datos se ofrece al usuario diferente información como gráficos de monitorización (diagramas de control, mapas de calor, etc.), recomendaciones lingüísticas sobre el nivel del estrés teórico que se detecta y predicciones sobre la evolución del estrés y del uso del tiempo durante los próximos días.

3. Trabajo Futuro

Como trabajo futuro se plantea desarrollar una aplicación para teléfonos móviles inteligentes que incluya las funcionalidades de la aplicación web actualmente desarrollada. Además se plantea que el propio usuario ofrezca su opinión sobre las recomendaciones recibidas y de esta manera se pueda retroalimentar la aplicación con este conocimiento. También sería interesante estudiar con detalle como el trabajo en diferentes turnos puede afectar al rendimiento y como herramientas de este tipo pueden ayudar a conseguir una mayor productividad.

Referencias

- [1] J. A. Olivás. *Contribución al Estudio Experimental de la Predicción basada en Categorías Deformables Borrosas*. Tesis Doctoral. Universidad de Castilla La Mancha, España, 2000
- [2] L. A. Zadeh. *A note on prototype set theory and fuzzy sets*. Cognition 12, pp. 291- 297, 1982.

Control de suspensión activa basado en lógica difusa

Javier López de la Nieta, Matilde Santos

Facultad de Informática, Universidad Complutense de Madrid, 28030-Madrid
jldorado@gmail.com, msantos@ucm.es

En este trabajo se propone un sistema de suspensión activa con un control basado en la lógica difusa. El sistema de suspensión utiliza el modelo de un cuarto de vehículo, esto es, dos grados de libertad. El controlador difuso diseñado tiene como entrada la diferencia de desplazamiento vertical entre el chasis del vehículo y la rueda, proporcionando como señal de control la fuerza ejercida por el actuador. Tanto el diseño del controlador como la implementación del modelo se han realizado mediante la herramienta MATLAB/Simulink.

En la actualidad las empresas de automoción dedican gran parte de sus recursos al desarrollo e investigación de nuevas tecnologías que permitan mejorar la seguridad en la conducción. Y cada vez con más frecuencia estos nuevos sistemas de seguridad se encuentran basados en técnicas de control inteligente, debido a la complejidad y a la no linealidad que estos sistemas presentan. Por su parte, la suspensión de un vehículo como sistema de seguridad proporciona una mejor maniobrabilidad, mejora el confort en la conducción, favorece el contacto entre el neumático y la carretera, y reduce la transmisión de las perturbaciones del asfalto al conductor.

Todo este conjunto de características suponen un reto muy interesante para los desarrolladores e ingenieros de control de sistemas de suspensión porque son sistemas no-lineales que dependen de parámetros externos como son el perfil de la carretera o el coeficiente de rozamiento del asfalto, entre otros. Además son sistemas que pueden contar con múltiples objetivos de control, como puede ser lograr una buena maniobrabilidad y al mismo tiempo un buen confort para los pasajeros.

Los sistemas de suspensión se pueden clasificar en suspensión pasiva, suspensión activa o suspensión semi-activa. Los sistemas de suspensión activa, objeto de este trabajo, son sistemas regulables mediante un actuador que proporciona la fuerza necesaria para compensar los movimientos de balanceo y cabeceo del vehículo, provocados por factores como la carretera. Son los más complejos pero su principal ventaja es el alto grado de seguridad que proporciona al conductor, ya que permite tener un control independiente de cada rueda.

El modelo seleccionado para simular el comportamiento de la suspensión activa del vehículo es un modelo simplificado que representa un cuarto del vehículo. Este modelo cuenta con una masa suspendida, que es la carga que soportan las suspensiones en reposo, es decir, la masa del chasis. Además cuenta también con la masa no suspendida, y que representa el resto de las masas, como por ejemplo las ruedas, la propia masa de las suspensiones, los neumáticos o los ejes. Los sistemas de suspensión están modelados como dos sistema resorte-amortiguador. Uno de ellos se encuentra entre las dos masas mencionadas anteriormente, simulando el sistema de suspensión real del coche. El otro sistema resorte amortiguación se encuentra entre la masa no suspendida y la carretera, con el fin de simular el contacto del neumático con las perturbaciones del asfalto.

Los sistemas de suspensión activa se caracterizan por contar con un actuador que aplica una determinada fuerza entre el chasis y las ruedas del vehículo, modificando así la

actuación del sistema de suspensión en función de las perturbaciones de la carretera. El objetivo de control en una suspensión activa consiste en mejorar la maniobrabilidad, incrementar el confort y reducir la potencia del controlador. La relación entre todas estas características no es inmediata, y en algunos casos es opuesta. Esto es, el incremento de la maniobrabilidad se logra mediante suspensiones más rígidas, mientras que para que las perturbaciones de la carretera no se transmitan con tanta fuerza al conductor, la suspensión debe ser ligera. Así se mejora el confort. Por lo tanto, el problema de control de una suspensión es multi-objetivo, aunque en este trabajo nos centraremos principalmente en mostrar las ventajas de una suspensión activa utilizando un control difuso en comparación con la suspensión pasiva.

El control de lógica difusa diseñado se compone de una entrada y una salida. La entrada es la derivada de la diferencia de desplazamiento vertical entre ambas masas, y la salida de control es la fuerza que aplica el actuador a las dos masas.

De acuerdo a los resultados experimentales obtenidos en diferentes estudios, se considera que para lograr unos resultados aceptables de confort, la sobreoscilación no debe superar el 5% del valor de la entrada, y el tiempo de asentamiento debe ser inferior a 10 segundos.

En los resultados de simulación se compara la respuesta obtenida por el controlador difuso con el sistema de suspensión pasiva. La sobre-oscilación aunque no se han alcanzado los límites aceptables de confort descritos. A pesar de ello sí que se ha logrado reducir el tiempo de asentamiento siendo éste menor de 10 segundos en todos los casos estudiados.

El siguiente paso es añadir otra variable de control, como la aceleración del desplazamiento vertical de la masa suspendida, para mejorar la respuesta de la suspensión.

A pesar de todo, el principal inconveniente que presentan los sistemas de suspensión activa es el elevado precio que supone su implementación, ya que se requiere de un actuador adicional que proporcione la fuerza de control. Por ello, aunque existen una gran variedad de técnicas de control inteligente que se pueden aplicar a estos sistemas, el objetivo que se deberían marcar las empresas de automoción para asegurarse la rentabilidad de estos sistemas sería la búsqueda de sensores y actuadores de bajo coste. Es por ello que a raíz de esta problemática surgieron los sistemas de suspensión semi-activa.

References

- [1] Ezeta, J. H., Mandow, A., & Cerezo, A. G. *Los sistemas de suspensión activa y semiactiva: una revisión*. Revista Iberoamericana de Automática e Informática Industrial. No. 10(2) (2013) 121–132.
- [2] Sakman, L. E., Guclu, R., & Yagiz, N. *Fuzzy logic control of vehicle suspensions with dry friction nonlinearity*. Sadhana. No. 30(5) (2005) 649–659.
- [3] Salem, M. M. M., & Aly, A. A.. *Fuzzy control of a quarter-car suspension system*. World Academy of Science, Engineering and Technology. No. 53(5) (2009) 258–263.
- [4] Changizi, N., Zare, A., Sheiie, N., & Moghadas, M. *Comparison of PID and fuzzy controller for automobile suspension system*. In: Applied Mechanics and Materials. No. 110 (2012) 671–676.
- [5] Abushaban, M. H., Abuhadrous, I. M., & Sabra, M. B. *A new fuzzy control strategy for active suspensions applied to a half car model*. Journal of Mechatronics. No. 1(2) (2013) 128–134.

Extended Set of Hesitant Fuzzy Linguistic Term Sets

Jordi Montserrat-Adell^{1,2}, Núria Agell², Mónica Sánchez¹, Francisco J. Ruiz¹

¹ *UPC-BarcelonaTech, Barcelona, Spain*

{jordi.montserrat-adell, monica.sanchez, francisco.javier.ruiz}@upc.edu

² *ESADE-Universitat Ramon Llull, Barcelona, Spain*

nuria.agell@esade.edu

In this paper, an extension of the set of Hesitant Fuzzy Linguistic Term Sets (HFLTS) is proposed. Different approaches involving linguistic assessments have been presented in the decision-making literature to deal with the uncertainty connate with human preferences [2, 3]. Furthermore, different levels of precision can be used to give more realistic assessments when uncertainty increases [4]. To model this kind of situations, HFLTS were introduced in [5]. A lattice structure considering the operations intersection and connected union was established in the set of HFLTS in [1].

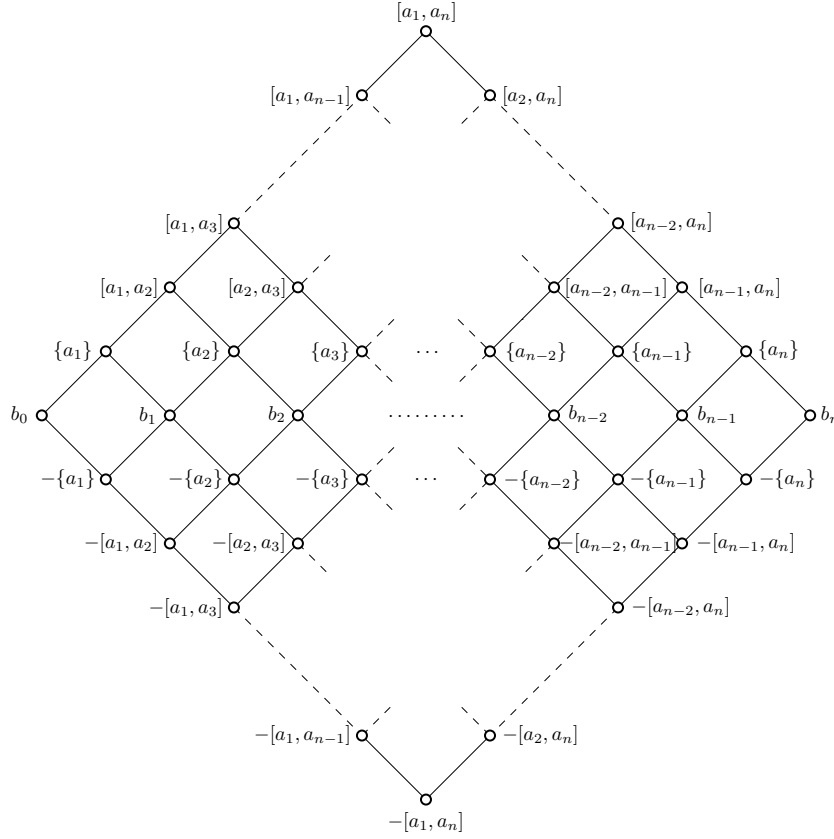
With the aim of distinguish among pairs of HFLTS with empty intersection, a new extension of this set is presented in this paper. To do so, the concept of negative HFLTS is introduced in order to define an *extended intersection*, $\sqcap$, that results either the part in common between two HFLTS or the gap between them in a single operation.

The introduction of this new concept makes us redefine the connected union of HFLTS in contemplation of the negative HFLTS. This new operation is called *extended connected union*, $\sqcup$. Properties satisfied by the extended connected union and the extended intersection are also studied in this paper.

In addition, the extended set of HFLTS is defined as follows: given a partition of a real interval defined by a set of landmarks $\mathcal{B}$, and considering $\mathcal{S}$ the set of ordered linguistic term sets associated to the subintervals of the partition, let $\mathcal{H}_{\mathcal{S}}$ be the set of all the HFLTS by means of $\mathcal{S}$. Then, the extended set of HFLTS is $\overline{\mathcal{H}_{\mathcal{S}}} = \mathcal{H}_{\mathcal{S}} \cup (-\mathcal{H}_{\mathcal{S}}) \cup \mathcal{B}$, where $-\mathcal{H}_{\mathcal{S}}$ is the set of negative HFLTS. Taking into consideration the two previous operations, $(\overline{\mathcal{H}_{\mathcal{S}}}, \sqcup, \sqcap)$ is a distributive lattice.

Figure 1 shows the structure of the lattice $(\overline{\mathcal{H}_{\mathcal{S}}}, \sqcup, \sqcap)$ in the case where $\mathcal{S} = \{a_1, a_2, \dots, a_n\}$ and $\mathcal{B} = \{b_0, b_1, \dots, b_n\}$ with b_i being the landmark between a_i and a_{i+1} if $i \neq 0, n$. The corresponding inclusion relation that can be seen in Figure 1 is also studied.

Finally, a distance between HFLTS is defined as a distance between elements in the lattice $(\overline{\mathcal{H}_{\mathcal{S}}}, \sqcup, \sqcap)$, based on the length of the shortest path from one element to another one within the lattice.

Figure 1: The lattice $(\overline{\mathcal{H}_S}, \sqcup, \sqcap)$.

References

- [1] Agell, N., Sánchez, M., Prats, F., Ruiz, F.J.: *Computing distances between decision makers using hesitant fuzzy linguistic term sets*. Frontiers in Artificial Intelligence and Applications. **277** (2015) 71-79.
- [2] Herrera, F., Herrera-Viedma, E., Martínez, L.: *A fuzzy linguistic methodology to deal with unbalanced linguistic term sets*. IEEE Transactions on Fuzzy Systems. **16** (2) (2008) 354-370.
- [3] Herrera-Viedma, E., Herrera, F., Chiclana, F.: *A consensus model for multiperson decision making with different preference structures*. IEEE Transactions on Systems, Man and Cybernetics, Part A: Systems and Humans. **32** (2002) 394-402.
- [4] Prats, F., Roselló, L., Sánchez, M., Agell, N.: *Using L-fuzzy sets to introduce information theory into qualitative reasoning* Fuzzy Sets and Systems **236** (1) (2014) 73-90.
- [5] Rodríguez, R.M., Martínez, L., Herrera, F.: *Hesitant fuzzy linguistic term sets for decision making*. IEEE Transactions on Fuzzy Systems. **20** (1) (2012) 109-119.

Storing heterogeneous information in fuzzy ontologies using multi-granular fuzzy linguistic methods

Juan Antonio Morente-Molinera¹, Ignacio Javier Pérez Gálvez², Francisco Javier Cabrerizo³, Sergio Alonso⁴, Enrique Herrera-Viedma¹

¹ *Department of Computer Science and Artificial Intelligence, University of Granada*
jamoren@decsai.ugr.es, viedma@decsai.ugr.es

² *Department of Information Science, University of Cádiz*
ignaciojavier.perez@uca.es

³ *Department of Department of Software Engineering and Computer Systems, National Distance Education University (UNED)*
cabrerizo@issi.uned.es

⁴ *Department of Languages and Computer Systems, University of Granada*
zerjioi@ugr.es

Ontologies and fuzzy ontologies are information representation structures that have become quite popular in the recent literature [1, 2, 3]. Thanks to them, it is possible to represent information in an organized way. They also allow users to carry out queries in order to retrieve certain pieces of information. Fuzzy ontologies differ to crisp ontologies in the fact that they are able to represent imprecise information thanks to the use of the fuzzy sets [4] and linguistic modelling [5] computational environments.

There is plenty of literature about building ontologies using external information in order to overcome certain challenges [1, 2, 3]. Nevertheless, there is a lack of papers exposing how to deal with the gathered information in a general way that allow us to create good quality ontologies with any kind of information. It is a fact that, in order to build a proper ontology, it is necessary to look for information in different resources. It is easy to note that the gathered information is heterogeneous and difficult to deal with due to the fact that several representation means are used for its representation.

We propose to solve this issue by using multi-granular fuzzy linguistic methods. Thanks to this kind of methods, it is possible to transform information expressed using different representation means into data expressed by a unique one. This way, it is possible to create a fuzzy ontology using the scheme showed in Figure 1. First, information is uniformed and expressed using the same means using multi-granular fuzzy linguistic information. Afterwards, fuzzy ontology can be built and users can make queries and retrieve the required information.

This paper has been developed with the financing of FEDER funds in TIN2013-40658-P and Andalusian Excellence Project TIC-5991.

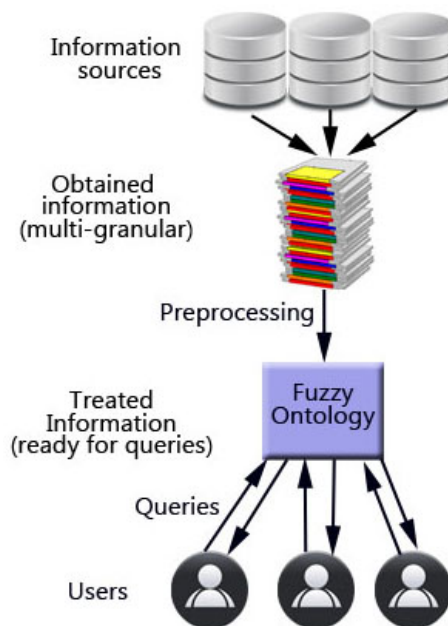


Figure 1: Fuzzy ontology building process using multi-granular fuzzy linguistic modelling.

References

- [1] S. L. Subramanian et al.: *Integration of extracellular RNA profiling data using metadata biomedical ontologies and Linked Data technologies*. Journal of extracellular vesicles. **4** (2015) 1–6.
- [2] T. F. Hayamizu, R. A. Baldock and M. Ringwald: *Mouse anatomy ontologies: enhancements and tools for exploring and integrating biomedical data*. Mammalian Genome. **26**(9-10) (2015) 422–430.
- [3] N. D. Rodríguez, M. P. Cuéllar, J. Lilius and M. D. Calvo-Flores: *A fuzzy ontology for semantic modelling and recognition of human behaviour*. Knowledge-Based Systems. **66** (2014) 46–60.
- [4] L. A. Zadeh: *Fuzzy sets as a basis for a theory of possibility*. Fuzzy sets and systems. **1**(1) (1978) 3–28.
- [5] L. A. Zadeh: *The concept of a linguistic variable and its application to approximate reasoning-I*. Information sciences. **8**(3) (Year) 199–249.
- [6] F. Herrera, E. Herrera-Viedma, L. Martínez: *A fusion approach for managing multi-granularity linguistic term sets in decision making*. Fuzzy Sets and Systems. **114** (2000) 43–58.

Convergence and parameter learning of Fuzzy Cognitive Maps

Gonzalo Nápoles^{1,2}, Elpiniki Papageorgiou^{1,3}, Rafael Bello², Koen Vanhoof¹

¹ Faculty of Business Economics, Hasselt University, Belgium

² Department of Computer Sciences, Central University of Las Villas, Cuba

³ Department of Computer Engineering, Technological Education Institute of Central Greece, Greece

gnapoles@uclv.edu.cu

Fuzzy Cognitive Maps (FCM) are recurrent neural networks where both neurons and causal weights have a precise meaning for the modeled problem [1]. Unlike to feed-forward artificial neural networks, FCM support backward connections that regularly form cycles in the causal graph, which represent the underlying relation among variables. Equation (1) formalizes the updating rule of FCM, where $A_i^{(t+1)}$ is the activation value of the i th neuron at cycle $t + 1$, f_i is the transformation function, w_{ji} is the weight between cause neuron C_j and effect neuron C_i and M is the number of concepts in the system. The updating rule is repeated until (i) the system converges to a fixed-point attractor or (ii) a maximal number of iterations is reached. The former implies that a hidden pattern was discovered [2] whereas the latter suggests that the system outputs are cyclic or completely chaotic.

$$A_i^{(t+1)} = f_i \left(\sum_{j=1}^M w_{ji} A_j^{(t)} + w_{ii} A_i^{(t)}, \lambda_i, h_i \right), i \neq j \quad (1)$$

Recently, Nápoles et al. [3][4] analyzed the convergence of sigmoid FCM used in pattern classification scenarios. From numerical simulations the authors concluded that the proper estimation of the sigmoid parameter λ_i leads to improved convergence performance, without altering the capability for recognizing new patterns. Equation (2) shows the sigmoid function used when transforming the activation value A_i of the i th neuron, where $\lambda_i > 0$ is the sigmoid inclination and $0 \leq h_i \leq 1$ represents the sigmoid offset.

$$f_i(A_i, \lambda_i, h_i) = 1 / (1 + e^{-(A_i - h_i)\lambda_i}) \quad (2)$$

In this paper we extend the above scheme by including the offset parameter h_i into the learning methodology in order to improve the system convergence. It should be highlighted that our methodology attempts enhancing the convergence on previously adjusted FCM, and therefore causal weights cannot be modified. As a second contribution, we studied the FCM convergence and the parameter learning on simulation environments. Equation (3) shows the error function to be minimized during the learning stage, where K is the number of training patterns, T denotes the number of discrete-time steps, M is the number of neurons, ω_t is the relevance of the current discrete-time step, while $A_{ik}^{(T)}$ is the expected response. The goal of this learning model is to reduce, for each training pattern, the squared difference between the response at each discrete-time step and the expected system response.

$$\min \rightarrow \phi(\lambda_1, \dots, \lambda_M, h_1, \dots, h_M) = \sum_{k=1}^K \sum_{i=1}^M \sum_{t=1}^T \frac{2\omega_t (A_{ik}^{(t)} - A_{ik}^{(T)})^2}{KM(T+1)} \quad (3)$$

To solve the real-parameter optimization problem associated to the Equation (3) we used an approach based on Particle Swarm Optimization [5]. For simulations, we generated 990 FCM with different topologies (i.e. M ranges from 2 to 100) and 9900 activation vectors, that is, 10 initial states for each map. The numerical results have shown that learning the sigmoid parameter leads to enhanced convergence performance in the sense that FCM are capable to produce responses which are similar to the desired output. It should be stated that we penalize those solutions leading to non-acceptable numerical outcomes (e.g. the modified FCM does not produce a reasonable response). Figure 1 illustrates the convergence improvements of 20 neurons across 50 discrete-time steps for the same activation vector.

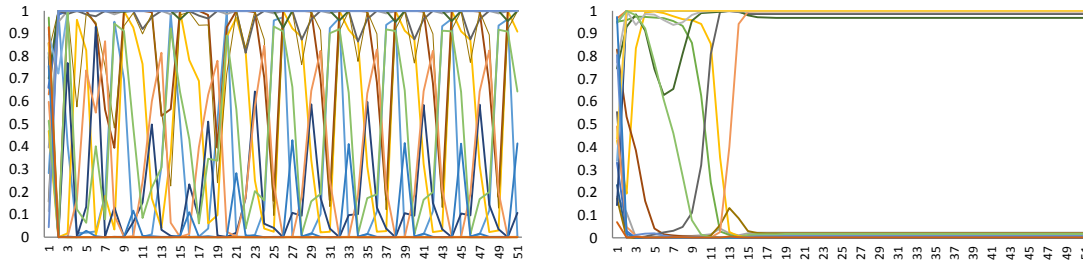


Figure 1: Activation values of 20 neurons for a) using fixed sigmoid parameters b) using the sigmoid parameters learned by the proposed algorithm.

From empirical simulations we can conclude that the estimation of sigmoid parameters may be convenient to improve the convergence on FCM where the causal relations cannot be modified (e.g. FCM directly constructed from experts). However, determining the feasibility threshold requires the intervention of experts since it regularly depends on the problem under investigation. In order to address this serious drawback we may consider a solution as feasible if the ordinal relation among map neurons is preserved. This can be determined by measuring the Hausdorff distance between the ranking of neurons for the original FCM, and the ranking of neurons for the FCM using the estimated parameters.

References

- [1] B. Kosko: *Fuzzy Cognitive Maps*. Int. Journal of Man-Machine Studies. **24** (1986) 65–75.
- [2] A.K. Tsadiras: *Comparing the inference capabilities of binary, trivalent and sigmoid fuzzy cognitive maps*. Information Sciences. **178** (2008) 3880–3894.
- [3] G. Nápoles, R. Bello, K. Vanhoof: *How to improve the convergence on Sigmoid Fuzzy Cognitive Maps?* Intelligent Data Analysis. **18** (2014) 77–88.
- [4] G. Nápoles, E. Papageorgiou, R. Bello, K. Vanhoof: *On the convergence of Sigmoid Fuzzy Cognitive Maps*. Information Sciences. Accepted (2016).
- [5] R. Poli, J. Kennedy, T. Blackwell: *Particle Swarm Optimization—an overview*. Swarm Intelligence. **1** (2007) 37–57.

A Fuzzy System to Assign Reviewers for Conference Papers Evaluation

Jennifer Nguyen¹, Germán Sánchez-Hernández², Núria Agell², Cecilio Angulo¹

¹ *GREC Research Group*

Industrial Engineering School, Universitat Politècnica de Catalunya

Barcelona, Spain

{jennifer.nguyen, cecilio.angulo}@upc.edu

² *GREC Research Group*

ESADE Business School, Ramon Llull University

Barcelona, Spain

{german.sanchez, nuria.agell}@esade.edu

Assigning papers to reviewers is a large and difficult task for conference chairs and scientific committees. Conferences can receive excessively high numbers of paper submissions, sometimes thousands of papers, with only hundreds of reviewers available. Often these assignments must be made within days after the submission deadline creating a huge burden on the conference chairs. Hence, making recommender systems helpful in simplifying this task.

Much research has been conducted regarding this topic in terms of applicability and methodology [1, 2]. In this paper, we focus our attention on aggregating multiple factors to determine reviewer expertise to analyze the match with conference manuscripts. Specifically, we elaborate on variables used to compute reviewer expertise. Although, some prior research focused on keyword terms from publications and supporting information through human intervention [3], we propose a more realistic approach to gathering expertise information using an automated method to acquire reviewer profiles from publicly available information.

Reviewers expertise in different domains is imprecise and multiple factors may contribute to the profile of a reviewer [4]. Therefore, we propose to apply a fuzzy approach to analyze reviewer expertise. Specifically, the methodology considered is based on fuzzy OWA (Order Weighted Average) aggregation functions to summarize information coming from different sources. General constraints for the RAP (Reviewer Assignment Problem) have been incorporated: (i) avoiding conflicts of interest between the reviewer and authors, (ii) each paper must have a minimum number of reviewers, and (iii) each reviewer load cannot exceed a certain number of papers.

To prove the applicability of the proposed methodology an illustrative example has been carried out. The dataset utilized comes from 17th International Conference of the Catalan Association for Artificial Intelligence (CCIA 2014) held in Barcelona, Spain. The data consists of 40 papers and 74 technical committee members. Each paper has a list of keywords associated with it chosen by its authors.

We created a profile for each of the CCIA technical committee members by combing popular global and local researcher websites like Aminer, Portal Recerca (portal for the research in Catalonia), and KTIA (a community of people interested in artificial intelligence). Each

reviewer was identified according to his name and organizational affiliation. All available information was collected for each reviewer and translated to English.

In order to fulfill the conflict of interest constraint, we scraped *dblp*'s website for bibliographic information. We obtained each reviewer's publication titles and co-authors. Next, we collected bibliographic information from *tdx* (a catalog of doctoral theses presented at Catalan universities), for reviewers theses and theses supervised in order to identify supervisors and former students.

We compared the publicly available variables that we obtained to the variables proposed in previous research on state-of-the-art RAP recommenders. Missing variables were added if it could be determined from the available information. Finally, 12 variables were used to determine a reviewer's expertise: PhD Thesis Year, PhD Supervised, Books, Book Chapters, Activity, Research Areas, Publications, Citations, H-index, Sociability, Diversity, and Professional Age.

To evaluate our model, first we compared our results to the assignments made by the conference chairs, then to a random assignment, and finally to a gold standard. This gold standard was determined by several external experts. First, the expert assigned each reviewer to the CCIA call for paper topics. There were 20 topics. The expert then read the keywords and abstract of each paper submitted and assigned relevant CCIA topics to each paper. This gave us a gold standard for all possible combinations of reviewers and papers.

Three main measures of performance are used: precision, recall, and coverage. The results of our study show a promising way of utilizing publicly available information to create expertise profiles and applying fuzzy logic to conference assignment problems.

References

- [1] Wang, Fan, Ning Shi, and Ben Chen: "*A comprehensive survey of the reviewer assignment problem.*" International Journal of Information Technology & Decision Making. **9** (4) (2010) 645–668.
- [2] Karimzadehgan, Maryam, ChengXiang Zhai, and Geneva Belford: "*Multi-aspect expertise matching for review assignment.*" Proceedings of the 17th ACM conference on Information and knowledge management. (2008) 1113–1122.
- [3] Charlin, Laurent, and Richard S. Zemel. "*The Toronto paper matching system: an automated paper-reviewer assignment system.*" International Conference on Machine Learning (ICML). (2013).
- [4] Tayal, Devendra Kumar, et al. "*New method for solving reviewer assignment problem using type-2 fuzzy sets and fuzzy functions.*" Applied intelligence. **40** (1) (2014): 54–73.

El cálculo modal de la borrosidad en el *Pharum Scientiarum* (1659) de Sebastián Izquierdo (1601-1681). Una posible prolongación de la lógica fuzzy de Lofti Zadeh.

Carlos Ortiz de Landázuri

Universidad de Navarra
cortiz@unav.es

Se trata de retrotraer la génesis de la *lógica fuzzy* a un momento previo a la propuesta de Lofti Zadeh (1975). Se analizan a este respecto la posibilidad de desarrollar un cálculo modal formalizado a partir de las propuestas sugeridas por Sebastián Izquierdo (1601-1681) en su *Pharum Scientiarum* (1659) a propósito de la *borrosidad u oscuridad* de los conceptos. A pesar de haber sido una propuesta muy postergada, sin embargo, introdujo unos criterios muy precisos para separar la oscuridad y ambigüedad - o de la correspondiente claridad y distinción -, de una noción respecto de la verdad y falsedad que solo se puede atribuir a las proposiciones. Además, se trató de una propuesta manifiestamente diferente al cálculo de las probabilidades objetivas de Pascal y de la lógica de Port Royal, o al cálculo bayesiano de las probabilidades subjetivas, ambos del siglo XVII. También sería una propuesta diferente de la propuesta cartesiana de una geometría analítica o de un cálculo combinatorio infinitesimal de carácter bivalente, al modo también propuesto posteriormente por Leibniz. En todos estos casos se pretendía formalizar directamente el correspondiente cálculo de proposiciones, sin abordar previamente la formalización de las llamadas nociones oscuras o borrosas, dando un paso a todas luces improcedente. Por su parte, también se diferenciaría de la *lógica de la vaguedad*, en que ya no se aplicaría sólo a nociones en sí mismas relativas, donde siempre se concibe según un más o un menos, sino también a nociones absolutas que, por las razones que sean, no pueden cuantificarse de un modo claro y distinto. En cualquier caso la mayor ventaja de la propuesta de Izquierdo residiría en justificar la noción de borrosidad o oscuridad de la correspondiente noción fuzzy en virtud de una *lógica modal*, concretamente en virtud de la noción de *posibilidad y necesidad* de un concepto. De este modo no habría necesidad de romper con un principio de bivalencia que siempre se refiere a la verdad o falsedad de una determinada proposición, siendo la crítica habitual más frecuente que se ha formulado a la lógica fuzzy posterior a Lofti Zadeh. De todos modos el mayor inconveniente de la propuesta de Izquierdo estribaría en la carencia de una adecuada cuantificación, aunque se trataría de un defecto relativamente fácil de subsanar. A este respecto la comunicación plantea la posibilidad de sustituir la noción de *más o menos verdadero* habitualmente usada en la lógica fuzzy por una noción cuantificada de *más o menos posible*, o *más o menos necesario*, como paso previo a la valoración del efectivo valor de verdad otorgado a los correspondientes puntos de corte de las respectivas proposiciones de la lógica fuzzy. Se trataría de una noción modal de posibilidad y necesidad que se adaptaría mejor al carácter borroso u oscuro que se debe asignar a los conceptos en la lógica fuzzy, aunque secundariamente también deberían poder remitirse a unas determinadas proposiciones fuzzy más o menos verdaderas que fueran efectivamente comprobables en la experiencia. Se trataría en cualquier caso de una herramienta propedéutica complementaria que permitiría ayudar a justificar los complejos procesos utilizados por la lógica fuzzy a la hora de

proponer unos determinados puntos de corte respecto de la correspondiente noción de verdad, sin venir acompañada en muchos casos de una justificación lógica adecuada al respecto.

Referencias

- [1] Cf. Haack, S.; *Deviant Logic, Fuzzy Logic. Beyond the Formalism*, The University of Chicago Press, Chicago, 1974, 1996.
- [2] Cf. AA. VV.; *Fuzzy Sets and Systems. International Journal of Soft Computing and Intelligence. Vol 94, N°1, 1998*, North-Holland, Amsterdam.
- [3] Cf. Zadeh, L.; ‘Fuzzy Sets’. *Information and Control* 8: 338-53, 1965; ‘Fuzzy Logic and Approximate Reasoning’, *Synthese* 30: 407-25, 1975.
- [4] Cf. Keefe, R.; Smith, P. (ed); *Vagueness. A Reader*, MIT, Cambridge (MA), 1996. Smith, B. (ed); *Vagueness, The Monist*, April 1998, 81, 2.
- [5] Izquierdo, S.; “*Pharus Scientiarum*”, Lyon, 1659.
- [6] Fuertes Herreros, J. L.; *La lógica como fundamentación del arte general del saber en Sebastian Izquierdo. Estudio del "Pharus scientiarum" (1659)*, Ed.Uni. Salamanca, 1981.
- [7] Fuertes Herreros, J. L.; *La lógica de Sebastián Izquierdo (1601-1681). Un intento precursor de la lógica moderna en el siglo XVII*, Anuario Filosófico 16 (1983), pp. 219-263.
- [8] Vona, P. Di; *I concetti trascendenti in Sebastian Izquierdo e nella scolastica del seicento*, Loffredo, Napoli, 1993
- [9] Ceñal Lorente, Ramón; El padre Sebastián Izquierdo S. I. (1601-1681) y su “Pharus Scientiarum”, *Revista de Filosofía* 1 (1942), pp. 127-154.
- [10] Abellán, J. L.; *Historia crítica del pensamiento español*, Vol. III, Espasa-Calpe, Madrid, 1981, pp. 253-263.
- [11] Dou Mas de Xerás, A.; Comentario a ‘La combinatoria’ de Sebastián Izquierdo del Padre Ramón Ceñal, S. I.; Real Academia de La Historia, Madrid, 1975.
- [12] Trillas, E.; *Conjuntos borrosos*, Vicens-Vives, Barcelona, 1988.
- [13] Trillas, E.; *Introducción a la lógica borrosa*, Ariel, Barcelona, 1995.
- [14] Grattan-Guinness, I.; *Del cálculo a la teoría de conjuntos, 1630-1910: Una introducción histórica*, Alianza, Madrid, 1984.
- [15] Trillas, E.; *Fuzzy Logic. An Introductory Course for Engineering Students*, Springer, Cham, 2015.

Una propuesta para el modelado del conocimiento del sentido común en la Programación Lógica

Clemente Rubio-Manzano¹, M. Pereira-Fariña²

¹ Universidad del Bío-Bío, Chile.
clrubio@ubiobio.cl

² Centro Singular de Investigación en Tecnoloxías da Información (CiTIUS), Universidade de Santiago de Compostela, España
martin.pereira@usc.es

La Programación Lógica (Prolog) permite el empleo de la lógica de primer orden como lenguaje de programación. Entre sus principales puntos fuertes, destacan su proximidad con el lenguaje natural y su capacidad para la representación de conocimiento y el razonamiento. Sin embargo, al estar basada en la lógica de primer orden, también presenta algunas limitaciones, entre las que destacamos su incapacidad para tratar la vaguedad intrínseca en el lenguaje o el fenómeno de la cuantificación (más allá de los operadores universal y existencial).

En este trabajo, abordamos una propuesta para modelar y representar automáticamente el conocimiento del sentido común en una base de conocimiento de Prolog. En particular, abordaremos el fenómeno de la sinonimia entre términos lingüísticos, un fenómeno típico de este tipo de conocimiento. A pesar de que los seres humanos podemos manejarlo fácilmente, es un problema complejo de abordar desde el punto de vista computacional.

En la bibliografía existen diversas propuestas que abordan este fenómeno [3], introduciendo las denominadas *Ecuaciones de Proximidad* (EP) (1),

$$w_1 \approx w_2 = \alpha \quad (1)$$

donde w_1 y w_2 son símbolos (constantes, función o predicado) de la base de conocimiento del programa Prolog y $\alpha \in [0, 1]$ denota el grado de similitud entre ellos, siendo 1 la máxima semejanza (i.e., igualdad) y 0 la nula semejanza entre ellos. Cuando las EPs cumplen las propiedades de similitud se pueden dar resultados de corrección y completud.

Usualmente, las EPs son definidas *ad hoc* por el diseñador del programa [3]. Desde nuestro punto de vista, esto presenta las siguientes limitaciones: i) los α obtenidos son altamente dependientes de la subjetividad del diseñador; ii) tienen que ser específicamente definidas para cada base de conocimiento; y, iii) cualquier actualización de la misma tiene que ser implementada manualmente.

Este trabajo ha sido realizado en colaboración con el grupo de investigación SOMOS (Software - MOdelling - Science) y parcialmente financiado por la Dirección de Investigación y la Facultad de Ciencias Empresariales de la Universidad del Bío-Bío a través de la ayuda 130415 GI/EF, por el programa Becas Iberoamérica. Jóvenes Profesores e Investigadores del Bando Santander, por el Fondo Europeo de Desarrollo Regional (ERDF/FEDER) a través de las ayudas CN2012/151 y GRC2014/030 de la Consellería de Educación de la Xunta de Galicia y del Ministerio de Economía y Competitividad del a través de la ayuda TIN2014-56633-C3-1-R.

En este trabajo, siguiendo las ideas descritas en [4], proponemos generar los valores de las EPs a partir del conocimiento de sentido común. Asumiendo la hipótesis débil de Shapir-Worf (i.e., los hablantes de una lengua, por hablarla, comparten una parte importante del conocimiento del sentido común), un *thesaurus* nos proporcionaría información de este tipo, dado que recoge como la gente efectivamente usa el lenguaje. En este caso, proponemos usar WordNet y alguna de las métricas basadas en él.

Por otra parte, dada la vaguedad intrínseca del lenguaje y los distintos usos que hacemos de las palabras (una veces son usadas en un sentido estricto mientras que otras se usan en sentido más laxo o metafórico), proponemos modelar los valores α de las EPs mediante relaciones (o conjuntos) borrosas intervalo-valoradas (IVFS) [1].

Los IVFSs, con respecto a los conjuntos borrosos típicos, tienen un mejor comportamiento en el análisis de un entorno con imprecisión. El conocimiento del sentido común es un claro ejemplo y los conjuntos borrosos típicos presentan algunas limitaciones [1]: i) no todos los individuos definen la misma función de pertenencia para una misma palabra; ii) no proporcionan una buena aproximación a la representación del significado de una palabra; y, iii) las palabras pueden significar diferentes cosas para diferentes personas, lo que genera un intervalo de incertidumbre. La aplicación de los IVFSs en el marco de Prolog no sólo nos permitiría aumentar los recursos expresivos y de inferencia del lenguaje, sino también modelar las diferentes creencias o percepciones compartidas por los hablantes de una comunidad (gracias a WordNet) a la hora de realizar un proceso de representación de conocimiento.

En este trabajo, usamos el sistema Bousi Prolog como entorno de programación, al que incorporamos los IVFSs junto con WordNet y métricas de similitud. A continuación, ilustramos mediante un sencillo ejemplo un programa y las EPs modeladas con esta propuesta, usando las métricas de Rada et al. y Wu and Palmer [2]:

```
practises~loves=[0.3,0.5].      loves(mary,mountaineering)
practises~likes=[0.3,0.5].     likes(john,football).
practises~plays=[0.3,0.5].     plays(peter,basketball).
sport~basketball=[0.2,0.8].
sport~football=[0.3,0.9].      healthy(X):- practises(X,sport).
sport~mountaineering=[0.1,0.5].
```

Referencias

- [1] H. Bustince: Interval-valued Fuzzy Sets in Soft Computing, International Journal of Computational Intelligence Systems, Vol.3, No. 2 pp. 215-222 (2010).
- [2] Hadj Taieb, M. A. and Ben Aouicha, M. and Ben Hamadou, A.: Ontology-based approach for measuring semantic similarity, Engineering Applications of Artificial Intelligence, 36, 238-261 (2014).
- [3] Rubio-Manzano, C., Julián-Iranzo, P: A Fuzzy linguistic prolog and its applications, Journal of Intelligent and Fuzzy Systems, 26(3), 1503-1516 (2014).
- [4] Rubio-Manzano, C., Julian-Iranzo, P: Reasoning with words-A first approximation. In FUZZ-IEEE, 569-574 (2014).

Uniform Interpolants in Nilpotent Minimum Logic

Diego Valota

Institut d'Investigació en Intel·ligència Artificial (IIIA), Consejo Superior de Investigaciones Científicas (CSIC) - Campus de la Universitat Autònoma de Barcelona s/n 08193 Bellaterra
diego@iiia.csic.es

Since Craig's seminal paper [5], interpolation properties have proven to be useful tools to understand what can be expressed or derived in a logical system. Besides theoretical interests, algorithms to compute *interpolants* are essential to applications such as formal verification or theorem proving (see [9, 10]). To cope with different situations, several forms of interpolation properties have been introduced in literature (deductive interpolation, uniform interpolation etc.). Although all these forms are equivalent in Boolean logic, the situation becomes much more complicated when moving to nonclassical logics.

The aim of this talk is to concretely show how to effectively compute *uniform* interpolants in Nilpotent Minimum logic.

Nilpotent Minimum (NM) logic is a well-studied case in the setting of *many-valued triangular norms based logic* (t-norms logics for short). In [6], Esteva and Godo have introduced the *Monoidal T-norm based Logic* (MTL), that is precisely the logic of all left-continuous t-norms and their residua [8]. NM logic lies in the hierarchy of schematic extensions of MTL and it is the logic of the *Nilpotent minimum t-norm* [7]. Adding *idempotency* axiom to MTL logic *Gödel-Dummett logic* arises. The algebraic semantic of MTL is given by the variety of *MTL algebras*. As Gödel-Dummett algebras are exactly the prelinear Heyting algebras, NM algebras are the prelinear Nelson algebras [4]. Hence, NM logic is to *Nelson logic* (constructive logic with strong negation) as Gödel-Dummett logic is to Intuitionistic logic.

In this talk, we exploit the concrete representation of free finitely generated Nilpotent Minimum algebras given by Aguzzoli and Gerla in [1] to give a constructive proof of uniform interpolation property of Nilpotent Minimum logic. We show how to obtain uniform interpolants in NM logic using the normal forms construction given in [1]. This technique has been applied to obtain *deductive interpolants* in [2] for Gödel-Dummett logic and then generalized in [3] for *Revised Drastic Product* logic (another extension of MTL).

References

- [1] S. Aguzzoli and B. Gerla. Normal Forms and Free Algebras for Some Extensions of MTL. *Fuzzy Sets and Systems*, 159(10):1131–1152, 2008.
- [2] S. Bova. Combinatorics of Interpolation in Gödel Logic. In *4th Conference on Topology, Algebra and Categories in Logic (TACL)*. 2009.
- [3] S. Bova and D. Valota. Finitely Generated RDP-Algebras: Spectral Duality, Finite Coproducts and Logical Properties. *Journal of Logic and Computation*, 2011.

The author acknowledges support from a Marie Curie INdAM-COFUND Outgoing Fellowship.

- [4] M. Busaniche and R. Cignoli. Constructive logic with strong negation as a substructural logic. *Journal of Logic and Computation*, 20(4):761–793, 2010.
- [5] W. Craig. Linear reasoning. A new form of the Herbrand-Gentzen theorem. *Journal of Symbolic Logic*, 22:250–268, 1957.
- [6] F. Esteva and L. Godo. Monoidal t-Norm Based Logic: Towards a Logic for Left-Continuous t-Norms. *Fuzzy Sets and Systems*, 124(3):271–288, 2001.
- [7] J. Fodor. Nilpotent Minimum and Related Connectives for Fuzzy Logic. In *Proceedings of FUZZ-IEEE'95*, pages 2077–2082. 1995.
- [8] S. Jenei and F. Montagna. A Proof of Standard Completeness for Esteva and Godo's Logic MTL. *Studia Logica*, 70(2):183–192, 2002.
- [9] Kenneth L. McMillan. Interpolation and sat-based model checking. In *Computer Aided Verification, 15th International Conference, CAV 2003, Boulder, CO, USA, July 8-12, 2003, Proceedings*, pages 1–13, 2003.
- [10] Kenneth L. McMillan. An interpolating theorem prover. *Theor. Comput. Sci.*, 345(1):101–121, 2005.

CONDICIONALIDAD Y CONTINUA LINGÜÍSTICOS

Enric Trillas , Gaspar Mayor

etrillasetrillas@gmail.com , gmayor@uib.es

Este artículo sólo debe ser considerado como una exploración, previa y abierta, acerca del razonamiento ordinario, de cada día o de sentido común y, en ningún caso pretende ser de tipo conclusivo o dogmático; no es un trabajo inscribible en lo que se entiende por lógica. Sólo contiene algunas reflexiones iniciales y soportadas empíricamente en un simbolismo matemático muy elemental. Supuesto que lo presentado fuese continuado por vías formales y experimentos controlados se podría, eventualmente, llegar a una teoría científica del razonamiento ordinario y, en particular, de la deducción, la abducción, la especulación y, en particular, la de tipo creativo.

CONDICIONALIDAD

De forma ingenua, representaremos simbólicamente por $p < q$ al enunciado lingüístico 'Si p , entonces q ', simplemente aceptándolo como una construcción mental que se reconoce empíricamente y que está en la base del razonamiento. Diremos que la relación ' $<$ ' entre enunciados es la 'relación empírica de inferencia natural', de la que y de partida sólo se aceptará que es reflexiva, es decir, que para todo enunciado p , es $p < p$. Con alguna frecuencia, existen ternas p, q, r para las cuales la relación de inferencia natural no es transitiva ($p < q$ y $q < r$: $p < r$) sino que, como en los ejemplos de las sensaciones y los sinónimos, es $p < q$ y $q < r$: $p \not< r$. También existen ternas para las cuales se verifica la propiedad $p \cdot q < r$, representando por $(\cdot)$ a la conjunción; se dirá entonces que ' $<$ ' verifica la 'transitividad restringida respecto de la conjunción'. Nótese que en tales casos, de ser $p < q$ y $q < r$, no se sigue necesariamente que deba ser $p < r$, a menos que, por ejemplo, $p \cdot q$ sea idéntico a p .

Debe observarse que nuestro planteamiento es previo a cualquier lógica en la cual pretenda formalizarse el razonamiento bajo leyes prefijadas. No coincide, por ejemplo, con el cálculo lógico clásico en el cual 'Si p , entonces q ' ($p \rightarrow q$) se considera con el mismo significado que $\text{no-}p$ o q ($p' + q$). El cálculo lógico clásico tiene muchas leyes (las del álgebra de Boole) que no pueden suponerse en el lenguaje natural como son, por ejemplo, la conmutatividad de la conjunción y la del reparto perfecto ($p = p \cdot q + p \cdot q'$, para todo par de enunciados p y q).

CONTINUA LINGÜÍSTICOS

Un par $(X, <)$ es un continuo lingüístico empírico - idea introducida por Henri Poincaré para distinguir el continuo matemático del físico o empírico - cuando existen ternas (p, q, r) tales que $p < q$, $q < r$, pero $p \not< r$, es decir, ternas no transitivas. De hecho, dados p y r , se dice que r es inferencialmente aislado de p cuando cualquier terna (p, q, r) , no es $<$ -transitiva. Por ello, la densidad relativa de elementos enlazados inferencialmente con p , tomado como punto de partida de la inferencia, juega un cierto papel en la posibilidad de que r no esté aislado de p . Saber que la situación inferencial es un continuo lingüístico, obliga a ir con mucho tiento para intentar alcanzar inferencialmente a r partiendo de p ; encontrar cadenas

En memoria del profesor Lluís A. Santaló (Girona, 1911- Buenos Aires, 2001).

de elementos $p_1, \dots, p_n$, tales que $p < p_1, p_1 < p_2, p_2 < p_3, \dots, p_{n-1} < p_n, p_n < r$ que, de ser ' $<$ ' transitiva, lleva a $p < r$ y que, cuando ' $<$ ' puede identificarse con la deducción formal, no es sino una prueba formal de r bajo la hipótesis p . Análogamente, puede llegarse a $p < r'$, con lo que r es una refutación de p ; por el contrario, la existencia de continuos lingüísticos puede permitir que sea $p \not< r'$, con lo cual r es una conjetura de p . La existencia de roturas o saltos, puede hacer que en lugar de seguirse r de p , sólo pueda conjeturarse en el sentido de no ser contradictorio con p . Los r tales que $r < p$, se llaman hipótesis o explicaciones empíricas de p , en tanto que aquellos r inferencialmente aislados de p , reciben el nombre de especulaciones de p y, obviamente, no pueden alcanzarse por cadenas inferenciales que partan de p . Dado p , no hay más que enunciados que se deducen de p , que explican p , que refutan p , o que son especulaciones a partir de p .

En el caso de que $<$ sólo sea transitiva restringida, es decir, que para toda terna (p, q, r) se verifique $p \cdot q < r$, lo que lleva a $p \cdot q < p$ y $p \cdot q < q$ a partir, respectivamente, de las ternas (p, q, p) y (p, q, q) , el problema se complica y está abierto, al tener que contemplar los posibles recorridos de p hasta r , o hasta r' , por la existencia de elementos entre ambos, por ejemplo, de un q tal que $p < q$ y $q < r$, verificando $p \cdot q < r$. Ello abre un posible camino a la búsqueda heurística de especulaciones de p , a través de enunciados q con los cuales, aun sin ser $p < r$, sea $p \cdot q < r$; es una posibilidad que al igual que casi todo lo aquí esbozado resta por explorar.

CONCLUSIÓN

El gran prodigio del lenguaje natural reside en la alta conectividad de las palabras reflejada en el razonamiento y que, gracias a la capacidad de abstracción y a cuanto los sentidos y el cerebro permiten coordinar lleva, finalmente, a conocer el mundo inmediato a través de conceptos. La representación simbólica de los conceptos permite la construcción de modelos matemáticos, con los cuales y además del razonamiento deductivo empírico, también el deductivo formal puede aplicarse a organizar el conocimiento y tanto intentar prever su comportamiento, como su futuro.

Como se insinúa en esta contribución, la básica operación racional de inferencia requiere para poder, por lo menos, deducir informalmente, la falta de continua lingüísticos entre antecedente y consecuente; el consecuente no puede estar aislado del antecedente, no puede haber roturas entre ambos sino puentes, o pruebas, o cadenas de enunciados ligándolos.

References

- [1] H. Poincaré : *La Science et l'Hypothèse*. ; Eds. La Bohème, Paris. (1992).
- [2] E. Trillas: *Glimpsing at Guessing*. Fuzzy Sets and Systems, 281: 32-43. (2015)
- [3] E. Trillas, Ll. Valverde: *An Inquiry into Indistinguishability Operators*. Aspects of Vagueness (Eds., H.J. Skala, et altri), Reidel, Dordrecht. (1984).
- [4] E. Trillas, R. Seising, *On Meaning and Measuring: A Philosophical and Historical View*. Agora, 34/2: 1-23.(2015)
- [5] D. Hofstadter, E. Sander: *Surfaces and Essences: Analogy as the fuel and fire of thinking*. Basic Books, New York. (2013).

Modelos de conjeturas y conexiones de Galois

Adolfo R. de Soto

Escuela de Ingenierías Industrial e Informática
Universidad de León

adolfo.rdesoto@unileon.es

Sin lugar a dudas la capacidad de abstracción es una de las características más sobresalientes del razonamiento humano. Ignorar o eliminar detalles no relevantes para focalizar la atención en lo realmente significativo de un problema bajo estudio se convierte en una herramienta fundamental a la hora de abordar cuestiones complejas. La complejidad se ve reducida si eliminamos interacciones no relevantes. En Inteligencia Artificial ha sido donde más esfuerzos se han realizado para encontrar modelos generales que dieran cuenta del fenómeno [4]. Sin embargo la tarea ha fracasado en cierto grado cuando se han buscado soluciones independientes del dominio y por ello los esfuerzos se han dirigido más hacia soluciones ad-hoc [6].

Uno de los mecanismos de modelización de la abstracción ha sido la utilización de conexiones de Galois [3, 2]. Recientemente su uso ha crecido en popularidad debido a que diversas teorías como Rough Sets o Formal Concept Analysis[11] son ejemplos de mecanismos de aproximación modelizados mediante conexiones de Galois. Los conjuntos Fuzzy son una herramienta para realizar abstracciones de la misma forma que el nombre común “árbol” abstrae todas las diferencias entre los diferentes árboles, dando lugar a un concepto. Cuando anulamos detalles, provocamos que elementos distintos sean considerados iguales, de esta forma el concepto de “indistinguibilidad” surge de forma natural en este contexto. Es conocido que un sistema de reglas fuzzy puede interpretarse como la aplicación de la regla composicional de inferencia a una indistinguibilidad[5]. El operador de la regla composicional de inferencia $\phi_E(\mu)(x) = \sup_{y \in X} T(E(x, y), \mu(y))$ definido a partir de una indistinguibilidad E es el coadjunto de una conexión de Galois cuyo adjunto es $\psi_E(\mu)(x) = \inf_{y \in X} I_T(E(x, y), \mu(y))$, definido a partir de la implicación residuada I_T . De esta forma un sistema de reglas fuzzy puede verse como una aproximación definida por una conexión de Galois.

Los modelos de Consecuencias, Hipótesis y Conjeturas [8, 9, 10] buscan dar cuenta de los procesos de razonamiento de sentido común en un marco formal basado principalmente en la regla básica de todo razonamiento, no contener contradicción. Es interesante hacer notar que la negación es una conexión de Galois en cualquier retículo pseudocomplementado[1]. En [7] se definen las BFA (Basic Fuzzy Algebra) como estructuras básicas que permiten analizar modelos CHC en el marco de los conjuntos fuzzy. En este trabajo se explora la variación de los conjuntos de Consecuencias, Hipótesis y Conjeturas cuando se define una conexión de Galois en un par de BFAs, representando la BFA imagen una abstracción de la BFA origen.

El autor reconoce el apoyo recibido por el Ministerio de Economía y Competitividad y los Fondos Europeos de Desarrollo Regional (ERDF/FEDER) bajo la ayuda TIN2014- 56633-C3-1-R.

Referencias

- [1] Garrett Birkhoff. Lattice theory. In *Colloquium Publications*, volume 25. Amer. Math. Soc., 3. edition, 1967.
- [2] P. Cousot and R. Cousot. Abstract interpretation: a unified lattice model for static analysis of programs by construction or approximation of fixpoints. In *4th ACM Symposium on Principles of Programming Languages*, pages 269–282, Jun 1977.
- [3] M Ern , J Koslowski, A Melton, and GE Strecker. A primer on galois connections. *Annals of the New York Academy of Sciences*, 704(Papers on General Topology and Applications):103–125, 1993.
- [4] Fausto Giunchiglia and Toby Walsh. A theory of abstraction. *Artificial Intelligence*, 57(2-3):323–389, Jun 1992.
- [5] Rudolf Kruse, Joan E Gebhardt, and F Klowon. *Foundations of fuzzy systems*. John Wiley & Sons, Inc., 1994.
- [6] Lorenza Saitta and Jean-Daniel Zucker. *Abstraction in artificial intelligence and complex systems*. Springer, 2013.
- [7] Enric Trillas. A model for “crisp reasoning” with fuzzy sets. *International Journal of Intelligent Systems*, 27(10):859–872, 2012.
- [8] Enric Trillas, Susana Cubillo, and Elena Casti neira. On conjectures in orthocomplemented lattices. *Artificial Intelligence*, 117(2):255 – 275, 2000.
- [9] Enric Trillas, Itziar Garc a-Honrado, and Ana Pradera. Consequences and conjectures in preordered sets. *Information Sciences*, 180(19):3573–3588, 2010.
- [10] Enric Trillas, Ana Pradera, and Angel  lvarez. On the reducibility of hypotheses and consequences. *Information Sciences*, 179(23):3957–3963, 2009.
- [11] Y.Y. Yao. A comparative study of formal concept analysis and rough set theory in data analysis. In *Rough sets and current trends in computing*, pages 59–68. Springer, Sep 2004.

Razonamiento difuso mediante modelos CHC y representación por niveles

Daniel Sánchez

Dept. Ciencias de la Computación e I.A., Universidad de Granada
 daniel@decsai.ugr.es

Los modelos CHC (*Conjeturas, Hipótesis y Consecuencias*) formalizan, en base a operadores de consecuencias, tipos de razonamiento ordinario distintos al puramente deductivo, como el abductivo y el especulativo que, junto al razonamiento por analogía (como caso particular, el razonamiento basado en casos) constituyen la mayor parte del razonamiento que realizamos [1]. En los modelos CHC, las sentencias lógicas y proposiciones propias del razonamiento humano pueden representarse por medio de las denominadas *Álgebras Flexibles Básicas* (Basic Flexible Algebras, BFAs), de las cuales retículos con negación, por ejemplo ortoretículos y álgebras de De Morgan, así como álgebras estándares de conjuntos difusos, son casos particulares.

Dado un conjunto crisp y no vacío L de sentencias, un cuerpo de conocimiento (conjunto de sentencias libre de incompatibilidades, donde el concepto de incompatibilidad puede definirse de distintas formas) $P \subset L$, y un operador de consecuencias de Tarski C [3], e independientemente del tipo de álgebra empleado, los modelos CHC generan una partición crisp de L en cuatro subconjuntos:

- *Consecuencias* de P en L , $C(P)$, expresiones que se derivan como consecuencias lógicas del cuerpo de conocimiento de manera necesaria y segura.
- *Hipótesis* de P en L , $Hyp(P)$, expresiones contingentes que son “explicaciones” tentativas del cuerpo de conocimiento, en el sentido de que todas las expresiones de éste último son consecuencias lógicas de la hipótesis.
- *Especulaciones* de P en L , $Sp(P)$, no tienen relación en términos de consecuencia lógica con el cuerpo de conocimiento, siendo su única relación que no son incompatibles con el mismo. Son elucubraciones contingentes con respecto al cuerpo de conocimiento.
- *Refutaciones* de P en L , $Ref(P)$, se caracterizan por ser incompatibles con el cuerpo de conocimiento, es decir, la conjunción de las expresiones que forman el cuerpo de conocimiento y una refutación respecto al mismo dan lugar a una *contradicción*.

El conjunto de sentencias de L que no es incompatible con P se denomina *conjeturas* de P en L y se nota $Conj(P)$, siendo $Conj(P) = C(P) \cup Hyp(P) \cup Sp(P)$ y $Conj(P) \cap Ref(P) = \emptyset$.

Este trabajo ha sido desarrollado dentro del proyecto TIN2014-58227-P, financiado por el Ministerio de Economía y Competitividad y por el Fondo Europeo de Desarrollo Regional (FEDER).

En este trabajo extendemos el marco de los modelos CHC a la *representación por niveles* de la borrosidad [2]. Una representación por niveles (RL, del inglés *representation by levels*) sobre X no es un subconjunto difuso de X en general, sino que se caracteriza por una función $\rho : (0, 1] \rightarrow \{0, 1\}^X$. Entre las características de esta representación podemos destacar:

- Contrariamente al conjunto de alfa-cortes de un conjunto difuso, $\alpha > \beta$ no implica $\rho(\alpha) \subseteq \rho(\beta)$.
- Las operaciones crisp en $\{0, 1\}^X$ se extienden al caso de RLs operando en cada nivel independientemente. En el caso de las operaciones usuales sobre conjuntos de unión, intersección y diferencia, el conjunto de RLs sobre X forma un álgebra de Boole con estas operaciones [2], contrariamente a las álgebras estándar de conjuntos difusos.
- Los conjuntos difusos pueden interpretarse como medidas de pertenencia de elementos de X a una RL [2]. En general, dado un conjunto difuso F , existe más de una RL cuyas medidas de pertenencia de los distintos elementos de X es igual a F , entre las que destaca la RL dada por $\rho_F(\alpha) = F_\alpha$, siendo F_α el α -corte de F .

Sea L un conjunto de expresiones y sea ρ_L una RL sobre L que representa el grado de cumplimiento de cada sentencia de L bajo una cierta interpretación, $\rho_L : (0, 1] \rightarrow \{0, 1\}^L$. Sea $P \subset L$ y $\rho_P = P \cap \rho_L$. Diremos que P es un cuerpo de conocimiento si $\rho_P(\alpha)$ es un cuerpo de conocimiento $\forall \alpha \in (0, 1]$. Utilizando modelos CHC tenemos que, para cada $\alpha \in (0, 1]$, podemos obtener una partición crisp de L en base a P en cuatro conjuntos $C(\rho_P(\alpha))$, $Hyp(\rho_P(\alpha))$, $Sp(\rho_P(\alpha))$, y $Ref(\rho_P(\alpha))$, pudiendo ocurrir en general que la partición sea distinta en cada nivel, es decir, que por ejemplo para un $l \in L$ sea $l \in C(\rho_P(\alpha))$ y $l \notin C(\rho_P(\beta))$ para $\alpha \neq \beta$. Por tanto, una sentencia $l \in L$ será en general consecuencia, hipótesis, especulación o refutación de P con cierto grado.

Como consecuencia, $C(P)$, $Hyp(P)$, $Sp(P)$ y $Ref(P)$ no son conjuntos crisp. Es directo proponer RLs de cada uno de los conjuntos anteriores: en el caso de las consecuencias, $\rho_{C(P)}(\alpha) = C(\rho_P(\alpha))$, y de manera similar para el resto. También es posible obtener a partir de cada RL un conjunto difuso midiendo el grado de pertenencia de cada sentencia $l \in L$ al RL, obteniéndose el conjunto difuso $\tilde{C}(P) : L \rightarrow [0, 1]$ a partir de $\rho_{C(P)}$, y de forma similar $\tilde{Hyp}(P)$, $\tilde{Sp}(P)$, y $\tilde{Ref}(P)$ a partir de sus respectivas RLs. Se puede demostrar además que estos cuatro conjuntos difusos forman una partición difusa de L en el sentido de Ruspini.

Referencias

- [1] Itziar García-Honrado and Enric Trillas. On an attempt to formalize guessing. In Rudolf Seising and Veronica Sanz, editors, *Soft Computing in Humanities and Social Sciences*, volume 273 of *Studies in Fuzziness and Soft Computing*, pages 237–255. Springer, 2012.
- [2] D. Sánchez, M. Delgado, M.A. Vila, and J. Chamorro-Martínez. On a non-nested level-based representation of fuzziness. *Fuzzy Sets and Systems*, 192(1):159–175, 2012.
- [3] Enric Trillas, Itziar García-Honrado, and Ana Pradera. Consequences and conjectures in preordered sets. *Inf. Sci.*, 180(19):3573–3588, 2010.

Enunciados condicionales y enunciados causales, precisos e imprecisos: analogías y diferencias.

Alejandro Sobrino

Facultad de Filosofía. Universidad de Santiago de Compostela, España.

alejandro.sobrino@usc.es

Los enunciados condicionales transmiten a veces contenido causal, y los enunciados causales son expresados frecuentemente en un formato condicional. Así, *Si tomo un gramo de cianuro, muero* puede ser expresado en forma causal diciendo que *Tomar un gramo de cianuro causa la muerte*. Pero el enunciado *Si Oslo es la capital de Noruega, Madrid es la capital de España*, no puede ser parafraseada como *Que Oslo sea la capital de Noruega causa que Madrid sea la capital de España*. Causalidad y condicionalidad presentan analogías, pero también diferencias. En este trabajo nos dedicaremos a explorarlas.

Una primera tipología de los enunciados condicionales distingue a los condicionales materiales de los condicionales estrictos. Los condicionales materiales expresan una relación de condicionalidad real o ficticia entre un antecedente y un consecuente y se formalizan así: $p \rightarrow q$. Los condicionales estrictos revelan una relación de necesidad entre antecedente y consecuente y se formalizan con el operador modal de necesidad: $\Box(p \rightarrow q)$ que significa: es necesario que $p \rightarrow q$. Esto equivale a decir que no hay ningún universo de discurso imaginable (ningún mundo posible) que haga al antecedente verdadero y falso al consecuente. Los condicionales estrictos están más próximos a expresar contenido causal que los condicionales materiales, pero no lo hacen siempre, porque los enunciados causales indican una relación real entre causa y efecto que en los condicionales estrictos es posible evitar haciendo de modo artificial al antecedente necesariamente falso (p. ej., afirmando que $2+2=6$) o al consecuente necesariamente verdadero (con un truco similar al anterior).

Los enunciados condicionales pueden ser también indicativos o subjuntivos. Si son subjuntivos, pueden ser de tiempo pasado o de tiempo futuro. Los más frecuentes son los de tiempo pasado y se les denomina también contrafácticos, porque verbalizan la negación de un hecho anteriormente afirmado. De los enunciados contrafácticos se dice que sirven para diagnosticar causalidad, porque sugieren que el efecto no se hubiera dado si la causa no hubiese ocurrido. Un contrafáctico puede expresar: (i) que la causa, de facto, no se da, (ii) en el caso pluricausal, que se puede interrogar sobre las consecuencia que tendría sobre el efecto intervenir sobre algún factor causal alternativo al considerado.

Dependiendo de la estructura del mecanismo causal, podemos conjeturar como se transforma la dependencia o independencia entre los diversos factores causales después de intervenir en alguno de ellos en una sucesión causal. Estos son los escenarios estándar: (1) Conexión serial ($A \rightarrow B \rightarrow C$): A causa B y B causa C. Una intervención sobre B deshabilita el posible influjo de A sobre C; esto es, hace a C independiente de A. (2) Causa común (*confounder*) $A \leftarrow B \rightarrow C$. Al ser A y C efecto de una misma causa (B), debería mostrar dependencia o asociación. Pero si se interviene sobre B, saber de A no añade nada a C; A y C son ahora independientes. (3) Causa múltiple (*collider*) ($A \rightarrow B \leftarrow C$). B es causado por A y C. En principio, A y C son independientes. Pero si se sabe que B ocurre y A es la causa, abre la comunicación con C (ahora se sabe que la causa no es C), mostrándose así dependiente de A.

Los condicionales, con o sin contenido causal, se usan en las reglas de inferencia lógicas como el MP: $(a \rightarrow b \wedge a) \rightarrow b$ o el MT: $(a \rightarrow b \wedge \neg b) \rightarrow \neg a$. Dos variantes de estas reglas se conocen como falacias de negación del antecedente NA: $(a \rightarrow b \wedge \neg a) \rightarrow \neg b$ y

afirmación del consecuente AC: $(a \rightarrow b \wedge b) \rightarrow a$. Mientras en una lógica inferencial los dos primeros esquemas son válidos y los dos últimos incorrectos, en una lógica causal se pone en cuestión la validez del MT como esquema correcto de inferencia causal. Además, si se considera el contenido de los argumentos y no sólo su forma, la veracidad de los diferentes esquemas se ve reforzada o mermada por la existencia o no de alternativas causales; es decir, de causas potenciales que deshabilitan o confirman al enlace causal. En general, la presencia de factores que deshabilitan afecta a la credibilidad de los argumentos MP y NA, y la presencia de causas alternativas a los argumentos MT y AC.

Los condicionales se pueden expresar también en términos de relaciones de orden o de relaciones conjuntistas: $p \rightarrow q \equiv p \leq q \equiv p \subseteq q$. La teoría de retículos y la teoría de conjuntos son ambos modelos, como la lógica proposicional, del álgebra de Boole. En el caso clásico, la cardinalidad de los conjuntos es precisa y las relaciones de orden o conjuntistas, determinadas. Pero hay condicionales con contenido causal que responden a una relación borrosa, como en *Si vistes un traje gris, irás de oscuro*. En ese caso, ‘gris’ y ‘oscuro’ son conjuntos borrosos y la relación de orden admite grados. En el ámbito de la imprecisión, el análisis causal de las oraciones condicionales con predicados vagos se ha vinculado tradicionalmente a la inclusión borrosa, cuya definición recuerda a la de la probabilidad condicional.

En el caso unicausal, la definición de inclusión de Zadeh ($A \subseteq B$ si y sólo si $\mu_A(x) \leq \mu_B(x)$) es aceptada, aunque parece rara, porque siempre tiene solución precisa. Pero, en situaciones pluricausales, antecedente y consecuente no son siempre completamente comparables. E. Sánchez propuso a la cardinalidad escalar para calcular la inclusión borrosa entre conjuntos. Usando esa definición, Hegalson y Jobe ofrecieron, en un contexto médico, un cálculo de cómo se afectan mutuamente varias causas para obtener el efecto deseado [1]

Como se dijo, los condicionales pueden formar parte de las reglas de inferencia como el MP. En el caso borroso, el MP generalizado (MPG) es la regla de inferencia por antonomasia. Permite, de las premisas ‘Si X es P entonces Y es Q’ y ‘X es P* (P* similar, pero no idéntico a P), inferir Y es Q*’. Trillas [2] planteó una versión del MPG en términos de una relación de orden:

$$\text{grado}('x \text{ es } P^*' \text{ y } (x, y) \text{ es } C(P, Q)) \leq \text{grado}(Y \text{ es } Q^*)$$

siendo y , una conjunción modelada por una t-norma y $C(P, Q)$ una función de condicional. Dependiendo de cuál sea la t-norma y la función de condicional escogidas, la relación de orden de arriba se satisfará o no. En general, tendrá tantas soluciones como modelos de y y C satisfagan esa desigualdad.

Un tipo de razonamiento aproximado es el razonamiento por analogía. Se caracteriza porque hay un argumento fuente, (*source*), ya conocido, que se recupera para hacer otro razonamiento, (*target*), por semejanza con él. La analogía puede ser de forma o de contenido. En la analogía por contenido, la causalidad juega un papel relevante, porque como relación de orden superior que es, fortalece la fuerza de los argumentos basados en una relación de analogía.

El análisis de estas y otras cuestiones debe contribuir a clarificar la relación entre enunciados causales y enunciados condicionales tanto en el caso preciso como aproximado.

Referencias:

[1] Helgason, C. M. & Jobe, T. H. (2007), ‘Stroke is a Dynamic Process Best Captured Using a Fuzzy Logic Based Scientific Approach to Information and Causation’, *International Journal of Computing in Medical Sciences and Image Processing*, Vol. 1, No. 1, 1, 5-9.

[2] Trillas, E., *Diálogos a la frágil luz de la razón inexacta*. Inédito.

Ecuaciones funcionales relacionadas con conjuntos fuzzy, ordenaciones representables y quasi-métricas ponderadas.

María Jesús Campión¹, Raquel G. Catalán¹, Esteban Indurain¹, Gustavo Ochoa¹, Oscar Valero²

¹ *INAMAT y Departamento de Matemáticas. Universidad Pública de Navarra*
mjcampion@unavarra.es, raquel.garcia@unavarra.es, steiner@unavara.es,
ochoa@unavarra.es

² *Departament de Ciències Matemàtiques i Informàtica. Universitat de les Illes Balears*
o.valero@uib.es

En los últimos años, se ha demostrado que las quasi-métricas ponderadas son útiles a la hora de modelizar muchos procesos en Ciencias de la Computación. En el trabajo [5], partiendo de la definición de quasi-métrica ponderada se observa que se inducen de manera natural varias ecuaciones funcionales. De hecho, las quasi-métricas ponderadas se caracterizan como las quasi-métricas que satisfacen cierta ecuación funcional de invarianza de circuitos. Además, la función de disimetría de una quasi-métrica ponderada satisface la ecuación funcional de Sincov e induce un preorden total, diferente, en general, del orden de especialización inducido directamente por la quasi-métrica dada. Esto nos permitió, en dicho trabajo establecer un link entre nociones aparentemente dispares: i) quasi-métricas ponderadas, ii) funciones reales que satisfacen la ecuación funcional de Sincov, y iii) preórdenes totales que se representan a través de una función de utilidad real.

A su vez, en el trabajo [7] se muestra que el concepto de subconjunto fuzzy se puede considerar desde un punto de vista alternativo, concretamente, la función característica de un subconjunto fuzzy se puede reinterpretar en términos de la ecuación funcional de Sincov y por tanto relacionar con la existencia de ciertos preórdenes totales representables.

Así, considerando los comentarios anteriores, analizaremos la relación entre varias teorías que, de algún modo, se pueden considerar equivalentes, porque podemos usar cada una de ellas para interpretar las otras. Mostraremos los resultados obtenidos como consecuencia de dicho análisis, así como aplicaciones y cuestiones abiertas que nos planteamos.

Referencias

- [1] Abrísqueta, F.J., Candeal, J.C., Catalán, R.G., De Miguel, J.R. and Induráin, E. *Generalized Abel functional equations and numerical representability of semiorders*, Publ. Math. Debrecen **78** (2011) 557-568.
- [2] Aczél, J. *Lectures on Functional Equations and their Applications*, (Academic Press, New York and San Diego, 1966).
- [3] Aczél, J. *A Short Course on Functional Equations Based upon Recent Applications to the Social and Behavioral Sciences*, (D. Reidel, Dordrecht, 1987)

- [4] Campión, M.J., Candeal, J.C., Induráin, E. : *Representability of binary relations through fuzzy numbers*. Fuzzy Sets and Systems **157** (2006) 1–19.
- [5] Campión, M.J., Induráin, E., Ochoa, G., Valero, O.: *Functional equations related to weightable quasi-metrics*. Hacettepe Journal of Mathematics and Statistics. **44** (4) (2015) 775–787.
- [6] Campión, M.J., Catalán, R. G., Garay, J., Induráin, E.: *Functional equations related to fuzzy sets and representable orderings*. Journal of Intelligent & Fuzzy Systems. **29**(3) (2015) 1065–1077.
- [7] Campión, M.J., Catalán, R. G., Induráin, E., Ochoa, G.: *Reinterpreting a fuzzy subset by means of a Sincov's functional equation*. Journal of Intelligent & Fuzzy Systems. **27** (2014) 367–375.
- [8] Gronau, D. *A remark on Sincov's functional equation*, Not. S. Afr. Math. Soc. **31**, 1-8, 2000.
- [9] Hitzler, P. and Seda, A.K. *Generalized distance functions in the theory of computation*, The Computer Journal **53**, 443-464, 2010.
- [10] Künzi, H.P.A. and Vajner V. *Weighted quasi-metrics*, Ann. NY Acad. Sci. **728**, 64-77, 1994.
- [11] Llull-Chavarría, J. and Valero, O. *An application of generalized complexity spaces to denotational semantics via domain of words*, LNCS **5457**, 530-541, 2009.
- [12] Romaguera, S., Sánchez-Pérez, E.A. and Valero, O. *Quasi-normed monoids and quasi-metrics*, Publ. Math. Debrecen **62**, 53-69, 2003.
- [13] Romaguera, S., Sapena, A. and Valero, O. *Quasi-uniform isomorphisms in fuzzy quasi-metric spaces, bicompletion and D-completion*, Acta Math. Hungar. **114**, 49-60, 2007.
- [14] Romaguera, S., Tirado, P. and Valero, O. *New results on mathematical foundations of asymptotic complexity analysis of algorithms via complexity spaces*, Int. J. Comput. Math. **89**, 1728-1741, 2012.
- [15] Romaguera, S. and Valero, O. *Asymptotic complexity analysis and denotational semantics for recursive programs based on complexity spaces*, in: Semantics-Advances in Theories and Mathematical Models, Afzal M.T. (Ed.), (InTech Open Science, Rijeka, 2012), 99-120.
- [16] Seda, A.K. *Quasi-metrics and the Semantics of Logic Programs*, Fund. Inform. **29**, 97-117, 1997.
- [17] Wilson, W.A. *On quasi-metric spaces*, Amer. J. Math. **53**, 675-684, 1931.
- [18] Zadeh, L.A. (1965). *Fuzzy Sets. Information and Control*, 8, 338–353.

Conjuntos Fuzzy y la Ecuación de Separabilidad

María Jesús Campión¹, Raquel G. Catalán¹, José Garay², Esteban Indurain¹

¹ INAMAT y Departamento de Matemáticas Universidad Pública de Navarra
mjcampion@unavarra.es, raquel.garcia@unavarra.es, steiner@unavara.es

² Departamento de Matemáticas Universidad de Zaragoza
jgaray@unizar.es

Las ecuaciones funcionales aparecen de forma natural cuando tratamos con conjuntos difusos. La simple definición de un subconjunto fuzzy $S \subset X$ de un conjunto X a través de la función de pertenencia $\mu_S : X \rightarrow [0, 1]$ origina, a través de la función $F(x, y) = \mu_S(y) - \mu_S(x)$ una solución de la ecuación de Sincov $F(x, y) + F(y, z) = F(x, z)$. Recíprocamente, las soluciones acotadas de una ecuación de Sincov pueden ser descritas a través de funciones de pertenencia de subconjuntos difusos. Las relaciones presentes entre estos subconjuntos, las soluciones de la ecuación de Sincov y los preórdenes completos representables ligados a las funciones de pertenencia de los subconjuntos difusos fueron estudiadas en [JIFS2014].

Esta relación entre subconjuntos fuzzy y ecuaciones funcionales fue posteriormente extendida y presentada en el ESTYLF2014, mostrando allí cómo los resultados obtenidos para la ecuación de Sincov podían adaptarse para una de las generalizaciones naturales de dicha ecuación, en la que la nulidad de las soluciones sobre la diagonal de $X \times X$ no es impuesta, aunque se mantiene la necesidad de ser constante: $F(x, y) + F(y, z) = F(x, z) + \lambda$. Estas soluciones están relacionadas con los semiórdenes representables.

El siguiente paso obvio, y anunciado en ESTYLF2014, era tratar de ver hasta dónde los resultados obtenidos podían extenderse para ilustrar el comportamiento de las soluciones de la ecuación de Separabilidad $F(x, y) + F(y, z) = F(x, z) + F(y, y)$ en términos de subconjuntos difusos y de los órdenes intervalo representables. Este estudio se desarrolló en gran medida en [JIFS2015] y en el mismo se muestra cómo ligados a cada solución de la ecuación de Separabilidad aparecen dos subconjuntos Fuzzy S, T , entrelazados, anidados los soportes de sus funciones de pertenencia μ_S, μ_T , hasta el punto de que el supremo del soporte de cualquiera de ellas viene determinado por el ínfimo del soporte de la otra y por las cotas de la función solución de la ecuación de Separabilidad. Llegar a entender su comportamiento mutuo en relación con las características de la solución de la ecuación funcional no podía hacerse con un estudio similar al realizado para las ecuaciones funcionales más restringidas (Sincov y Sincov Generalizada) y supuso, entre otras cosas, ampliar los parámetros de acotación de dicha solución de la ecuación funcional.

Nos proponemos además reinterpretar los variados resultados obtenidos para las funciones μ_S y μ_T asociadas a una solución acotada de la ecuación de Separabilidad, cuando exigimos las condiciones de ser constante sobre la diagonal de $X \times X$ o nula. En este sentido son especialmente interesantes los resultados sobre la unicidad de representación.

Agradecemos al departamento de Matemáticas de la UpNa.

Referencias

- [1] M.J. Campión, R.G. Catalán, E. Indurain, G. Ochoa. *Reinterpreting a fuzzy subset by means of a Sincov's functional equation*. Journal of Intelligent Fuzzy Systems. **27** (2014) 367-375.
- [2] M.J. Campión, R.G. Catalán, J. Garay, E. Indurain *Functional Equations related to fuzzy sets and representable orderings*. Journal of Intelligent Fuzzy Systems. **No** (2015) ??–??.

Characterizations of lattice OWA operators

Laura De Miguel¹, Daniel Paternain², Gustavo Ochoa³, Inmaculada Lizasoain⁴

¹ *Departamento de Automática y Computación. Universidad Pública de Navarra, 31006 Pamplona*

laura.demiguel@unavarra.es

² *Departamento de Automática y Computación. Universidad Pública de Navarra. ISC Institute of Smart Cities, 31006 Pamplona*

daniel.paternain@unavarra.es

³ *Departamento de Matemáticas. Universidad Pública de Navarra, 31006 Pamplona*

ochoa@unavarra.es

⁴ *Departamento de Matemáticas. INAMAT Institute for Advanced Materials. Universidad Pública de Navarra, 31006 Pamplona*

ilizasoain@unavarra.es

OWA (ordered weighted average) operators, defined by Yager in [6] for real values in the unit interval, have been described in different ways [4] as:

1. Weighted arithmetic means of the given values previously ordered [6].
2. Choquet integrals with respect to symmetric capacities [2].
3. Symmetric comonotone additive aggregation functions (from results in [2] and [5]).

In [3], OWA operators were extended from $[0, 1]$ to a complete lattice $(L, \leq_L, T, S)$, where T and S denote respectively a t-norm and a t-conorm defined on L . In this paper we analyze some generalizations of each description of OWA operators given above to this new setting.

1. The generalization of this item to a complete lattice $(L, \leq_L, T, S)$ as above, is simply the definition:

Definition 1 ([3]) *For each distributive weighting vector $\alpha = (\alpha_1, \dots, \alpha_n) \in L^n$, the OWA operator $F_\alpha : L^n \rightarrow L$ is a function assigning to any $(a_1, \dots, a_n) \in L^n$, the following lattice weighted average:*

$$F_\alpha(a_1, \dots, a_n) = S[T(\alpha_1, b_1), \dots, T(\alpha_n, b_n)],$$

where $b_1 \geq_L b_2 \geq_L \dots \geq_L b_n$ is the following rearrangement of the given values, $a_1, \dots, a_n \in L$:

$$\begin{aligned} b_1 &= a_1 \vee \dots \vee a_n \in L; \quad b_2 = \bigvee \{a_i \wedge a_j \mid i < j \in \{1, \dots, n\}\} \in L; \dots; \\ b_k &= \bigvee \{a_{j_1} \wedge \dots \wedge a_{j_k} \mid j_1 < \dots < j_k \in \{1, \dots, n\}\} \in L; \dots; \\ b_n &= a_1 \wedge \dots \wedge a_n \in L. \end{aligned}$$

This work has been partially supported by Spanish project TIN2013-40765-P

2. In this case, we consider a complete distributive lattice $(L, \leq_L, \wedge, \vee)$ endowed with the t-norm given by the meet and the t-conorm given by the join. The generalization of description 2. to these lattices is established in the following result:

Theorem 1 (Th. 4.8 and Th. 4.9 in [3]) *Let $(L, \leq_L, \wedge, \vee)$ be a complete distributive lattice, $P = \{1, \dots, n\}$ and $\mathcal{P}$ the set of all the subsets of P . If $\mu : \mathcal{P} \rightarrow L$ is an L -fuzzy measure on P , then the discrete Sugeno integral $S_\mu : L^n \rightarrow L$ given by*

$$S_\mu(a_1, \dots, a_n) = \bigvee \{ \bigwedge \{a_i \mid i \in X\} \wedge \mu(X) ; X \in \mathcal{P} \} \quad (a_1, \dots, a_n) \in L^n$$

is an OWA operator if and only if μ is symmetric, i. e., $\mu(X) = \mu(Y)$ whenever $\text{card} X = \text{card} Y$.

3. The concept of symmetric aggregation function is generalized to the context of complete lattices in a natural way. Now, the following result can be shown:

Theorem 2 *Let $(L, \leq_L, \wedge, \vee)$ be a complete distributive lattice. Then, the following assertions are equivalent for any $F : L^n \rightarrow L$:*

- (a) *F is a lattice OWA operator.*
- (b) *F is a symmetric aggregation function which is both $\vee_L$ -homogeneous and $\wedge_L$ -homogeneous.*

These different ways of looking at lattice OWA operators allows us to use a wide variety of tools, taken from different theoretical fields, to work with this kind of lattice operators.

References

- [1] M. Couceiro, J.-L. Marichal. *Characterizations of discrete Sugeno integrals as polynomial functions over distributive lattices*. Fuzzy Sets and Systems **161** (2010) 694–707.
- [2] M. Grabisch. *Fuzzy integral in multicriteria decision making*. Fuzzy Sets and Systems **69** (1995) 279–298.
- [3] I. Lizasoain, C. Moreno. *OWA operators defined on complete lattices*. Fuzzy Sets and Systems **224** (2013) 36–52.
- [4] R. Mesiar, A. Stupňanová, R. Yager. *Generalizations of OWA operators*. Proceedings of IFSA-EUSFLAT (2015).
- [5] D. Schmeidler. *Integral representation without additivity*. Proc. Amer. Math **97** (1986) 255–270.
- [6] R. R. Yager. *On ordered weighting averaging aggregation operators in multicriteria decision-making*. IEEE Transaction on Systems, Man and Cybernetics **18** (1988) 183–190.
- [7] R. R. Yager. *Families of OWA operators*. Fuzzy Sets and Systems **59** (1993) 125–148.

Convolución de funciones entre retículos

L. De Miguel¹, H. Bustince¹, B. De Baets²

¹ *Departamento de Automática y Computación, Universidad Pública de Navarra*

² *Department of Mathematical Modelling, Statistics and Bioinformatics, Ghent University*

La convolución es una operación matemática que transforma dos funciones con el mismo conjunto dominio e imagen en una tercera función que, en cierta medida, resume la información de ambas funciones. En particular, la convolución puede ser utilizada para construir distintas funciones de agregación [1, 4, 2].

Dadas dos funciones cualesquiera $f, g : U \mapsto V$, una definición general de la operación de convolución [3] viene dada por la ecuación:

$$(f \bullet g)(u) = \nabla_{x \circ y = u} (f(x) \blacktriangle g(y))$$

donde $\circ$ es una operación binaria en U , $\blacktriangle$ es una operación binaria en V y ∇ es otra operación en V que puede ser definida para cualquier número de variables de entrada.

En particular en este trabajo consideramos el conjunto de funciones cuyo dominio es un retículo acotado L_1 y su conjunto imagen un retículo completo y completamente distributivo L_2 . Es decir, consideramos el conjunto,

$$\mathbb{F} = \{f \mid f : L_1 \longrightarrow L_2\}.$$

Mediante la convolución definimos las funciones $\sqcup$ and $\sqcap$, dadas por, para todo $f_1, f_2 \in \mathbb{F}$:

$$f_1 \sqcup f_2(z) = \bigvee_{x \vee y = z} (f_1(x) \wedge f_2(y)) \quad (1)$$

$$f_1 \sqcap f_2(z) = \bigvee_{x \wedge y = z} (f_1(x) \wedge f_2(y)). \quad (2)$$

donde $\vee$ and $\wedge$ son la unión e intersección de los retículos L_1 o L_2 .

El propósito del trabajo es estudiar qué estructura algebraica generan estas operaciones sobre el conjunto $\mathbb{F}$. Nótese que un estudio similar para conjuntos difusos de tipo 2, cuando $L_1 = [0, 1]$ y $L_2 = [0, 1]$ puede encontrarse en [3]. Sin embargo, el paso del retículo $[0, 1]$ a un retículo general cualquiera no resulta inmediato porque algunas propiedades, como la idempotencia no se cumplen. Ello nos obliga a restringir el conjunto de funciones a aquellas funciones que cumplen ciertas propiedades. En particular, consideramos los siguientes conjuntos:

El trabajo ha sido financiado por los proyectos TIN2013-40765 y el servicio de investigación de la Universidad Pública de Navarra

- i) $\mathbb{N} = \{f \in \mathbb{F} \mid \bigvee_{x \in L_1} f(x) = 1\}$;
- ii) $\mathbb{C} = \{f \in \mathbb{F} \mid \text{para todo } x_1 \leq x_2 \leq x_3, f(x_2) \geq f(x_1) \wedge f(x_3)\}$;
- iii) $\mathbb{I} = \{f \in \mathbb{F} \mid \text{para todo } x_1, x_2 \in L_1, f(x_1 \vee x_2) \geq f(x_1) \wedge f(x_2) \text{ y } f(x_1 \wedge x_2) \geq f(x_1) \wedge f(x_2)\}$;
- iv) $\mathbb{S} = \{f \in \mathbb{F} \mid \text{existe } x^* \in L_1 \text{ tal que } f(x^*) = 1 \text{ y para todo } x \neq x^*, f(x) = 0\}$.

Teniendo en cuenta estos conjuntos de funciones se pueden demostrar los siguientes resultados.

Teorema 1 *Las operaciones $\sqcup$ y $\sqcap$ dadas en las Ecuaciones (1) y (2) generan un retículo en cualquier conjunto cerrado del conjunto $\mathbb{C} \cap \mathbb{I} \cap \mathbb{N}$.*

Nótese que puede demostrarse que $\mathbb{C} \cap \mathbb{I} \cap \mathbb{N}$ no es un conjunto cerrado y por lo tanto, las operaciones no generan un retículo en $\mathbb{C} \cap \mathbb{I} \cap \mathbb{N}$.

Proposición 2 *Se pueden demostrar los dos siguientes resultados.*

- i) $\mathbb{S} \subset \mathbb{C} \cap \mathbb{I} \cap \mathbb{N}$ es un subconjunto cerrado.
- ii) Si L_1 es un retículo distributivo entonces $\mathbb{C} \cap \mathbb{I} \cap \mathbb{N}$ es un conjunto cerrado.

Del Teorema 1 y la Proposición 2 se infiere lo siguiente.

Corolario 3 *Las operaciones $\sqcup$ y $\sqcap$ dadas en las Ecuaciones. (1) y (2) generan un retículo en $\mathbb{S}$.*

Corolario 4 *Sea L_1 un retículo distributivo. Entonces las operaciones dadas en las Ecuaciones (1) y (2) generan un retículo en $\mathbb{S}$. Además, el retículo generado es distributivo.*

Referencias

- [1] P. Hernandez, S. Cubillo, S. y C. Torres-Blanc: *On T-Norms for Type-2 Fuzzy Sets*. IEEE Transactions on Fuzzy Systems, **23**, no 4 (2015), pp.1155-1163.
- [2] D. Li: *Type-2 triangular norms and their residual operators*. Information Sciences, **317** (2015), 259–277.
- [3] C.L. Walker y E.A. Walker: *The algebra of fuzzy truth values*. Fuzzy Sets and Systems, **149** (2005), 309–347.
- [4] C.L. Walker y E.A. Walker: *Sets with type-2 operations*. International Journal of Approximate Reasoning **50** (2009), 63–71.

Relaciones binarias definidas a través de ecuaciones funcionales

María Jesús Campión¹, Raquel G. Catalán¹, Laura De Miguel², Esteban Induráin¹

¹ *Universidad Pública de Navarra. Departamento de Matemáticas. 31006 Pamplona, Spain.*
mjesus.campion@unavarra.es, raquel.garcia@unavarra.es,
steiner@unavarra.es

² *Universidad Pública de Navarra. Departamento de Automática y Computación. 31006 Pamplona, Spain.*
laura.demiguel@unavarra.es

Partiendo de un conjunto X no vacío, consideramos aplicaciones bivalentes definidas en el producto cartesiano del conjunto por sí mismo, y tomando valores en la recta real. Siendo $F : X \times X \rightarrow \mathbb{R}$ una tal aplicación bivalente, genérica, podemos definir a partir de ella, de manera natural, una relación binaria sobre el conjunto X , relación que denotaremos $\mathcal{R}_F$, definida mediante $x\mathcal{R}_F y \Leftrightarrow F(x, y) > 0$ para cualquier par $(x, y) \in X^2$. La intuición que hay detrás de esta idea es que F , y por ende su relación asociada $\mathcal{R}_F$, establecen una preferencia entre los elementos de la pareja (x, y) de manera que si $F(x, y)$ es estrictamente positivo podemos interpretar, en cierto modo, que “el elemento x es preferido al elemento y ”. Sea cual sea la interpretación que hagamos aquí, el hecho interesante que analizaremos es el siguiente: Cuando F satisface alguna ecuación funcional estándar, la correspondiente relación binaria $\mathcal{R}_F$ debe satisfacer ciertas propiedades estructurales adicionales. He aquí varios ejemplos:

- 1) Si $F(x, x) = 0$ para todo $x \in X$, entonces $\mathcal{R}_F$ es irreflexiva.
- 2) Si $F(x, y) + F(y, x) = 0$ para todo $x, y \in X$ (lo que se denomina ecuación de hemisimetría), entonces $\mathcal{R}_F$ es asimétrica.
- 3) Si $F(x, y) + F(y, z) + F(z, x) = 0$ para todo $x, y, z \in X$, o equivalentemente si $F(x, y) + F(y, z) = F(x, z)$ para todo $x, y, z \in X$ (ecuación denominada de Sincov), entonces $\mathcal{R}_F$ es asimétrica y negativamente transitiva.
- 4) Si $F(x, y) + F(y, z) = F(x, z) + F(y, y)$ para todo $x, y, z \in X$ (conocida como ecuación de la separabilidad), entonces bajo ciertas condiciones adicionales $\mathcal{R}_F$ es un orden intervalo.

El problema recíproco al anterior, del que solamente se conocen ciertas soluciones parciales, consiste en identificar aquellas relaciones binarias $\mathcal{R}$ para las que se puede definir una adecuada función bivalente F tal que $\mathcal{R}$ y $\mathcal{R}_F$ coincidan. Por último, este tipo de técnicas e ideas basadas en ecuaciones funcionales pueden extenderse también al caso de relaciones binarias fuzzy. En nuestra comunicación analizaremos el “estado del arte” en relación con este tipo de cuestiones.

Este trabajo ha sido parcialmente financiado a través de los proyectos MTM2012-37894-C02-02, TIN2013-47605-P y TIN2011-29520 (M.E.C., España) y el Servicio de Investigación de la Universidad Pública de Navarra.

References

- [1] J.C. Candeal, E. Induráin, E. Olóriz: *Aplicaciones bivariantes que representan semi-grupos ordenados*. *Rev. R. Acad. Cien. Exact. Fis. Nat. (Esp)* **90** (3) (1996), 157–170.
- [2] E. Olóriz, J.C. Candeal, E. Induráin: *Representability of interval orders*. *J. Econom. Theory* **78** (1998), 219–227.
- [3] G. Bosi, J.C. Candeal, E. Induráin, E. Olóriz, M. Zudaire: *Numerical representations of interval orders*. *Order* **18** (2001), 171–190.
- [4] J.C. Candeal, E. Induráin, M. Zudaire: *Numerical representability of semiorders*. *Math. Social Sci.* **43** (2002), 61–77.
- [5] M. J. Campión, J.C. Candeal, E. Induráin: *Representability of binary relations through fuzzy numbers*. *Fuzzy Sets and Systems* **157** (2006), 1–19.
- [6] M. J. Campión, J.C. Candeal, E. Induráin, M. Zudaire: *Continuous representability of semiorders*. *J. Math. Psych.* **52** (2008), 48–54.
- [7] F.J. Abrísqueta, J.C. Candeal, E. Induráin, M. Zudaire: *Scott-Suppes representability of semiorders: internal conditions*. *Math. Social Sci.* **57** (2009), 245–261.
- [8] J.C. Candeal, E. Induráin: *Semiorders and thresholds of utility discrimination: Solving the Scott-Suppes representability problem*. *J. Math. Psych.* **54** (2010), 485–490.
- [9] F.J. Abrísqueta, J.C. Candeal, R.G. Catalán, J.R. De Miguel, E. Induráin: *Generalized Abel functional equations and numerical representability of semiorders*. *Publ. Math. Debrecen* **78** (3-4) (2011), 557–568.
- [10] M.J. Campión, J.C. Candeal, R. G. Catalán, J. R. De Miguel, E. Induráin and J. A. Molina: *Aggregation of preferences in crisp and fuzzy settings: functional equations leading to possibility results*. *Internat. J. Uncertain. Fuzziness Knowledge-Based Systems* **19** (1) (2011), 89–114.
- [11] J.C. Candeal, A. Estevan, J. Gutiérrez García, E. Induráin: *Semiorders with separability properties*. *J. Math. Psych.* **56** (12) (2012), 444–451.
- [12] F.J. Abrísqueta et al.: *New trends on the numerical representability of semiordered structures*. *Mathware Soft Comput.* **19**(1) (2012), 25–37.
- [13] M.J.Campión, R.G. Catalán, E. Induráin, G. Ochoa: *Reinterpreting a fuzzy subset by means of a Sincov's functional equation*. *J. Intell. Fuzzy Systems* **27** (2014), 367–375.
- [14] M.J. Campión, E. Induráin, G. Ochoa, O. Valero: *Functional equations related to weightable quasi-metrics*. *Hacet. J. Math. Stat.* **44** (4) (2015), 775–787.
- [15] M.J.Campión, R.G. Catalán, J. Garay, E. Induráin: *(Functional equations related to fuzzy sets and representable orderings)*. *J. Intell. Fuzzy Systems* **29** (3) (2015), 1065–1077.

Point-sensitive aggregation operators: functional equations

Juan Carlos Candeal¹, Esteban Induráin²

² *Universidad Pública de Navarra. Departamento de Matemáticas. 31006 Pamplona, Spain.*
steiner@unavarra.es

¹ *Universidad de Zaragoza. Departamento de Fundamentos del Análisis Económico. 50005 Zaragoza, Spain.*
candeal@unizar.es

We study aggregation operators that are point-sensitive, so that their values strongly depend on the point where the functions to be aggregated are defined, as well as on the values of those functions at that point. This analysis gives rise to consider several functional equations that appear in a natural way.

Given two abstract nonempty sets X and Y , suppose that we have at hand a finite collection of maps $\{f_1, f_2, \dots, f_n\}$ from the set X into the set Y . Then, any new map $f_{n+1} : X \rightarrow Y$ obtained somehow from the given ones $f_1, f_2, \dots, f_n$ is usually called an *aggregation* of those maps.

Three main situations can be considered at this stage:

- *Case 1:* Given an element $x \in X$, to obtain $f_{n+1}(x)$ we should know the maps $f_1, \dots, f_n$ as a whole, maybe in all the points of its domain, or compulsorily *at least in several points different from the point $x \in X$ of reference*.
- *Case 2:* Given an element $x \in X$, in order to get $f_{n+1}(x)$ *we only need to have at hand the values of $f_1, \dots, f_n$ at that point x* . That is: $f_{n+1}(x)$ directly comes from the n -tuple $(f_1(x), \dots, f_n(x)) \in Y^n$, and we may then assume that there exists a map $G : Y^n \rightarrow Y$ such that $f_{n+1}(x) = G(f_1(x), \dots, f_n(x))$ holds true for every $x \in X$.
- *Case 3:* Given an element $x \in X$, in order to get $f_{n+1}(x)$ *we need to know the values $f_1, \dots, f_n$ at that point x , as well as the given point $x \in X$* . That is: in this case, $f_{n+1}(x)$ perhaps cannot be directly obtained only from $(f_1(x), \dots, f_n(x))$. Instead, it can actually be got if we know $(x, f_1(x), \dots, f_n(x)) \in X \times Y^n$. The knowledge of the point x could be decisive at this stage. Moreover, we may then assume that there exists a map $H : X \times Y^n \rightarrow Y$ such that $f_{n+1}(x) = H(x, f_1(x), \dots, f_n(x))$ holds true for every $x \in X$.

These nuances are essential in our approach throughout the present paper. The second case corresponds to the so-called *pointwise aggregation* of maps, whereas the third case corresponds to *weakly pointwise aggregation* of maps (also known as *point-sensitive aggregation*, to which we pay an special attention here).

The authors acknowledge financial support from the Ministry of Economy and Competitiveness of Spain under the research projects ECO2012-34828, MTM2012-37894-C02-02 and TIN2013-40765-P.

Notice that here the knowledge of the point x is not necessary if we already know the n -tuple $(f_1(x), \dots, f_n(x))$. In other words, if $x \neq t \in X$ are such that $f_i(x) = f_i(t)$ ($i = 1, \dots, n$) then a fortiori $f_{n+1}(x) = f_{n+1}(t)$.

References

- [1] J.C. Candeal: *Social evaluation functionals: a gateway to continuity in social choice*. Soc. Choice Welf. **44(2)** (2015) 369–388.
- [2] J.C. Candeal,: *Aggregation operators, comparison meaningfulness and social choice*. *J. Math. Psych.* (Accepted, to appear, 2015).
- [3] L. De Miguel, M.J. Campión, J.C. Candeal, E. Induráin, D. Paternain: *First approach of type-2 fuzzy sets via fusion operators*. Proceedings of the IPMU 2014 Congress, in A. Laurent et al. (eds.), Part III, CCIS **444** (2014), 355–363. Montpellier, France.
- [4] H. Tizhoosh: *Image thresholding using type II fuzzy sets*. *Pattern Recognition* **38** (2005), 2363–2372.

Una introducción a las *FNI*-implicaciones

I. Aguiló, J. Suñer, J. Torrens

Universitat de les Illes Balears

Departamento de Ciencias Matemáticas e Informática

Crta. de Valldemossa, km. 7,5, 07122 Palma

{isabel.aguiló, jaume.sunyer, jts224}@uib.es

Las funciones de implicación borrosas se usan habitualmente para modelar los condicionales borrosos y también en los procesos de inferencia borrosa, por lo que resultan ser de vital importancia en el razonamiento aproximado y el control borroso [3]. El estudio de las funciones de implicación ha conocido un gran auge en las últimas décadas y ello ha permitido no sólo encontrar nuevas aplicaciones de las mismas sino también considerables avances en el campo desde el punto de vista teórico.

En este marco teórico, una de las líneas de investigación más estudiadas es la construcción de nuevas implicaciones a partir de otras dadas o a partir de otro tipo de operadores lógicos (ver capítulos 1 y 4 en [2]). En esta comunicación presentamos un método de construcción de implicaciones a partir de tres operadores lógicos como sigue.

Proposición 1 Sea F una función de agregación disyuntiva (en el sentido de que $F(1, 0) = F(0, 1) = 1$), N una negación e I una función de implicación borrosa. La función I_{FNI} dada por:

$$I_{FNI}(x, y) = F(N(x), I(x, y)) \quad \text{para todos } x, y \in [0, 1], \quad (1)$$

es siempre una función de implicación borrosa.

Definición 2 A una función de implicación borrosa I_{FNI} construida mediante la ecuación (1) la llamaremos una *FNI*-implicación. Cuando F sea una *t*-conorma S , diremos que I_{SNI} es una *SNI*-implicación.

En [4] se dejó como problema abierto el análisis de este tipo de implicaciones. Así, en este trabajo presentamos una introducción al estudio de las *FNI*-implicaciones y lo hacemos con un buen número de ejemplos que muestran que, en realidad, cualquier implicación puede obtenerse como una *FNI*-implicación (de hecho como una *SNI*-implicación). Además, estos ejemplos también muestran como muchas implicaciones, no previamente conocidas, pueden obtenerse mediante este método.

Por otra parte, la importancia de estos métodos de construcción radica en el hecho de que las implicaciones resultantes preserven las propiedades de las implicaciones iniciales. En este sentido se demuestra en este trabajo que las *FNI*-implicaciones preservan siempre la propiedad (*IP*) y que, con algunas condiciones sobre F y N , preservan también (*CB*), (*OP*), (*NP*) y (*CO*) entre otras (para el significado de todas estas propiedades y otras usuales de las funciones de implicación, ver por ejemplo [3]).

Este trabajo ha sido subvencionado por el proyecto TIN2013-42795-P.

Otra característica importante de este método es que las FNI -implicaciones se pueden interpretar como un método para modificar una función de implicación dada, I , que no verifica una determinada propiedad con el objetivo de que la I_{FNI} resultante sí la verifique. Todo ello de forma similar al caso de las contrapositivaciones que modifican una implicación para que verifique la contraposición respecto de una negación (ver [3] y [1]). En este sentido, en este trabajo se demuestra que, eligiendo adecuadamente la función de agregación F y la negación N , las FNI -implicaciones resultantes pueden verificar propiedades adicionales que no verifican las implicaciones inicialmente consideradas. De este modo, se pueden obtener FNI -implicaciones que verifican por ejemplo (IP) , (CO) o (SNP) , a partir de implicaciones que no verifican dichas propiedades.

Finalmente, queremos hacer notar que este trabajo es solo una introducción a las FNI -implicaciones y que éstas presentan diferentes líneas de trabajo abiertas. Dado que las FNI -implicaciones dependen de tres operadores diferentes, se pueden investigar desde diversas direcciones diferentes:

1. Se pueden fijar por ejemplo los operadores I y F obteniendo de esta forma implicaciones I_{FNI} que dependen de una negación N y estudiar las propiedades de estas implicaciones en función de las propiedades de N . Un ejemplo de esta línea se puede ver en [4] que es el origen del presente trabajo.
2. De manera similar se pueden fijar los operadores I, N obteniendo implicaciones I_{FNI} que dependen de una función de agregación F (o de una t -conorma S) y estudiar sus propiedades en función de las de F .
3. Finalmente, fijando los operadores N, F se obtienen funciones de implicación a partir de otra dada y se pueden analizar las propiedades que se preservan por esta construcción. Más aún, se puede investigar qué propiedades adicionales verifica la nueva implicación que no verificaba la inicial.

Referencias

- [1] I. Aguiló, J. Suñer, J. Torrens: *New types of contrapositivisation of fuzzy implications with respect to fuzzy negations*. Information Science. **322** (2015) 223–226.
- [2] M. Baczyński, G. Beliakov, H. Bustince Sola, A. Pradera (Eds.): *Advances in Fuzzy Implication Functions*. Studies in Fuzziness and Soft Computing, vol. 300. Springer, Berlin Heidelberg, 2013.
- [3] M. Baczyński, B. Jayaram: *Fuzzy Implications*. Studies in Fuzziness and Soft Computing, vol. 231. Springer, Berlin Heidelberg, 2008.
- [4] Y. Shi, B. Van Gasse, D. Ruan and E. Kerre: *Fuzzy implications: Classification and a new class*. In [2] (2013) pp. 31–53.

Negaciones para un orden admisible en conjuntos fuzzy valorados en intervalos

María Asiain¹, Humberto Bustince², Javier Fernández³, Razko Mesiar⁴, Anna Kolesárová⁵

¹ *Departamento de Matemáticas, Universidad Pública de Navarra,
Campus Arrosadía s/n, P.O. Box 31006, Pamplona, Spain*
asiain@unavarra.es

² *Departamento de Automática y Computación, Institute of Smart Cities, Universidad
Pública de Navarra,
Campus Arrosadía s/n, P.O. Box 31006, Pamplona, Spain*
bustince@unavarra.es

³ *Departamento de Automática y Computación, Universidad Pública de Navarra,
Campus Arrosadía s/n, P.O. Box 31006, Pamplona, Spain*
fcojavier.fernandez@unavarra.es

^{4,5} *Department of Mathematics and Descriptive Geometry Faculty of Civil Engineering,
Slovak University of Technology in Bratislava, Radlinského 11, 81368 Bratislava, Slovak
Republic*
radko.mesiar@math.sk, anna.kolesarova@math.sk

Los intervalos aparecen con frecuencia en problemas relativos a conjuntos fuzzy, especialmente cuando el grado de conocimiento de un dato tiene cierta incertidumbre que hace necesario atribuir a cada elemento un valor comprendido entre dos valores dados sin determinar exactamente el valor fijo exacto porque éste no es conocido por los expertos. En este sentido son numerosos los trabajos publicados que utilizan conjuntos fuzzy valorados en intervalos.

Es sencillo dotar a los intervalos de un orden parcial definido mediante un orden en $\mathbb{R}$ aplicado a los puntos extremos del intervalo. Como no existe un orden total definido de forma natural en los intervalos, definir un orden total en ellos requiere otras técnicas. El orden total más conocido es el de Xu y Yager. En [2] H. Bustince, J. Fernández, R. Mesiar y A. Kolesárová definen órdenes lineales para intervalos mediante funciones de agregación y estos órdenes admisibles permiten extender a los conjuntos fuzzy valorados en intervalos las herramientas matemáticas usuales utilizadas en los conjuntos fuzzy ordinarios. En [1] H. Bustince, E. Barrenechea y M. Pagola estudian los órdenes $K_{\alpha,\beta}$, de los que son casos particulares el orden de Xu y Yager y los órdenes lexicográficos.

En la literatura, sin embargo, no aparecen funciones de negación compatibles con estos órdenes, por ello resulta útil encontrar dichas funciones de negación, y más si se trata de negaciones fuertes. En el trabajo, utilizando distintas negaciones fuertes en $\mathbb{R}$ y atendiendo a la amplitud de los intervalos, definimos negaciones y negaciones fuertes para los órdenes $K_{\alpha,\beta}$, en particular para el orden de Xu y Yager y los órdenes lexicográficos.

La bibliografía más relevante en este trabajo es la que mencionamos a continuación:

Referencias

- [1] H. Bustince, E. Barrenechea, M. Pagola, Generations of Interval-Valued Fuzzy and Atanassov's Intuitionistic Fuzzy Connectives and from K_α Operators: Laws for Conjunctions and Disjunctions, *Amplitude, International Journal of Intelligent Systems* 23 (2008) 680–714.
- [2] H. Bustince, J. Fernández, A. Kolesárová, R. Mesiar, Generation of linear orders for intervals by means of aggregation functions, *Fuzzy Sets and Systems* 220 (2013) 69–77.
- [3] H. Bustince, E. Barrenechea, M. Pagola, J. Fernández, Z. Xu, B. Bedregal, J. Montero, H. Hagrais, F. Herrera and B. De Baets, A Historical Account of Types of Fuzzy Sets and Their Relationships, *IEEE Transactions on Fuzzy Systems* 24 (1) (2016) 174–194.

La noción de función de pre-agregación

Humberto Bustince^{1,2}, Aranzazu Jurio, Daniel Paternain, Mikel Galar, Jose Antonio Sanz, Mikel Elkano, Giancarlo Lucca¹ Radko Mesiar³, Anna Kolesárová⁴, y Graçaliz Dimuro⁵

¹ *Universidad Publica de Navarra
Campus Arrosadía s/n, 31006 Pamplona
Spain*

E-mail: { bustince, aranzazu.jurio, daniel.paternain, mikel.galar, joseantonio.sanz, mikel.elkano, }@unavarra.es

² *Institute of Smart Cities
Campus Arrosadía s/n, 31006 Pamplona
Spain*

³ *Slovak University of Technology, Radlinskeho 11,
Bratislava, Slovakia
E-mail: mesiar@math.sk*

⁴ *Institute of Information Engineering, Automation and Mathematics,
Slovak University of Technology,
SK-812 37 Bratislava, Slovakia E-mail: anna.kolesarova@stuba.sk* ⁵ *Centro de Ciências
Computacionais,
Universidade Federal do Rio Grande,
Rio Grande, Brazil, E-mail: gracalizdimuro@furg.br*

La monotonía es una de las características fundamentales de las funciones de agregación ([3]), tanto en el hipercubo unidad como cuando se consideran extensiones de los conjuntos difusos. Es un requisito natural en muchos campos de aplicación, desde la toma de decisión al procesamiento de imagen. Sin embargo, en algunos casos se trata de una condición muy exigente. Por ejemplo, consideremos el caso de la moda en procesamiento de imagen. Es un operador que es muy apropiado para eliminar determinados tipos de ruido. Sin embargo, la moda no es monótona, luego no puede ser considerada como una función de agregación ([1]).

En esta charla, basándonos en la idea de monotonía direccional [2], presentamos la noción de función de pre-agregación [4], como una función que satisface las mismas condiciones de frontera que una función de agregación usual, pero a la que solo se exige monotonía a lo largo de una dirección fijada. De esta forma, las funciones de pre-agregación incluyen muchos ejemplos de funciones que no son agregaciones, como es el caso, en particular, de la moda.

Discutiremos tres métodos de construcción de pre-agregaciones diferentes y presentaremos diversos ejemplos de funciones de pre-agregación que no son agregaciones. En particular, uno de los métodos de construcción considerados se basa en la definición usual de integral Choquet discreta, reemplazando el producto por otras funciones apropiadas. Las funciones de pre-agregación construidas de esta forma pueden utilizarse para diseñar sistemas de clasificación basados en reglas difusas que mejoran los resultados que aparecen en la literatura.

Este trabajo ha sido parcialmente financiado por los proyectos TIN2013-40765-P, APVV-14-0013 y 481283/2013-7, 306970/2013-9, 232827/2014-1, 307681/2012-2 y 233950/2014-1 del CNPQ.

Referencias

- [1] H. Bustince, E. Barrenechea, M. Pagola, J. Fernandez *Interval-valued fuzzy sets constructed from matrices: Application to edge detection*. Fuzzy Set and Systems. **160** (2009) 1819–1840.
- [2] H. Bustince, A. Kolesárová, J. Fernandez, and R. Mesiar, *Directional monotonicity of fusion functions*, European Journal of Operational Research, **244** (2015) 300–308.
- [3] T. Calvo, A. Kolesárová, M. Komorníková and R. Mesiar, Aggregation operators: properties, classes and construction methods. In T. Calvo, G. Mayor and R. Mesiar (Eds.): *Aggregation Operators New Trends and Applications* Physica-Verlag, Heidelberg,(2002) 3–104.
- [4] G. Lucca, J. Sanz, G. P. Dimuro, B. Bedregal, R. Mesiar, A. Kolesárová, and H. Bustince, *Pre-aggregation functions: construction and an application*, IEEE Transactions on Fuzzy Systems, in press, DOI: 10.1109/TFUZZ.2015.2453020.

Funciones monótonas direccionalmente ordenadas

Humberto Bustince, Edurne Barrenechea, Miguel Pagola, Javier Fernandez, Laura de Miguel, Carlos Lopez-Molina^{1,2}, Julio Lafuente, Mikel Sesma¹ Radko Mesiar³ y Anna Kolesárová⁴

¹ *Universidad Publica de Navarra
Campus Arrosadía s/n, 31006 Pamplona
Spain*

*E-mail: { bustince, edurne.barrenechea,miguel.pagola, fcojavier.fernandez,
laura.demiguel,carlos.lopez, lafuente}@unavarra.es*

² *Institute of Smart Cities
Campus Arrosadía s/n, 31006 Pamplona
Spain*

³ *Slovak University of Technology, Radlinskeho 11,
Bratislava, Slovakia
E-mail: mesiar@math.sk*

⁴ *Institute of Information Engineering, Automation and Mathematics,
Slovak University of Technology,
SK-812 37 Bratislava, Slovakia E-mail: anna.kolesarova@stuba.sk*

En los últimos tiempos existe un gran interés por el estudio de formas generalizadas de monotonía. [2], de tal modo que sea posible definir y/o cubrir muchas funciones que, a pesar de no ser funciones de agregación debido a la falta de monotonía usual, sí que resultan de gran interés para aplicaciones en los más diversos campos, como, por ejemplo, el procesamiento de imagen, la clasificación o la toma de decisión.

Un paso muy importante en esta dirección lo ha constituido la introducción del concepto de función de pre-agregación [3]. Una función de pre-agregación es una función que satisface las mismas condiciones de contorno que una función de agregación pero que, en lugar de ser creciente en el sentido usual, solo es creciente a lo largo de una dirección fijada. Las funciones de pre-agregación se han mostrado muy útiles en problemas de clasificación [3]. Además, dentro de las funciones de pre-agregación es posible incluir algunos casos muy relevantes que no encajan dentro del marco de las agregaciones, como es el caso de la moda.

La exigencia de monotonía a lo largo de una dirección fija que es la misma en todos los puntos del hipercubo unidad puede resultar, sin embargo, todavía demasiado estricta para algunas aplicaciones. En procesamiento de imagen, por ejemplo, cuando se trata de hallar los bordes en una imagen dada, las direcciones relevantes pueden cambiar de un píxel (punto) a otro [1].

Teniendo en cuenta esta observación, entre otras, en este trabajo presentamos la noción de función de pre-agregación ordenada direccionalmente, como sigue.

Este trabajo ha sido parcialmente financiado por los proyectos TIN2013-40765-P y APVV-14-0013.

Definición 1 Sea $F : [0, 1]^n \rightarrow [0, 1]$ una función y sea $\vec{r} \neq \vec{0}$ un vector n dimensional. F es direccionalmente ordenada $\vec{r}$ - creciente si para cualquier $\mathbf{x} \in [0, 1]^n$, para cualquier $c > 0$ y para cualquier permutación $\sigma : \{1, \dots, n\} \rightarrow \{1, \dots, n\}$ con $x_{\sigma(1)} \geq \dots \geq x_{\sigma(n)}$ y tal que

$$1 \geq x_{\sigma(1)} + cr_1 \geq \dots \geq x_{\sigma(n)} + cr_n \geq 0$$

se verifica que

$$F(\mathbf{x} + c\vec{r}_{\sigma^{-1}}) \geq F(\mathbf{x})$$

donde $\vec{r}_{\sigma^{-1}} = (r_{\sigma^{-1}(1)}, \dots, r_{\sigma^{-1}(n)})$.

Análogamente, F es direccionalmente ordenada $\vec{r}$ - decreciente si para cualquier $\mathbf{x} \in [0, 1]^n$, para cualquier $c > 0$ y para cualquier permutación $\sigma : \{1, \dots, n\} \rightarrow \{1, \dots, n\}$ con $x_{\sigma(1)} \geq \dots \geq x_{\sigma(n)}$ y tal que

$$1 \geq x_{\sigma(1)} + cr_1 \geq \dots \geq x_{\sigma(n)} + cr_n \geq 0$$

se verifica que

$$F(\mathbf{x} + c\vec{r}_{\sigma^{-1}}) \leq F(\mathbf{x})$$

donde $\vec{r}_{\sigma^{-1}} = (r_{\sigma^{-1}(1)}, \dots, r_{\sigma^{-1}(n)})$

De este modo la dirección a lo largo de la que se exige monotonía puede cambiar de un punto a otro, lo que permite tener en cuenta información relevante en un entorno de dicho punto. Esta noción da lugar a aplicaciones muy interesantes, especialmente en procesamiento de imagen, ya que permite construir detectores de bordes muy competitivos, como mostraremos en la charla.

Referencias

- [1] H. Bustince, E. Barrenechea, M. Pagola, J. Fernandez *Interval-valued fuzzy sets constructed from matrices: Application to edge detection*. Fuzzy Set and Systems. **160** (2009) 1819–1840.
- [2] H. Bustince, A. Kolesárová, J. Fernandez, and R. Mesiar, *Directional monotonicity of fusion functions*, European Journal of Operational Research, **244** (2015) 300–308.
- [3] G. Lucca, J. Sanz, G. P. Dimuro, B. Bedregal, R. Mesiar, A. Kolesárová, and H. Bustince, *Pre-aggregation functions: construction and an application*, IEEE Transactions on Fuzzy Systems, in press, DOI: 10.1109/TFUZZ.2015.2453020.

La clase $(\leq^\omega)_{\omega \in L}$ de órdenes de actividad en retículos distributivos $(L, \leq)$ y su familia de nulnormas asociada $(\wedge^\omega)_{\omega \in L}$

Ramón Fuentes-González¹
 Universidad Pública de Navarra
 rfuentes@unavarra.es

El marco general para el tratamiento de imágenes mediante la teoría de Morfología Matemática [1,2,3,4] se fundamenta a su vez en la teoría de retículos completos y distributivos $(L, \leq, \wedge, \vee)$ en los que los puntos $A \in L$ representan imágenes y el orden $\leq$ la inclusión entre ellas. En ese modelo, dada una imagen $A \in L$ se obtiene nuevas imágenes $\varphi(A) \in L$ mediante filtros morfológicos, que vienen dados por operadores $\varphi : L \rightarrow L$ que cumplen ciertas propiedades. (La identidad i en L tal que $i(A) = A \ \forall A \in L$ es uno de ellos). La clase $\mathfrak{F} \subseteq L^L$ de estos operadores es un retículo completo distributivo $(\mathfrak{F}, \preceq)$ con el orden: $\varphi_1 \preceq \varphi_2 \Leftrightarrow \varphi_1(A) \leq \varphi_2(A) \ \forall A \in L$. En este retículo $(\mathfrak{F}, \preceq)$, y asociada a cada $\omega \in \mathfrak{F}$, se introduce una nueva relación de orden auxiliar $\preceq^\omega$ en $\mathfrak{F}$. Por ejemplo, si $\omega = i$, la expresión $\varphi_1 \preceq^i \varphi_2$ traduce la situación: “el operador φ_2 es más activo que el operador φ_1 ”. En esa línea, para cualquier $\omega \in \mathfrak{F}$, la relación $\preceq^\omega$ se llama en la literatura el *orden de actividad relativo a ω* [2,3]. El estudio de los conjuntos ordenados $(\mathfrak{F}, \preceq^\omega)$ de operadores con los órdenes de actividad es un tópico importante en Morfología Matemática [4].

En [5] se ha utilizado el concepto de orden de actividad en un contexto diferente: el relativo a la medición de la ambigüedad y la borrosidad de sub-intervalos cerrados de $[0,1]$. En ese mismo estudio se demuestra que el orden de actividad recupera la definición de “orden aguzado” (*sharpened order*) [6] asociado al concepto de entropía en entornos borrosos.

En este trabajo continuamos el estudio de estos órdenes de actividad en un contexto general, proporcionando una versión equivalente y más simple que la inicial y también analizando relaciones entre esos órdenes y otros conceptos de la lógica borrosa.

En primer lugar, definimos relaciones de actividad en conjuntos parcialmente ordenados $(P, \leq)$ dirigidos y filtrados. Esas relaciones $\preceq^\omega$ con $\omega \in P$ resultan ser pre-órdenes. Si en P existe elemento mínimo 0 y/o elemento máximo 1 , entonces $\preceq^0$ coincide con el orden $\leq$ en P y $\preceq^1$ coincide con el dual $\geq$ del orden en P .

En [5] se prueba que si $(P, \leq)$ es un retículo distributivo $(L, \leq)$, las relaciones $\preceq^\omega$ sí que son relaciones de orden para cualquier $\omega \in L$ y además, para cada ω el par $(L, \preceq^\omega, \wedge^\omega)$ es semi-retículo con operador ínfimo $\wedge^\omega$ asociado al orden $\preceq^\omega$. En este caso, la familia $(\preceq^\omega)_{\omega \in L}$ constituye la clase de los *órdenes de actividad en el retículo distributivo* $(L, \leq)$. Se completa aquí ese resultado demostrando que si $(L, \leq)$ es completo, Broweriano y dual-Broweriano entonces el semi-retículo $(L, \preceq^\omega, \wedge^\omega)$ resulta ser completo, (es decir, existe $\inf^\omega M$ para todo subconjunto de L $M \neq \emptyset$).

¹ Miembro honorario del grupo de investigación de la Universidad Pública de Navarra: “Adquisición de conocimiento y minería de datos, funciones especiales y métodos numéricos avanzados”.

En un retículo distributivo $(L, \leq, \wedge, \vee)$, la relación de actividad $\preceq^\omega$ es tal que $\alpha \preceq^\omega \beta$ es equivalente a $(\omega, \alpha, \beta) \in B$, siendo $B \subseteq L^3$ la relación ternaria “intermediación” (*betweenness*) [7,8]. También, en estos casos, el elemento $\alpha \wedge^\omega \beta \in L$ coincide con el valor $m(\omega, \alpha, \beta)$ del operador $m: L^3 \rightarrow L$ conocido como *mediana* [9].

El operador ínfimo $\wedge^\omega$ en $(L, \preceq^\omega)$ es un *operador nulnorma* [10,11] en el retículo $(L, \leq)$. Analizamos los casos en los que $\wedge^\omega$ tiene operador residuo $\leadsto^\omega$.

También se caracteriza el orden de actividad en el retículo producto $(\prod_{j \in J} L_j, \leq)$ de una familia $((L_j, \leq))_{j \in J}$ de retículos distributivos.

Se estudia el orden de actividad en casos distinguidos de retículos: las Álgebras de Boole completas $(L, \leq, \wedge, \vee, ^c, 0, 1)$, demostrando que en cada una de éstas, la familia de órdenes $(\preceq^\omega)_{\omega \in L}$ tiene asociada otra nueva de operadores supremo: $(\vee^\omega)_{\omega \in L}$. Se obtiene así una familia de Álgebras de Boole completas $((L, \preceq^\omega, \wedge^\omega, \vee^\omega, ^c, \omega, \omega^c))_{\omega \in L}$, todas isomorfas al álgebra inicial $(L, \leq, \wedge, \vee, ^c, 0, 1)$. Se prueba también que ahora cada operador nulnorma $\wedge^\omega$ es *uninorma* [12] en $(L, \leq, \wedge, \vee, ^c, 0, 1)$.

Finalmente interpretamos tanto la relación de orden $\preceq^\omega$, como la ley $\wedge^\omega$ y su residuo $\leadsto^\omega$ en el álgebra de Boole del conjunto $\wp(E)$ de las partes de un referencial ordinario E .

Palabras clave Orden de actividad, uninormas, operadores nulnorma en retículos.

Referencias

- [1] J. Serra: *Image Analysis and Mathematical Morphology Vol 1*. Academic Press; (1984).
- [2] J. Serra (Editor): *Image Analysis and Mathematical Morphology Vol 2: Theoretical Advances*. Academic Press; 1st edition (1988).
- [3] H.J.A.M. Heijmans, R. Keshet, *Inf-semilattice approach to self-dual morphology*, Journal of Mathematical Imaging and Vision **Vol 17** (2002) 55–80.
- [4] J. Serra, L. Vincent: *An overview of morphological filtering*. Circuits, Systems and Signal Processing **Vol 11 No 1** (1992) 47–108.
- [5] C. Alcalde, A. Burusco, R. Fuentes-González: *Ambiguity and fuzzyness measures defined on the set of closed intervals in $[0, 1]$* . Information Sciences **Vol 177 Issue 7** (2007) 1687–1698.
- [6] A. De Luca, S. Termini. *A definition of a nonprobabilistic entropy in the setting of fuzzy sets theory*. Information and Control **Vol 20 Issue 4** (1972) 301–312.
- [7] J.M. Cibulskis: *A characterization of the lattice orderings on a set which induce a given betweenness*. Journal of the London Mathematical Society **Vol 1 (1)** (1969):480–482.
- [8] E. Pitcher and M.F. Smiley, *Transitives of betweenness*, Trans. Amer. Math. Soc. **Vol 52** (1942) 95–114.
- [9] M. Barbut: *Médiane, distributivité, éloignements*. Math. et Sci. Humaines **Vol 70** (1980) 5–31.
- [10] T. Calvo, B. de Baets, J.C. Fodor: *Te functional equations of Alsina and Frank for uninorms and nullnorms*. Fuzzy Sets and Systems **Vol 120** (2001) 15–24.
- [11] F. Karaçal, M.A. İnce, R. Mesiar: *Medians and nullnorms in bounded lattices*. Fuzzy Sets and Systems **Vol 289** (2016) 74–81.
- [12] R.R. Yager, A. Rybalov: *Uninorm aggregation operators*. Fuzzy Sets and Systems. **Vol 80** (1996) 11–120.

Hesitant fuzzy relations and their transitive closure

L. Garmendia¹, R. González del Campo¹, J. Recasens²

¹ *F. Informática, Universidad Complutense de Madrid, Spain*
lgarmend@fdi.ucm.es, rgonzale@estad.ucm.es

² *Universitat Politècnica de Catalunya, Spain*
j.recasens@upc.edu

Abstract

Hesitant fuzzy relations and some of their operations are defined.

A first definition of T-transitive property and its closure is defined.

Keywords: hesitant fuzzy sets, T-transitivity, transitive closure

Introduction

Hesitant fuzzy sets are a generalization of fuzzy sets that allow us to assign a few or several membership degrees to the elements of a (finite) universe, introducing the concept of possible hesitation of their membership degrees.

Operations and partial order $\leq_H$ on Lists.

Let $L \leq [0, 1]$ be a finite sorted list of membership degrees in $[0, 1]$, from now on just noted L . Sorted lists are Abstract Data Types (ADT) well known in computer programming sciences that allow at least to ask for their length, whether they are empty, and add and remove elements keeping the order of the elements.

Definition 1. Let L be a sorted list of l membership degrees in $[0, 1]$ and n be a natural number. Then $L \times n$ is a sorted list of $l \times n$ membership degrees adding every membership degree n times.

Example: $\{0.4, 0.8, 1\} \times 3 = \{0.4, 0.4, 0.4, 0.8, 0.8, 0.8, 1, 1, 1\}$

Definition 2. Let L be a sorted list of l membership degrees in $[0, 1]$, then $L[i]$ is the i th element of the list L , for $i = 1..l$.

Definition 3. Let L_1 and L_2 be lists of length l_1 and l_2 . Then $L_1 \leq_H L_2$ if and only if $L_1 \times (l/l_1)[i] \leq L_2 \times (l/l_2)[i]$ for all $i = 1..mcm(l_1, l_2) = l$. $\leq_H$ is a partial order on the set of Lists.

Definition 4. Let T be a triangular norm. The intersection T_H of 2 lists L_1 and L_2 of length l_1 and l_2 is a list $T_H(L_1, L_2)$ of length $l = mcm(l_1, l_2)$ such that $T_H(L_1, L_2)[i] = T(L_1 \times (l/l_1)[i], L_2 \times (l/l_2)[i])$ for all $i = 1..l$

Hesitant fuzzy sets and hesitant fuzzy sets bags

Let E be a finite universe of discourse.

Definition 5. A hesitant fuzzy set μ_H on a universe E is a mapping $\mu_H: E \rightarrow H$, where H is the set of sorted lists of membership degrees in $[0, 1]$.

Let l_i be the length of the list $L_i = \mu_H(e_i)$. So l_i is the number of membership degrees of e_i .

Let l_H be the mcm of all l_i for all i in $1..\text{card}(E)$.

Definition 6. A hesitant fuzzy set bag B_μ of a hesitant fuzzy set μ_H on a universe E is a mapping

$B_\mu: E \rightarrow L_B$ where L_B is the set of sorted lists of length l_H and $B_\mu(e_i) = L_i \times (l_H/l_i)$

Example: Let $E = \{e_1, e_2\}$, Let $\mu_H = \{\{0.4, 0.6\}, \{0.5, 0.7, 1\}\}$

Then $B_\mu = \{\{0.4, 0.4, 0.4, 0.6, 0.6, 0.6\}, \{0.5, 0.5, 0.7, 0.7, 1, 1\}\}$

Note that $\{0.4, 0.6\} \leq_H \{0.5, 0.7, 1\}$ because $\{0.4, 0.4, 0.4, 0.6, 0.6, 0.6\} \leq \{0.5, 0.5, 0.7, 0.7, 1, 1\}$.

Hesitant fuzzy relations and hesitant fuzzy relation bags

Definition 7. A hesitant fuzzy relation R_H on a universe E is a mapping $R_H: E \times E \rightarrow H$ where H is the set of a sorted lists of membership degrees in $[0, 1]$.

Let l_{ij} be the length of the list $L_{ij} = R_H(e_i, e_j)$.

Let l_R be the mcm of all l_{ij} for all i, j in $1.\text{card}(E)$.

Definition 8. A hesitant fuzzy relation bag B_R of a hesitant fuzzy relation R_H is a mapping $B_R: E \times E \rightarrow L_B$ where L_B is the set of a sorted lists of length l_R and $B_R(e_i, e_j) = L_{ij} \times (l_R/l_{ij})$

T-Transitive property and T-transitive closure of Hesitant fuzzy relations

Definition 9. Let T be a triangular norm. A hesitant fuzzy relation R_H on a universe E is T -transitive if and only if $T_H(R_H(a, b), R_H(b, c)) \leq_H R_H(a, c)$ for all a, b, c in E .

Definition 10. Let T be a triangular norm. Let R_H be hesitant fuzzy relation on a universe E . The T transitive closure of R_H is the $\leq_H$ -lowest relation R_H^T that is T -transitive and $\leq_H$ -contains R_H

Theorem 1. The T -transitive closure of a hesitant fuzzy relation always exists and it is unique.

Lemma 1. The T -transitive closure of a hesitant is computable by a finite algorithm.

Example: Let $E = \{e_1, e_2\}$, Let $R_H = \begin{pmatrix} 1 & \{0.5, 0.7, 1\} \\ \{0.4, 0.8\} & 0 \end{pmatrix}$

Then $B_R = \begin{pmatrix} \{1, 1, 1, 1, 1\} & \{0.5, 0.5, 0.7, 0.7, 1, 1\} \\ \{0.4, 0.4, 0.4, 0.8, 0.8, 0.8\} & \{0, 0, 0, 0, 0, 0\} \end{pmatrix}$ and $l_R = 6$.

Then R_H^{Min} , R_H^{Prod} , R_H^W (closures for the t-norms min, product and Łukasiewicz) are

$$R_H^{\text{Min}} = \begin{pmatrix} \{1, 1, 1, 1, 1\} & \{0.5, 0.5, 0.7, 0.7, 1, 1\} \\ \{0.4, 0.4, 0.4, 0.8, 0.8, 0.8\} & \{0.4, 0.4, 0.4, 0.7, 0.8, 0.8\} \end{pmatrix}$$

$$R_H^{\text{Prod}} = \begin{pmatrix} \{1, 1, 1, 1, 1\} & \{0.5, 0.5, 0.7, 0.7, 1, 1\} \\ \{0.4, 0.4, 0.4, 0.8, 0.8, 0.8\} & \{0.2, 0.2, 0.28, 0.56, 0.8, 0.8\} \end{pmatrix}$$

$$R_H^W = \begin{pmatrix} \{1, 1, 1, 1, 1\} & \{0.5, 0.5, 0.7, 0.7, 1, 1\} \\ \{0.4, 0.4, 0.4, 0.6, 0.6, 0.6\} & \{0, 0, 0.1, 0.5, 0.8, 0.8\} \end{pmatrix}$$

Forcing that all lists are comparable is a strong condition. A weaker definition is possible:

Definition 11. Let T be a triangular norm. A hesitant fuzzy relation R_H on a universe E is weak T -transitive if and only if $T_H(R_H(a, b), R_H(b, c)) \leq_H R_H(a, c)$ or $T_H(R_H(a, b), R_H(b, c))$ is not comparable to $R_H(a, c)$ for all a, b, c in E .

Theorem 2. The weak T -transitive closure of a hesitant fuzzy relation does not exist.

References

- [1] J. Recasens. *Indistinguishability Operators. Modelling Fuzzy Equalities and Fuzzy Equivalence Relations*. Studies in Fuzziness and Soft Computing. Springer. 2011.
- [2] V. Torra. *Hesitant fuzzy sets*. International Journal of Intelligent Systems **25-6** (2010) 529–539.
- [3] R. R.Yager. *On the theory of bags*. International Journal of General Systems, **13-1** (1986) 23–37.

Agregaciones computables

Ramón González del Campo¹, Luis Garmendia, Javier Montero, J. Tinguaro Rodríguez, Daniel Gómez

Universidad Complutense de Madrid

rgonzale@estad.ucm.es, lgarmend@fdi.ucm.es, monty@mat.ucm.es,
jtrodrig@mat.ucm.es, dagomez@estad.ucm.es

En este artículo se define el concepto de agregación computable como una alternativa a la definición clásica de agregación [2, 3, 4], que puede naturalmente extenderse a listas [1] de la siguiente forma:

Definición 1. Una agregación generalizada sobre un retículo completo de elementos de tipo T es una función $Ag : T^n \rightarrow T$ que satisface:

$$1. \ Ag(\underbrace{0_{\mathcal{T}}, 0_{\mathcal{T}}, \dots, 0_{\mathcal{T}}}_{n \text{ times}}) = 0_{\mathcal{T}} \text{ y } Ag(\underbrace{1_{\mathcal{T}}, 1_{\mathcal{T}}, \dots, 1_{\mathcal{T}}}_{n \text{ times}}) = 1_{\mathcal{T}}$$

donde $0_{\mathcal{T}}$ y $1_{\mathcal{T}}$ son los menores y mayores elementos del retículo $\mathcal{T}$ respectivamente.

2. Ag es monótona.

Definición 2. Sea $L < T >$ una lista de elementos de tipo T . Una agregación computable es una agregación generalizada $Ag : L < T > \rightarrow T$ tal que $Ag(L) = P(L)$ donde P es un programa que verifica las propiedades de las agregaciones generalizadas.

Ejemplo 1. La media aritmética es una agregación computable.

Sea $L = \{x_1, \dots, x_n\}$ una lista de valores $[0, 1]$. El siguiente programa escrito en C++ calcula la expresión $\frac{1}{n} \sum_{i=1}^n x_i$:

```
float arithmetic_mean(float L[], int n){
    float acum;
    acum=0;
    for(i=0; i++; i<n)
        acum=acum+L[i];
    return acum/n;
}
```

En muchas ocasiones la implementación más inmediata de una agregación nos remite a algoritmos de un determinado paradigma de programación. En este sentido podemos hablar de agregaciones procedurales, funcionales, lógicas, etc.

Definición 3. Sea $L = \{x_1, \dots, x_n\}$ una lista de valores $[0, 1]$. Una agregación computable es procedural $Ag : L < T > \rightarrow T$ si existe un programa procedural P tal que $Ag(L) = P(L)$ y Ag es una agregación generalizada.

Agradecimientos: este artículo ha sido posible debido al apoyo de los proyectos TIN2015-66471-P y S2013/ICE-2845.

De forma similar se definen las agregaciones procedurales, funcionales y lógicas. Así, por ejemplo:

- El programa del ejemplo 1 es un programa procedural que implementa la media aritmética.
- La media aritmética es una agregación funcional: si L es una lista de valores en $[0, 1]$, el siguiente programa escrito en Haskell calcula la media aritmética de n elementos:

```
suma [] = 0
suma (l:ls) = l + suma ls
long [] = 0
long (l:ls) = 1 + long ls
media l = (suma l) / (long l)
```

Definición 4. Una agregación map-reducible es una agregación computable que puede ser programada utilizando un algoritmo del paradigma de programación map-reduce.

Así,

- La media aritmética es una agregación map-reducible ya que la suma de los valores puede ser descompuesta y procesada en varios nodos.
- La mediana no es una agregación map-reducible, pues su cómputo no permite su descomposición y procesamiento en varios nodos.

Definición 5. Una agregación computable Ag es abordable con una complejidad $\Theta(t(n))$ si existe un programa P con complejidad $\Theta(t(n))$ y $Ag(L) = P(L)$ donde $t(n)$ es el tiempo de cómputo de $P(L)$ y Ag es una agregación generalizada.

Así, por ejemplo,

- La media aritmética es abordable con una complejidad lineal.
- La media geométrica es abordable con una complejidad lineal.
- Todo operador F de tipo OWA es abordable con una complejidad $n * \log(n)$.

Referencias

- [1] G. Barnett and L. Del Tonga. *Data Structures and Algorithms*. DotNetSlackers, 2008.
- [2] G. Beliakov, A. Pradera, and T. Calvo. *Aggregations Functions: A guide for Practitioners*. Springer, 2007.
- [3] D. Gómez and J. Montero. A discussion on aggregations operators. *Kybernetika*, 40:107–120, 2004.
- [4] R. González del Campo, L. Garmendia, and J. Montero. Expansible computable aggregation rules. In *Proceedings of the 2015 International Conference on Intelligent Systems and Knowledge Engineering, ISKE 2015, Taipei, Taiwan*, pages 8–11, November 2015.

On some properties of SUOWA operators

Bonifacio Llamazares

*Departamento de Economía Aplicada, Instituto de Matemáticas (IMUVA), and BORDA
Research Unit, Universidad de Valladolid, Spain.
boni@eco.uva.es*

In the literature of aggregation functions, weighted means and OWA operators (Yager [9]) play a major role in aggregating information. Although both families of functions are defined by means of weighting vectors, their behavior is quite different. This is because in the case of weighted means, the weighting vectors weight each information source in relation to their reliability, while in the case of OWA operators they weight the values according to their ordering.

Despite the fact that both families of functions allow to solve a wide range of problems, there are situations where both the importance of information sources and the importance of values had to be taken into account. Different aggregation operators have appeared in the literature to deal with this kind of problems. A usual approach is to consider families of functions parametrized by two weighting vectors, one for the weighted mean and the other one for the OWA type aggregation, that generalize weighted means and OWA operators in the following sense: A weighted mean (or a OWA operator) is obtained when the other weighting vector has a “neutral” behavior; that is, it is $(1/n, \dots, 1/n)$ (see Llamazares [3] for an analysis of some functions that generalize the weighted means and the OWA operators in this sense). Two of the solutions having better properties are the weighted OWA (WOWA) operator, proposed by Torra [8], and the semi-uniform based ordered weighted averaging (SUOWA) operator, introduced by Llamazares [4].

WOWA and SUOWA operators are continuous, monotonic, idempotent, compensative and homogeneous of degree 1 functions. These good properties are due to the fact that both families of functions are Choquet integrals with respect to normalized capacities (on the Choquet integral, see Choquet [1]). In the case of SUOWA operators, their capacities are the monotonic cover of certain games, which are defined by using the capacities associated with the weighted means and the OWA operators and “assembling” these values through semi-uniforms with neutral element $1/n$ (we recall that semi-uniforms, introduced by Liu [2], are an extension of uniformity by dispensing with the symmetry and associativity properties).

Given that semi-uniforms play a fundamental role in the definition of SUOWA operators, two procedures to construct semi-uniforms have been introduced by Llamazares [6]. The first one, based on the notion of ordinal sums of aggregation operators, allows us to obtain continuous semi-uniforms. The second one, based in the combination of several semi-uniforms, allows us to construct new semi-uniforms that keep some properties of the former semi-uniforms.

This work is partially supported by the Spanish Ministry of Economy and Competitiveness (Project ECO2012-32178) and the Junta de Castilla y León (Consejería de Educación, Project VA066U13).

Although WOWA and SUOWA operators share many properties, they are different classes of aggregation operators as it has been pointed out by Llamazares [4]. For this reason, it seems interesting to analyze their behavior from other points of view. In this sense, Llamazares [7] has studied WOWA and SUOWA operators regarding some simple cases of weighting vectors, the capacities from which they are building, the weights affecting the components of each vector and the values they return. Likewise, Llamazares [5] has proven that, in the case of two dimensions, any Choquet integral with respect to a normalized capacity can be expressed as a SUOWA operator.

Another interesting issue is to know the behavior of these functions by using several indices suggested in the literature: orness and andness degrees, importance and interaction indices, tolerance indices, dispersion indices, etc. The aim of this work is to analyze some specific cases of SUOWA operators with respect to various indices such as the orness, which measures the degree to which the aggregation is like an *or* operation; the Shapley value, which shows the global importance of each information source; and the veto and favor indices, which measure the degree to which an information source behaves like a veto or a favor.

References

- [1] G. Choquet: Theory of capacities. *Annales de l'Institut Fourier* **5** (1953) 131–295.
- [2] H.-W. Liu: Semi-uninorms and implications on a complete lattice. *Fuzzy Sets and Systems* **191** (2012) 72–82.
- [3] B. Llamazares: An analysis of some functions that generalizes weighted means and OWA operators. *International Journal of Intelligent Systems* **28** (2013) 380–393.
- [4] B. Llamazares: Constructing Choquet integral-based operators that generalize weighted means and OWA operators. *Information Fusion* **23** (2015) 131–138.
- [5] B. Llamazares: A study of SUOWA operators in two dimensions. *Mathematical Problems in Engineering* **2015**, Article ID 271491 (2015) 12 pages.
- [6] B. Llamazares: SUOWA operators: Constructing semi-uninorms and analyzing specific cases. *Fuzzy Sets and Systems* **287** (2016) 119–136.
- [7] B. Llamazares: A behavioral analysis of WOWA and SUOWA operators. *International Journal of Intelligent Systems* (forthcoming).
- [8] V. Torra: The weighted OWA operator. *International Journal of Intelligent Systems* **12** (1997) 153–166.
- [9] R. R. Yager: On ordered weighted averaging operators in multicriteria decision making. *IEEE Transactions on Systems, Man, and Cybernetics* **18** (1988) 183–190.

OWAs defining Multidistances and Dispersion Measures on $\mathbb{R}^k$

Javier Martín, Gaspar Mayor

*University of the Balearic Islands, Spain,
Cra. Valldemossa km 7.5, 07122 Palma, Spain
javier.martin@uib.es, gmayor@uib.es*

We consider in this work a class of both multidistances and dispersion measures defined by means of a family of OWA operators, on the set of lists of real numbers $\bigcup_{n \geq 1} (\mathbb{R}^k)^n$ equipped with the Euclidean distance.

Multidistances are a generalization of ordinary distances in order to measure how much separated are not only two elements of a set but any finite list. They are defined as follows.

Definition 1 [2] *We say that a function $D: \bigcup_{n \geq 1} X^n \rightarrow \mathbb{R}^+$ is a multidistance on a non empty set X when the following properties hold, for all n and $x_1, \dots, x_n, y \in X$:*

md1) $D(x_1, \dots, x_n) = 0$ if and only if $x_i = x_j$ for all $i, j = 1, \dots, n$.

md2) $D(x_1, \dots, x_n) = D(x_{\pi(1)}, \dots, x_{\pi(n)})$ for any permutation π of $1, \dots, n$,

md3) $D(x_1, \dots, x_n) \leq D(x_1, y) + \dots + D(x_n, y)$,

A remarkable class of these functions are the so called functionally expressible multidistances.

Definition 2 [3] *Let D be a multidistance on a set X and d an ordinary distance on the same set. We will say that D is functionally expressible from d and F , or (d, F) -functionally expressible, if there exist a symmetric function $F: \bigcup_{n \geq 1} (\mathbb{R}^+)^n \rightarrow \mathbb{R}^+$ such that:*

$$D(x_1, \dots, x_n) = F(d(x_1, x_2), \dots, d(x_i, x_j), \dots, d(x_{n-1}, x_n)) , \quad (1)$$

for all $n \geq 2$, $1 \leq i < j \leq n$ and $x_1, \dots, x_n \in X$.

We will deal with a definition of measures of dispersion based on the one given in [1], slightly modified.

Definition 3 *A function $\Delta: \bigcup_{n \geq 1} (\mathbb{R}^k)^n \rightarrow \mathbb{R}^+$ is said to be a c -dispersion measure, $c > 0$, when it fulfills these four conditions:*

$\Delta 1)$ $\Delta(x_1, \dots, x_n) = 0$ if and only if $x_i = x_j$ for all $i, j = 1, \dots, n$.

This work has been supported by the projects TIN2013-42795-P and TIN2014-56381-REDT (LODISCO) (Spanish Government)

$\Delta 2)$ Δ is symmetric: $\Delta(x_1, \dots, x_n) = \Delta(x_{\pi(1)}, \dots, x_{\pi(n)})$ for any permutation π of $\{1, \dots, n\}$.

$\Delta 3)$ Δ is invariant for isometries: $\Delta(x_1, \dots, x_n) = \Delta(\phi(x_1), \dots, \phi(x_n))$ for any isometry ϕ defined on $\mathbb{R}^k$.

$\Delta 4)$ $\Delta(ax_1, \dots, ax_n) = a^c \Delta(x_1, \dots, x_n)$, for all $a \in \mathbb{R}^+$.

Comparing Defs. 3 and 1 (with $X = \mathbb{R}^k$ and d the Euclidean distance), it can be observed that conditions $\Delta 1$, $\Delta 2$ are the same as $\text{md}1$, $\text{md}2$. Moreover, condition $\Delta 3$ holds for functionally expressible multidistances, and an additional condition can be given in order to fulfill $\Delta 4$.

Proposition 1 Let D be a (d, F) -functionally expressible multidistance on $\mathbb{R}^k$, where F satisfies this additional condition $F(at_1, \dots, at_n) = a^c F(t_1, \dots, t_n)$, with $c > 0$, for all $n \geq 1$, $t_1, \dots, t_n$ and $a \geq 0$. Then, D is a dispersion measure.

Let $W = (W_n)_{n \geq 1}$ be a family of OWAs. The weights of the OWA W_n will be denoted by $\omega_1^n, \dots, \omega_n^n$.

A function $D_W: \bigcup_{n \geq 1} (\mathbb{R}^k)^n \rightarrow \mathbb{R}^+$ can be defined in this way:

$$D_W(x_1, \dots, x_n) = \begin{cases} 0 & \text{if } n = 1, \\ W_n(\overbrace{d(x_1, x_2), \dots, d(x_{n-1}, x_n)}^{(n)_2}) & \text{if } n \geq 2, \end{cases} \quad (2)$$

for any list $(x_1, \dots, x_n) \in (\mathbb{R}^k)^n$.

The following result gives a necessary and sufficient condition for this kind of functions D_W to be multidistances.

Proposition 2 A function D_W , defined by the expression 2, is a multidistance if and only if the weights of the OWAs of the family W fulfill this condition:

$$\omega_1^{(n)_2} + \dots + \omega_{n-1}^{(n)_2} > 0, \quad \forall n \geq 2. \quad (3)$$

Functions D_W fulfilling 3 are called OWA-based multidistances. They are functionally expressible and also condition $\Delta 4$ holds for $c = 1$. Therefore, we can establish the main result of this work.

Proposition 3 If D_W is an OWA-based multidistance on $\mathbb{R}^k$, then it is also a 1-dispersion measure.

References

- [1] A. Kołacz, P. Grzegorzewski: Measures of Dispersion for Multidimensional Data. European Journal of Operational Research (2016), doi: 10.1016/j.ejor.2016.01.011.
- [2] J. Martín, G. Mayor: Multi-argument Distances. *Fuzzy Set. Syst.* **167**, 92–100 (2011)
- [3] J. Martín, G. Mayor, O. Valero: Functionally Expressible Multidistances. Proceedings of EUSFLAT–LFA 2011.

Relating absolute dispersion measures, multidistances and remainders of aggregation functions

Miguel Martínez-Panero, José Luis García-Lapresta, Luis Carlos Meneses

*PRESAD Research Group, IMUVA, Dept. de Economía Aplicada,
Universidad de Valladolid, Spain*

{panero, lapresta, lmeneses}@eco.uva.es

In this paper, we deal with several ways of formalizing the closeness among different mathematical objects (usually points in $\mathbb{R}^n$), such as dispersion measures, multidistances and remainders of some aggregation functions.

First, we propose a notion of absolute dispersion measure that is satisfied by some classical dispersion measures used in Descriptive Statistics. Before that, we introduce some notation.

Notation Let I be $[0, 1]$ or $\mathbb{R}$, and $\mathbb{I} = \bigcup_{n \in \mathbb{N}} I^n$. Vectors in I^n are denoted as $\mathbf{x} = (x_1, \dots, x_n)$, $\mathbf{0} = (0, \dots, 0)$ and $\mathbf{1} = (1, \dots, 1)$. Accordingly, $x \cdot \mathbf{1} = (x, \dots, x)$ for every $x \in I$.

Definition 1 A function $D : \mathbb{I} \rightarrow \mathbb{R}$ is an absolute dispersion measure if it satisfies the following conditions: (1) Positiveness ($D(\mathbf{x}) \geq 0$, for every $\mathbf{x} \in \mathbb{I}$); (2) Identity of indiscernibles ($D(\mathbf{x}) = 0 \Leftrightarrow x_1 = \dots = x_n$, for all $n \in \mathbb{N}$ and $\mathbf{x} \in I^n$); (3) Symmetry ($D(x_{\pi(1)}, \dots, x_{\pi(n)}) = D(\mathbf{x})$, for all $n \in \mathbb{N}$, $\mathbf{x} \in I^n$ and permutation π on $\{1, \dots, n\}$); (4) Invariance for translations (for all $\mathbf{x} \in I^n$ and $t \in \mathbb{R}$ such that $\mathbf{x} + t \cdot \mathbf{1} \in I^n$ it holds $D(\mathbf{x} + t \cdot \mathbf{1}) = D(\mathbf{x})$); (5) Invariance for replications (if $D(\overbrace{\mathbf{x}, \dots, \mathbf{x}}^m) = D(\mathbf{x})$ for every $\mathbf{x} \in \mathbb{I}$ and any number $m \in \mathbb{N}$ of replications of $\mathbf{x}$); and (6) Anti-self-duality (if $I = [0, 1]$ and for every $\mathbf{x} \in [0, 1]^n$ it holds $D(\mathbf{1} - \mathbf{x}) = D(\mathbf{x})$ (evenness if $I = \mathbb{R}$)).

Some classical measures commonly appearing in the literature, as the range, variance, standard deviation and mean deviation, as well as the absolute Gini index (the most popular measure of inequality in welfare economics) verify the above mentioned properties, becoming absolute dispersion measures. But others are not properly absolute dispersion measures, as the coefficient of variation (it vulnerates invariance for translations and anti-self-duality).

On the other hand, Martín and Mayor [2] have proposed the notion of multidistance generalizing the classical definition of distance between two points to any finite number of points by extending the usual triangle inequality.

Definition 2 A function $M : \mathbb{I} \rightarrow \mathbb{R}$ is a multidistance if it satisfies the following conditions: (1) Positiveness; (2) Identity of indiscernibles; (3) Symmetry; and (4) Generalized triangle inequality ($M(\mathbf{x}) \leq M(x_1, y) + \dots + M(x_n, y)$, for all $n \in \mathbb{N}$, $\mathbf{x} \in I^n$ and $y \in I$).

The authors gratefully acknowledge the comments of Javier Martín, Ana Pérez and Mercedes Prieto, as well as the funding support of the Spanish *Ministerio de Economía y Competitividad* (project ECO2012-32178) and *Consejería de Educación de la Junta de Castilla y León* (project VA066U13).

Some classes of multidistances introduced by Martín and Mayor [2] are:

- The *drastic multidistances*. Among them, the one defined as $D(\mathbf{x}) = \#\{x_1, \dots, x_n\} - 1$, trivially fulfills all the conditions of absolute dispersion measures.
- The *sum-based multidistances*. It is straightforward that the absolute Gini index is just twice a multidistance of this family. However, not all sum-based multidistances are absolute dispersion measures, because invariance of replications fails.
- The *OWA-based multidistances*. Among them, the maximum multidistance, that in our case coincides with the range, is an absolute dispersion measure. However, not all OWA-based multidistance are absolute dispersion measures (also invariance of replications fails).

And also introduced by Martín and Mayor [2], the *Fermat multidistance* $D_F : \mathbb{I} \longrightarrow \mathbb{R}$ is given by: $D_F(\mathbf{x}) = \min_{x \in I} \left\{ \sum_{i=1}^n |x_i - x| \right\}$ where $\mathbf{x} \in I^n$.

Proposition 1 *The Fermat multidistance is an absolute dispersion measure. Moreover, it is an iterated sum of ranges of the data, where in each term the extreme values are sequentially withdrawn.*

Thus, the Fermat multidistance inherits its condition of absolute dispersion measure from the range. And its is also true a sort of reciprocal: the above mentioned fact that the range is also a multidistance.

Finally, we take into account the dual decomposition of aggregation functions introduced by García-Lapresta and Marques Pereira [1]. Concretely, in this paper we focus our attention on remainders of exponentials means.

Definition 3 *Given an aggregation function $A : [0, 1]^n \longrightarrow [0, 1]$, the function $\tilde{A} : [0, 1]^n \longrightarrow \mathbb{R}$ defined as $\tilde{A}(\mathbf{x}) = \frac{A(\mathbf{x}) + A(\mathbf{1} - \mathbf{x}) - 1}{2}$ is called the remainder of A .*

Definition 4 *The exponential mean $A_\alpha : [0, 1]^n \longrightarrow [0, 1]$ is the aggregation function defined as $A_\alpha(\mathbf{x}) = \frac{1}{\alpha} \ln \frac{e^{\alpha x_1} + \dots + e^{\alpha x_n}}{n}$, where $\alpha \neq 0$.*

Proposition 2 *$\tilde{A}_\alpha$ is an absolute dispersion measure for every $\alpha > 0$.*

If $\alpha < 0$ then $\tilde{A}_{-\alpha}(\mathbf{x}) = -\tilde{A}_\alpha(\mathbf{x})$, and hence positiveness (required both for absolute dispersion measures and multidistances) is vulnerated. And even for $\alpha > 0$, $\tilde{A}_\alpha$ is not an multidistance (the triangle inequality does not hold).

As further research, a similar analysis with other aggregation functions such as quasarithmetic means or OWA operators can be developed.

References

- [1] García-Lapresta, J.L., Marques Pereira, R.A.: *The self-dual core and the anti-self-dual remainder of an aggregation operator*. Fuzzy Sets and Systems. **159** (2008) 47–62.
- [2] Martín, J., Mayor, G.: *Multi-argument distances*. Fuzzy Sets and Systems. **167** (2011) 92–100.

Estudio del Modus Ponens respecto de una uninorma

M. Mas, M. Monserrat, D. Ruiz-Aguilera, J. Torrens

Universitat de les Illes Balears

Departamento de Ciencias Matemáticas e Informática

Crta. de Valldemossa, km. 7,5, 07122 Palma

{mmg448, mma112, daniel.ruiz, jts224}@uib.es

Una de las aplicaciones más importantes de las funciones de implicación borrosa está centrada en el control borroso y el razonamiento aproximado, donde estos operadores son usados para modelar los condicionales borrosos así como en los procesos de inferencia, en los que la regla del Modus Ponens resulta esencial [1]. En estos casos, el Modus Ponens es aplicable cuando los operadores lógicos utilizados verifican la propiedad:

$$T(x, I(x, y)) \leq y \quad \text{para todos } x, y \in [0, 1], \quad (1)$$

donde T es habitualmente una t -norma (continua) e I una función de implicación borrosa. Por este motivo la propiedad del Modus Ponens es también conocida como T -condicionalidad.

La ecuación (1) ha sido extensamente estudiada para las clases de implicaciones más habituales, como las R -implicaciones, las (S, N) -implicaciones y las QL y D -implicaciones [1, 5, 6]. Recientemente, los autores la han estudiado también para las RU y (U, N) -implicaciones derivadas de uninormas [3].

En este trabajo pretendemos generalizar esta propiedad substituyendo la t -norma T por una uninorma U , introduciendo así la llamada U -condicionalidad. De esta forma, estudiamos la U -condicionalidad para una uninorma U cualquiera, pero vemos de inmediato que U ha de ser necesariamente conjuntiva. Vemos además que cualquier función de implicación borrosa, I , que sea un U -condicional ha de verificar determinadas propiedades, como por ejemplo que $I(1, y) < e$ para todo $y < 1$, siendo e el neutro de la uninorma. En particular, esto demuestra que las implicaciones más usuales, como las R , (S, N) , QL y D -implicaciones derivadas de t -normas y t -conormas, así como las implicaciones f y g generadas de Yager, no pueden ser U -condicionales.

Por el contrario, las implicaciones derivadas de uninormas sí son buenas candidatas a ser U -condicionales, así como las h -implicaciones (implicaciones construidas a partir de generadores de uninormas [4]). En este trabajo, nos limitamos a estudiar el caso de las RU -implicaciones, es decir, las implicaciones residuadas derivadas de uninormas. En concreto, lo hacemos en los casos en que la uninorma utilizada para construir la RU -implicación es una uninorma de $\mathcal{U}_{\min}$ o una uninorma idempotente (para las diversas clases de uninormas ver por ejemplo el reciente survey [2]).

El primer resultado que se ve, sin embargo, es general y viene dado en la siguiente proposición.

Este trabajo ha sido subvencionado por el proyecto TIN2013-42795-P.

Proposición 1 Sean U, U_0 dos uninormas con elementos neutros $e, e_0 \in]0, 1[$ respectivamente de forma que U es conjuntiva y U_0 verifica $U_0(0, y) = 0$ para todo $y < 1$. Sea I_{U_0} la implicación residuada derivada de U_0 . Si I_{U_0} es un U -condicional entonces ha de ser $e_0 \leq e$.

A continuación, para el caso en que U_0 es una uninorma de $\mathcal{U}_{\min}$, se ve que en realidad en este caso ha de ser $e_0 < e$ y la uninorma U tiene que venir dada por el mínimo en la región:

$$R_0 = [0, e_0] \times [e, 1] \cup [e, 1] \times [0, e_0]. \quad (2)$$

Esta condición necesaria, se ve que también es suficiente en el caso en que la t-conorma asociada a la uninorma U es el máximo y la t-norma asociada a U_0 es el mínimo. Por último, damos una caracterización completa en el caso general en que $U(e_0, e_0) = e_0$.

Por lo que respecta a las RU -implicaciones derivadas de uninormas idempotentes, $U_0 \equiv \langle g_0, e_0 \rangle_{\text{ide}}$ con $g_0(0) = 1$, se distinguen dos casos diferentes según que $g_0(e)$ sea igual a 0 o distinto de 0. En ambos casos, se da una caracterización completa de las RU -implicaciones I_{U_0} con U_0 idempotente, que son U -condicionales. A lo largo del trabajo, se van dando ejemplos ilustrativos de los resultados obtenidos.

Como trabajo futuro, pretendemos extender este estudio a las RU -implicaciones derivadas de otros tipos de uninormas, como las representables, las continuas en el cuadrado unidad abierto, las compensatorias, etc (ver [2]). Más aún, queremos tratar también con otros tipos de implicaciones como las (U, N) -implicaciones derivadas de uninormas disyuntivas ([1]) o las h y (h, e) -implicaciones recientemente introducidas en [4]. Finalmente, también valdría la pena estudiar una generalización del Modus Tollens a través de uninormas, similar a la estudiada aquí para el Modus Ponens.

Referencias

- [1] M. Baczyński, B. Jayaram: *Fuzzy Implications*. Studies in Fuzziness and Soft Computing, vol. 231. Springer, Berlin Heidelberg, (2008).
- [2] M. Mas, S. Massanet, D. Ruiz-Aguilera, J. Torrens: *A survey on the existing classes of uninorms*. Journal of Intelligent and Fuzzy Systems. **29** (2015) 1021–1037.
- [3] M. Mas, M. Monserrat, D. Ruiz-Aguilera, J. Torrens: *RU and (U, N) -implications satisfying Modus Ponens*. International Journal of Approximate Reasoning. (2016) <http://dx.doi.org/10.1016/j.ijar.2016.01.003>.
- [4] S. Massanet, J. Torrens: *On a new class of fuzzy implications: h -implications and generalizations*. Information Sciences. **181** (2011) 2111–2127.
- [5] E. Trillas, C. Alsina, A. Pradera: *On MPT-implication functions for fuzzy logic*. Revista de la Real Academia de Ciencias. Serie A. Matemáticas (RACSAM). **98(1)** (2004) 259–271.
- [6] E. Trillas, C. Alsina, E. Renedo, A. Pradera: *On contra-symmetry and MPT-conditionality in fuzzy logic*. International Journal of Intelligent Systems. **20** (2005) 313–326.

Sobre Implicaciones Racionales

Sebastia Massanet, Juan Vicente Riera, Daniel Ruiz-Aguilera

Universitat de les Illes Balears, Dept. Matemàtiques i Informàtica

Crta. Valldemossa km 7.5, 07122 Palma, Illes Balears

{s.massanet,jvicente.riera,daniel.ruiz}@uib.es

En estos últimos años, el estudio de las funciones de implicación borrosas se ha convertido en una de las líneas principales de investigación en lógica borrosa [1, 2]. En particular, la búsqueda de nuevas familias de implicaciones borrosas cuyas propiedades puedan responder a las necesidades específicas de cada aplicación se ha convertido en una constante. Habitualmente, estas nuevas familias han sido construidas mediante funciones de agregación, negaciones borrosas o generadores aditivos. Sin embargo, debido a que esta estrategia no tiene en cuenta la expresión final del operador, se pueden obtener funciones de implicación borrosas con expresiones complejas, que dificultan su uso en cualquier aplicación. En respuesta a esta consecuencia indeseada, se inició una línea de investigación centrada en la obtención de familias de implicaciones borrosas cuya expresión fuera simple y adecuada para su aplicación. Así, en [3] y [4], se introdujeron y estudiaron las implicaciones polinomiales y (OP)-polinomiales, aquéllas cuya expresión viene dada por un polinomio de dos variables. En este trabajo, se continúa en esta dirección y se introducen y estudian las implicaciones racionales, como aquéllas cuya expresión viene dada por el cociente de dos polinomios de dos variables.

Definición 1. Un operador $I : [0, 1]^2 \rightarrow [0, 1]$ es una *implicación racional borrosa de grado (n, m)* si es una función de implicación borrosa y su expresión viene dada por $I(x, y) = \frac{p(x, y)}{q(x, y)}$ donde p y q son polinomios de grados n y m respectivamente, sin factores comunes.

Cabe destacar que las implicaciones polinómicas, introducidas en [3], son un caso particular de las implicaciones racionales, coincidiendo con las de grado $(n, 0)$ con $n \geq 0$. Por otra parte, las implicaciones racionales siempre son continuas y por lo tanto, si satisfacen el principio de intercambio (EP), son de hecho (S,N)-implicaciones generadas a partir de una negación borrosa estricta N y por la t-conorma $S(x, y) = I(N^{-1}(x), y)$. También satisfacen la propiedad $I(x, y) = 1$ si, y sólo si, $x = 0$ e $y = 1$, siendo útiles para generar índices de equivalencia fuertes. Sin embargo, no verifican en ningún caso la propiedad de orden (OP) ni el principio de identidad (IP), por lo que nunca son R-implicaciones.

La caracterización de la familia de implicaciones racionales consiste en establecer condiciones en los coeficientes de los polinomios p y q para que la función $\frac{p(x, y)}{q(x, y)}$ sea de hecho una función de implicación borrosa. Aunque las condiciones de frontera permiten establecer algunas condiciones sobre los coeficientes, las propiedades de monotonía no son fácilmente expresables en términos de los coeficientes y por lo tanto, la caracterización general de esta familia es compleja. A pesar de todo, sí es posible la caracterización de las implicaciones racionales de algunos grados concretos.

Este trabajo ha sido financiado por el proyecto TIN2013-42795-P concedido por el Ministerio de Economía y Competitividad.

Teorema 1. Sea $I : [0, 1]^2 \rightarrow [0, 1]$ un operador. Entonces I es una función racional de :

- grado (1,2) $\Leftrightarrow I(x, y) = \frac{\alpha(1-x)+\beta y}{\alpha+\gamma x+\beta y+(-\alpha-\gamma)xy}$ donde $\alpha, \beta > 0$ y $\gamma > -\alpha$.
- grado (2,1) $\Leftrightarrow I(x, y) = \frac{\alpha(1-x)+\beta y+(\alpha+\beta)xy}{\alpha+\beta(x+y)}$ donde $\alpha > 0$, $\beta \neq 0$ y $\beta > -\frac{\alpha}{2}$

Mientras que algunas implicaciones racionales de grado (1,2) son (S,N)-implicaciones y en ese caso, S es una t-conorma de Hamacher y N es una negación de Sugeno, ninguna implicación racional de grado (2,1) es (S,N)-implicación. Por otro lado, no existen implicaciones racionales de grado (0,m) con $m \geq 0$ ni de grado (1,1). Además, en [3], se demostró que tampoco hay de grado (1,0) y se dieron las caracterizaciones de las de grado (2,0) y (3,0).

Vista la dificultad existente en la caracterización de todas las implicaciones racionales, se proponen dos métodos de construcción de implicaciones racionales a partir de otras implicaciones racionales. En concreto, la conjugación con automorfismos φ con expresiones racionales y la N -reciprocidad con negaciones N de expresión racional son dos métodos válidos para este fin.

Proposición 2. Sea I una implicación racional. Entonces:

- i) Si $\varphi : [0, 1] \rightarrow [0, 1]$ es un automorfismo racional, entonces $I_\varphi(x, y) = \varphi^{-1}(\varphi(x), \varphi(y))$ es una implicación racional.
- ii) Si N es una negación borrosa racional, entonces $I_N(x, y) = I(N(y), N(x))$ es una implicación racional.

Ambos métodos pueden construir implicaciones racionales de diferente grado a la original y por lo tanto, las caracterizaciones dadas para grados (1,2) y (2,1) permiten la construcción a través de estos métodos de implicaciones de grados superiores. De hecho, se demuestra que existen implicaciones racionales de grado (n, n) con $n \geq 2$ cualquiera a partir de la construcción de la (S,N)-implicación generada a partir de una t-conorma de Hamacher y la familia de negaciones borrosas $N_n(x) = 1 - x^n$.

Referencias

- [1] M. Baczyński and B. Jayaram. *Fuzzy Implications*, volume 231 of *Studies in Fuzziness and Soft Computing*. Springer, Berlin Heidelberg, 2008.
- [2] M. Baczyński, B. Jayaram, S. Massanet, and J. Torrens. Fuzzy implications: Past, present, and future. In J. Kacprzyk and W. Pedrycz, editors, *Springer Handbook of Computational Intelligence*, pages 183–202. Springer Berlin Heidelberg, 2015.
- [3] S. Massanet, J. V. Riera, and D. Ruiz-Aguilera. On fuzzy polynomial implications. In A. Laurent et al., editor, *Information Processing and Management of Uncertainty in Knowledge-Based Systems*, volume 442 of *Communications in Computer and Information Science*, pages 138–147. Springer International Publishing, 2014.
- [4] S. Massanet, J. V. Riera, and D. Ruiz-Aguilera. On (OP)-polynomial implications. In J. M. Alonso et al., editor, *Proc. of the IFSA-EUSFLAT-15*. Atlantis Press, 2015.

Sumas ordinales de funciones de implicación borrosas

S. Massanet, J.V. Riera, J. Torrens

Universitat de les Illes Balears

Departamento de Ciencias Matemáticas e Informática

Crta. de Valldemossa, km. 7,5, 07122 Palma

{s.massanet, jvicente.riera, jts224}@uib.es

Desde un punto de vista teórico, los conectivos lógicos juegan un papel fundamental en la teoría de los conjuntos borrosos y de la lógica borrosa. Usualmente estos conectivos son interpretados de forma funcionalmente expresable y entonces se representan mediante funciones definidas en el intervalo unidad $[0, 1]$. De este modo, las conjunciones y disyunciones se describen mediante t-normas y t-conormas, los complementos mediante negaciones borrosas y los condicionales borrosos mediante las funciones de implicación borrosas (ver [2, 4]).

Las funciones de implicación borrosas son esenciales en el razonamiento aproximado y el control borroso porque no solo se usan para representar los condicionales borrosos, sino que también se usan en los procesos de inferencia. Además, tienen una gran cantidad de aplicaciones en diferentes campos que van desde el control borroso y las medidas de inclusión hasta la morfología matemática borrosa y el tratamiento de imágenes. Por este motivo el interés por las funciones de implicación borrosas ha ido creciendo paulatinamente en las últimas décadas como muestran los artículos recopilatorios [3, 5] y los libros [1, 2] plenamente dedicados a este tipo de operadores.

Una línea de investigación intensamente estudiada es la búsqueda de métodos de construcción de nuevas implicaciones borrosas, muchos de los cuales han sido recopilados en los capítulos 1 y 4 del libro [1]. Pero incluso desde entonces, nuevos métodos han ido surgiendo entre los que cabe destacar la suma ordinal de funciones de implicación borrosas en [6], donde dicha suma ordinal está basada en las implicaciones residuadas. En esta comunicación presentamos un nuevo enfoque de las sumas ordinales pero, en esta ocasión, basado en las implicaciones materiales o (S, N) -implicaciones derivadas de sumas ordinales de t-conormas. De hecho, presentamos diferentes construcciones de sumas ordinales, una por cada posible negación borrosa N , como sigue.

Definición 1 Sea J un conjunto finito o numerable de índices, $\{(a_j, b_j)\}_{j \in J}$ una familia de subintervalos abiertos disjuntos de $[0, 1]$ y N una negación de forma que $c_j = N(b_j)$, $d_j = N(a_j)$ y

$$N_{I_j}(x) = \frac{N(a_j + (b_j - a_j)x) - c_j}{d_j - c_j}, \quad \text{para todo } x \in [0, 1]. \quad (1)$$

Consideremos $\{I_j\}_{j \in J}$ una familia de funciones de implicación borrosas, llamaremos N -suma ordinal de la familia $\{I_j\}_{j \in J}$ a la función $I : [0, 1]^2 \rightarrow [0, 1]$ dada por

$$I(x, y) = \begin{cases} c_j + (d_j - c_j) \cdot I_j\left(\frac{x - a_j}{b_j - a_j}, \frac{y - c_j}{d_j - c_j}\right) & \text{si } x \in [a_j, b_j] \text{ e } y \in [c_j, d_j], \\ \max\{N(x), y\} & \text{en cualquier otro caso,} \end{cases}$$

Este trabajo ha sido subvencionado por el proyecto TIN2013-42795-P.

y la denotaremos por $I = \langle N, (a_j, b_j, I_j)_{j \in J} \rangle$.

Los primeros resultados que se ven para esta definición, es que $I = \langle N, (a_j, b_j, I_j)_{j \in J} \rangle$ es una función de implicación siempre que todas las I_j con $b_j < 1$ verifican la condición de frontera $I_j(x, y) \geq y$ para todo $y \in [0, 1]$ y que, dada una familia cualquiera de implicaciones $\{I_j\}_{j \in J}$, siempre se puede considerar una negación adecuada de manera que se pueda hacer la N -suma ordinal de dichas implicaciones.

El interés de estos métodos de construcción radica en que sean capaces de preservar las propiedades de las implicaciones originales. De esta forma, la comunicación continúa con el estudio de qué propiedades se preservan mediante la construcción de N -sumas ordinales o, si no se preservan en general, bajo qué condiciones sí lo hacen. En este sentido, se ve que la negación natural asociada a una implicación N -suma ordinal es siempre la propia N y se analizan con detalle la neutralidad por la izquierda, la condición de frontera, la continuidad, la contraposición (y la contraposición por la derecha) respecto de una negación y el principio de intercambio.

Damos también algunos ejemplos ilustrando el método de construcción viendo en particular que la (S, N) -implicación derivada de una negación N y de una t -conorma suma ordinal de t -conormas $S = \langle (a_j, b_j, S_j)_{j \in J} \rangle$, coincide con la N -suma ordinal de las (S, N) -implicaciones de las t -conormas S_j .

Como trabajo futuro, planeamos estudiar la preservación, mediante la N -suma ordinal, de otras propiedades adicionales de las implicaciones como el principio del orden, el principio de identidad, la ley de importación o incluso las propiedades distributivas. Además, sería interesante caracterizar las intersecciones de las implicaciones dadas por N -sumas ordinales con las familias de implicaciones más habituales.

Referencias

- [1] M. Baczyński, G. Beliakov, H. Bustince Sola, A. Pradera (Eds.): *Advances in Fuzzy Implication Functions*. Studies in Fuzziness and Soft Computing, vol. 300, Springer, Berlin Heidelberg, 2013.
- [2] M. Baczyński, B. Jayaram: *Fuzzy Implications*. Studies in Fuzziness and Soft Computing, vol. 231, Springer, Berlin Heidelberg, 2008.
- [3] M. Baczyński, B. Jayaram, S. Massanet, J. Torrens: *Fuzzy Implications: Past, Present, and Future*. In J. Kacprzyk and W. Pedrycz (eds.), *Springer Handbook of Computational Intelligence*, Springer Berlin Heidelberg, (2015) pp. 183–202.
- [4] E.P. Klement, R. Mesiar, E. Pap: *Triangular norms*. Kluwer Academic Publishers, Dordrecht, 2000.
- [5] M. Mas, M. Monserrat, J. Torrens, E. Trillas: *A survey on fuzzy implication functions*. IEEE Transactions on Fuzzy Systems. **15** (2007) 1107–1121.
- [6] Y. Su, A. Xie, H. Liu: *On ordinal sum implications*. Information Sciences. **293** (2015) 251–262.

Qualitative orness for OWA operators defined on lattice intervals

Gustavo Ochoa¹, Daniel Paternain², Inmaculada Lizasoain³, Humberto Bustince²

¹ *Departamento de Matemáticas. Universidad Pública de Navarra, 31006 Pamplona*

² *Departamento de Automática y Computación. Universidad Pública de Navarra. ISC
Institute of Smart Cities, 31006 Pamplona*

³ *Departamento de Matemáticas. INAMAT Institute for Advanced Materials. Universidad
Pública de Navarra, 31006 Pamplona*

{ochoa, daniel.paternain, ilizasoain, bustince}@unavarra.es

OWA (ordered weighted average) operators, defined by Yager in [4], constitute a kind of aggregation functions of real values very appreciated in decision making problems or data mining processes.

The great variety of possible weights to be chosen provide a wide kind of real OWA operators, from that given by the maximum (OR-operator) to that defined by the minimum (AND-operator). In order to classify these operators, an orness measure was defined by Yager in [5]. This way, for each weighting vector $(\alpha_1, \dots, \alpha_n) \in [0, 1]^n$, the corresponding OWA operator F_α is assigned a numerical value between 0 and 1 according to its less or greater similarity with the OR-operator:

$$\text{orness}(F_\alpha) = \sum_{i=1}^n \frac{n-i}{n-1} \alpha_i. \quad (1)$$

Notice that this numerical value is calculated as a *weighted average* of $(\alpha_1, \dots, \alpha_n)$, this time by means of the quantities $(\frac{n-1}{n-1}, \frac{n-2}{n-1}, \dots, \frac{1}{n-1}, 0)$, i.e., the first weights $\alpha_1, \alpha_2, \dots$ are given greater ponderations than the last ones.

In [1], OWA operators were extended to the setting of lattice values. Later, in [2], a qualitative orness has been introduced for OWA (ordered weighted average) operators defined on finite lattices $(L, \leq_L, T, S)$ endowed with a t-norm T and a t-conorm S . The orness measure is based on its position with respect to the OR-operator as in the real case. Its formulation involves the calculation of some elements $(d_1, \dots, d_n)$ in the lattice that play the role of the numbers $(\frac{n-1}{n-1}, \frac{n-2}{n-1}, \dots, \frac{1}{n-1}, 0)$ occurring in the real case. Afterwards,

$$\text{orness}(F_\alpha) = S(T(\alpha_1, d_1), \dots, T(\alpha_n, d_n)).$$

The aim of this paper is extending that concept of a qualitative orness measure to the case of the finite lattice $(L^m, \leq_{L^m}, \mathbb{T}, \mathbb{S})$, consisting of all m -ary intervals $[a_1, \dots, a_m]$ with $a_1 \leq_L \dots \leq_L a_m$ belonging to a finite bounded lattice L . It is clear that L^m is finite, but it

This work has been partially supported by Spanish project TIN2013-40765-P

is not easy to calculate the elements $(d_1, \dots, d_n)$ in L^m , which are necessary to obtain the qualitative orness of each m -ary interval OWA operator.

In this paper, those elements are calculated in two especial cases: when L is a distributive and complemented finite lattice and when L is a finite chain. In both cases, it is necessary to calculate previously the following elements $\{b_1, \dots, b_l\}$ in L^m , where $l = |L^m|$:

$$\begin{aligned} b_1 &= a_1 \vee \dots \vee a_l \in L^m; b_2 = \bigvee \{a_i \wedge a_j \mid i < j \in \{1, \dots, l\}\} \in L^m; \dots; \\ b_k &= \bigvee \{a_{j_1} \wedge \dots \wedge a_{j_k} \mid j_1 < \dots < j_k \in \{1, \dots, l\}\} \in L^m; \dots; \\ b_l &= a_1 \wedge \dots \wedge a_l \in L^m. \end{aligned}$$

Finally, by scaling these elements b_i , the target elements $(d_1, \dots, d_n)$ will be obtained.

Theorem 1 *Let L be a distributive and complemented finite lattice with a set T of p atoms and consider the lattice L^m for some $m \geq 1$. For each $k \in \{1, \dots, (m+1)^p = |L^m|\}$, call b_k the join of all the possible meets of k different elements in L^m . If we write $k = i(m+1)^{p-1} + j$ with $i \in \{0, 1, \dots, m\}$ and $j \in \{1, \dots, (m+1)^{p-1}\}$, then*

$$b_k = e_i(1) = [0_L, \dots, 0_L^{(i)}, 1_L, \dots, 1_L].$$

Theorem 2 *Let L be a finite chain with $|L| = n+1$. For each $k \in \{1, 2, \dots, CR(n+1, m) = |L^m|\}$, call $b_k \in L^m$ the m -ary interval which is the join of all the meets of k m -ary different intervals. Then $b_k = [r_1, r_2, \dots, r_m]$, where for each $i \in \{1, \dots, m\}$,*

$$r_i = \max\{r \in \{0, 1, \dots, n\} \mid k \leq \sum_{j=0}^{i-1} CR(r, j) \cdot CR(n-r+1, m-j)\}.$$

References

- [1] I. Lizasoain, C. Moreno. *OWA operators defined on complete lattices*. Fuzzy Sets and Systems **224** (2013) 36–52.
- [2] G. Ochoa, I. Lizasoain, D. Paternain, H. Bustince, N. R. Pal. *Qualitative orness for OWA operators*. Submitted to Information Fusion.
- [3] D. Paternain, G. Ochoa, I. Lizasoain, H. Bustince, R. Mesiar: *Quantitative orness for OWA operators*. Information Fusion **30** (2016) 27–35. 23–37.
- [4] R. R. Yager. *On ordered weighting averaging aggregation operators in multicriteria decision-making*. IEEE Transaction on Systems, Man and Cybernetics **18** (1988) 183–190.
- [5] R. R. Yager. *Families of OWA operators*. Fuzzy Sets and Systems **59** (1993) 125–148.

Órdenes Parciales para Hesitant Fuzzy Sets

L. Garmendia¹, R. González del Campo², J. Recasens³

¹ DISIA, Universidad Complutense de Madrid, Spain
lgarmend@fdi.ucm.es

² DSIC, Universidad Complutense de Madrid, Spain
rgonzale@estad.ucm.es

³ 3 Universitat Politecnica de Catalunya, Spain
j.recasens@upc.edu

Resumen

Se define un nuevo orden parcial $\leq_{\mathbb{H}}$ en el conjunto $\mathbb{H}$ de subconjuntos finitos del intervalo unidad. Con este preorden $\leq_{\mathbb{H}}$ no es un retículo, pero sí lo son algunos subconjuntos de $\mathbb{H}$ interesantes. $\mathbb{H}$ se puede inyectar de forma natural en un retículo $(\overline{\mathbb{B}}, \leq_{\overline{\mathbb{B}}})$. En $(\mathbb{H}, \leq_{\mathbb{H}})$ y $(\overline{\mathbb{B}}, \leq_{\overline{\mathbb{B}}})$ se pueden definir t-normas.

Keywords: hesitant fuzzy sets, subconjuntos finitos del intervalo unidad, orden parcial.

Los hesitant fuzzy sets generalizan el concepto de conjunto borroso al permitir la posibilidad de asignar más de un valor del intervalo unidad a los objetos del universo de discurso. Son útiles cuando hay dudas (*hesitation*) al asignar valores numéricos a los objetos [1], [2].

Un tema relevante es cómo comparar hesitant fuzzy sets, lo que lleva a la cuestión de definir relaciones de orden en el conjunto $\mathbb{H}$ de los subconjuntos finitos de $[0, 1]$. Hay una manera natural de comparar subconjuntos con el mismo cardinal: si $A = \{a_1, a_2, \dots, a_n\}$ y $B = \{b_1, b_2, \dots, b_n\}$ (asumiendo como haremos durante todo el trabajo que los elementos de los subconjuntos están ordenados ($a_1 < a_2 < \dots < a_n$ y $b_1 < b_2 < \dots < b_n$)), la comparación punto a punto parece adecuada ($A \leq_{\mathbb{H}} B$ si, y sólo si, $a_i \leq b_i$ para todo $i = 1, 2, \dots, n$). El problema surge al intentar comparar conjuntos de diferente cardinal.

Dado que comparar dos conjuntos del mismo cardinal es una tarea fácil, la idea para comparar dos conjuntos A y B de diferentes cardinales n y m es repetir sus elementos convenientemente para obtener dos sucesiones de la misma longitud.

Definición 1. Sea $A = \{a_1, a_2, \dots, a_n\} \in \mathbb{H}$ y $r \in \mathbb{N}$. $A_r \in [0, 1]^{rn}$ es el vector de rn coordenadas definido por $A_r = (\underbrace{a_1, \dots, a_1}_{r \text{ veces}}, \underbrace{a_2, \dots, a_2}_{r \text{ veces}}, \dots, \underbrace{a_n, \dots, a_n}_{r \text{ veces}})$.

Definición 2. Sean $A = \{a_1, a_2, \dots, a_n\}$ y $B = \{b_1, b_2, \dots, b_m\}$ subconjuntos finitos del intervalo unidad y $\text{mcm}(n, m)$ el mínimo común múltiplo de n y m . Reescribiendo $A_{\frac{\text{mcm}(n, m)}{n}} = (c_1, c_2, \dots, c_{\text{mcm}(n, m)})$ y $B_{\frac{\text{mcm}(n, m)}{m}} = (d_1, d_2, \dots, d_{\text{mcm}(n, m)})$,

$A \leq_{\mathbb{H}} B$ si, y sólo si $c_i \leq d_i$ para todo $i = 1, 2, \dots, \text{mcm}(n, m)$.

Este trabajo ha sido parcialmente financiado por el proyecto TIN 2012-32484.

Ejemplo 3. $A = \{0,2,0,4\} \leq_{\mathbb{H}} B = \{0,3,0,5,0,8\}$, porque $\text{mcm}(2,3) = 6$,

$$A_3 = (0,2,0,2,0,2,0,4,0,4,0,4) \quad B_2 = (0,3,0,3,0,5,0,5,0,8,0,8)$$

y $0,2 \leq 0,3$, $0,2 \leq 0,5$, $0,4 \leq 0,5$, $0,4 \leq 0,8$.

Proposición 4. $\leq_{\mathbb{H}}$ es un orden parcial en $\mathbb{H}$ con $\{0\}$ y $\{1\}$ cota inferior y superior respectivamente.

Este orden parcial respeta el orden puntual en $[0, 1]$ y la ordenación usual de los intervalos de $[0, 1]$. Otro ejemplo: si $A = \{a_1, a_2\}$ y $B = \{b_1, b_2, b_3\}$, entonces $A \leq_{\mathbb{H}} B$ si y sólo si, $a_1 \leq b_1$ y $a_2 \leq b_2$, y $B \leq_{\mathbb{H}} A$ si, y sólo si, $b_2 \leq a_1$ y $b_3 \leq a_2$.

$(\mathbb{H}, \leq_{\mathbb{H}})$ no es un retículo. Por ejemplo, los conjuntos $A := \{0,2,0,4\}$ y $B = \{0,1,0,3,0,6\}$ carecen de supremo y de ínfimo. Sin embargo, si uno de los cardinales de A y B es múltiplo del otro, entonces estos conjuntos poseen supremo e ínfimo. De este hecho se sigue el siguiente resultado.

Proposición 5. Sea $R = r_1, r_2, \dots, r_n, \dots$ una sucesión de números naturales con r_{i+1} múltiplo de r_i para cada $i \geq 0$, $\mathbb{H}_{r_i}$ el conjunto de subconjuntos finitos de $[0, 1]$ de cardinal r_i y $\mathbb{H}_R = \bigcup_{i=1}^{\infty} \mathbb{H}_{r_i}$. Entonces $(\mathbb{H}_R, \leq_{\mathbb{H}})$ es un retículo.

$(\mathbb{H}, \leq_{\mathbb{H}})$ no es un retículo porque le "faltan" objetos. Si añadimos los conjuntos del intervalo unidad con algunos elementos repetidos, los llamados bolsas o multiconjuntos [3], obtenemos un retículo $(\mathbb{B}, \leq_{\mathbb{B}})$ en el que $(\mathbb{H}, \leq_{\mathbb{H}})$ está inmerso.

A partir de una t-norma usual (i.e., definida en $[0, 1]$) se pueden generar sendas t-normas en $(\mathbb{H}, \leq_{\mathbb{H}})$ y en $(\mathbb{B}, \leq_{\mathbb{B}})$.

Proposición 6. Sean $A = \{a_1, \dots, a_n\}$ y $B = \{b_1, b_2, \dots, b_m\}$ dos subconjuntos finitos de $[0, 1]$ y T una t-norma tal que $a < a'$ y $b < b'$ implica $T(a, b) < T(a', b')$. Escribamos $A_{\frac{\text{lcm}(n,m)}{n}} = (c_1, c_2, \dots, c_{\text{mcm}(n,m)})$ y $B_{\frac{\text{lcm}(n,m)}{m}} = (d_1, d_2, \dots, d_{\text{mcm}(n,m)})$. Entonces

$$\mathbb{T}(A, B) = \{T(c_1, d_1), T(c_2, d_2), \dots, T(c_{\text{mcm}(n,m)}, d_{\text{mcm}(n,m)})\}$$

es una t-norma en $(\mathbb{H}, \leq_{\mathbb{H}})$.

Referencias

- [1] V. Torra. Hesitant fuzzy sets. International Journal of Intelligent Systems, 25(6), 529–539, 2010.
- [2] V. Torra and Y. Narukawa. On hesitant fuzzy sets and decision. En 18th IEEE International Conference on Fuzzy Systems, 1378–1382, 2009.
- [3] R:R: Yager. On the theory of bags. International Journal of General Systems, 13(1), 23–37, 1986

A fuzzy approach to the problem of group identification

J. C. R. Alcantud¹, R. de Andrés Calle²

¹ *BORDA Research Unit and Multidisciplinary Institute of Enterprise (IME),
University of Salamanca, Spain.*

² *BORDA Research Unit, PRESAD Research Group and
Multidisciplinary Institute of Enterprise (IME),
University of Salamanca, Spain.*

{jcr,rocioac}@usal.es

In this contribution we consider the problem of self-selection of individuals in a group with respect to a given vague or imprecise attribute, e.g. ‘being Jewish’, ‘belonging to a nationality’, ‘being African-American’ or ‘belonging to a racial group’, etc. Our purpose is to extend the study of the *collective identity* in *Social Choice Theory* by allowing the natural possibility of partial membership degrees of the individuals to the groups. Then, we define the problem of group identification in a fuzzy environment.

Following Kasher and Rubinstein’s approach [6], we consider the collective identity problem as an aggregation problem of individual assessments within a group. We provide evidences that extending the analysis of the problem of collective identity by the introduction of partial memberships (cf., Alcantud and de Andrés Calle [2]) permits to better capture the nuances of preference modeling. In the fuzzy framework that we propose in this contribution not only we gain generality, but also new families of aggregators can be employed to determine the social outcome. Particularly, besides the few traditional aggregation rules from the literature we can now utilize the discrete Choquet integral, ordered weighted averaging operators (OWA), and weighted OWA operators (WOWA) as new tools to determine the degree of memberships of the individuals to the collective identity.

We show that it is still possible to preserve the nature of earlier investigations in Social Choice (e.g., the axiomatic treatment of the models: Alcalde and M. Vorsatz, [1], Alcantud, de Andrés Calle and Cascón [3, 4, 5]) and to avoid restrictions imposed by the crisp viewpoint of the problem. We prove these benefits by characterizing the rules that evaluate each individuals membership to the collectivity by a discrete Choquet integral, and by producing a positive escape from Kasher and Rubinstein’s impossibility theorem.

References

- [1] J. Alcalde and M. Vorsatz. Measuring the cohesiveness of preferences: an axiomatic analysis. *Soc Choice Welf*, 41:965–988, 2013.
- [2] J. C. R. Alcantud and R. de Andrés Calle. A fuzzy viewpoint of consensus measures in social choice. In *Actas del XVII Congreso Español sobre Tecnologías y Lógica Fuzzy (ESTYLF 14)*, pages 87–92, 2014.

Financial support by the Spanish Ministerio de Ciencia e Innovación under Project ECO2012–32178 (R. de Andrés Calle) is gratefully acknowledged.

- [3] J. C. R. Alcantud, R. de Andrés Calle, and J. M. Cascón. Consensus and the act of voting. *Studies in Microeconomics*, 1:1–22, 2013.
- [4] J. C. R. Alcantud, R. de Andrés Calle, and J. M. Cascón. On measures of cohesiveness under dichotomous opinions: some characterizations of approval consensus measures. *Information Sciences*, 240:45–55, 2013.
- [5] J. C. R. Alcantud, R. de Andrés Calle, and J. M. Cascón. A unifying model to measure consensus solutions in a society. *Mathematical and Computer Modelling*, 57:1876–1883, 2013.
- [6] A. Kasher and A. Rubinstein. On the question ‘Who is a J?’. *Logique & Analyse*, 160:385–395, 1997.

Ajuste de dosis de insulina mediante redes neuronales y control borroso

Ana García Álvarez ¹, Matilde Santos ¹

¹ *Facultad de Informática, Universidad Complutense de Madrid, 28030-Madrid*
angarez@live.com, msantos@ucm.es

El páncreas de los pacientes con diabetes de tipo I no libera la insulina necesaria para metabolizar la glucosa ingerida a través de los alimentos. Por ello requieren del suministro artificial de insulina para poder vivir. Uno de los tratamientos más habituales para esta enfermedad es la inyección subcutánea de insulina, frecuentemente diferenciada en dos fases. Una inyección diaria de acción lenta y varias inyecciones de acción rápida para neutralizar la glucosa ingerida durante cada una de las comidas principales del día.

Si bien es el médico endocrino el encargado de definir tanto el tratamiento como una estimación de las dosis a inyectar para cada una de las fases y acciones, el metabolismo de cada paciente va a condicionar la cantidad de insulina necesaria, y debe ser este último el que decida aumentar o disminuir la dosis en torno a la pauta facilitada por el doctor, en función de la hora del día, el nivel de actividad y estrés al que se va a ver sometido o el tipo de comida que va a ingerir. Sin embargo, a nivel paciente sólo hay un reducido conjunto de herramientas de recomendación de dosis de insulina.

Dada la incertidumbre en el ajuste de la dosis de insulina, debido principalmente al carácter no lineal de comportamiento metabólico del cuerpo humano, en este trabajo se plantea el desarrollo de un mecanismo inteligente para la asistencia al paciente en dicho ajuste, con lo cual se pretende conseguir evitar la fluctuación fuera de márgenes deseados del nivel de glucosa, principalmente en el caso de valores de riesgo de sufrir hipoglucemia.

El presente estudio se enfoca a la definición de la dosis de insulina recomendada para un paciente determinado en función de un conjunto de datos previos. Se ha partido de información real facilitada por un paciente anónimo. El tratamiento que sigue consta de una inyección diaria de insulina lenta Lantus (7 unidades) y tres inyecciones de insulina rápida Humalog (una previa a cada comida). Dispone de una pauta de inyección de dosis basada en el nivel de glucosa en sangre a medir por medio del instrumento Bayer Contour Next USB Meter. Se ha trabajado con mediciones tomadas entre los meses de febrero y mayo de 2014. El perfil de actividad/estrés del paciente varía en función del día de la semana. Realiza actividad deportiva intensa dos días a la semana y moderada los tres restantes. El fin de semana no realiza ningún tipo de actividad. El paciente no se somete a ningún tipo de dieta. De manera voluntaria selecciona alimentos sin azúcares añadidos y controla la ingestión de hidratos de carbono (pastas, arroces, patatas, pan, etc).

La identificación de la evolución de esta variable en el paciente se ha realizado mediante el uso de redes neuronales, mientras que el control de la dosis de insulina se ha diseñado mediante control difuso.

Para la identificación se ha seleccionado una red feed-forward con dos capas, una única capa oculta de 10 neuronas, aunque se probaron configuraciones más complejas (varias capas ocultas). Se ha aplicado el mecanismo de entrenamiento Levenberg-Marquardt, el criterio del error cuadrático medio y el mecanismo de separación de datos aleatorio. La información de entrada es un vector el cual especifica la franja horaria (asignación discreta para desayuno, comida o cena), el nivel de actividad (asignación discreta para nulo,

moderado o intenso), el nivel de glucemia medido o el número de unidades de insulina inyectado. La información de salida es el estado glucémico en la franja horaria siguiente. Los resultados obtenidos son aceptables excepto para los casos de niveles de glucemia bajos, que son sin duda el principal objetivo de la herramienta.

Se ha diseñado un sistema de control difuso de dosis de insulina para sustitución del mecanismo de decisión por parte del enfermo con tres variables de entrada (glucometría, nivel de actividad y franja horaria), que provienen de la red neuronal, y una de salida (el incremento en la dosis de insulina sobre la pauta en función de la opinión del enfermo sobre su estado metabólico).

Uno de los principales inconvenientes del problema de identificación del sistema, si no el fundamental, ha sido la limitación de datos facilitados por parte del paciente. Desarrollar una herramienta de obtención de datos poco costosa para el paciente (estilo aplicación móvil donde registre su nivel de actividad, estrés o tipo de alimentación), permitiría disponer de mayor cantidad de información a nutrir a la red neuronal). A su vez, conseguir mayores medidas de niveles de glucosa sería también beneficioso. El paciente sería más proclive a realizar mayor número de mediciones en caso de disponer de un dispositivo de control no invasivo existente en el mercado (el actual exige realizar un pinchazo para obtención de una muestra de sangre, motivo principal por el cual reduce el número de tomas de medida).

References

- [1] Baghdadi, G., & Nasrabadi, A. M. (2007, August). *Controlling blood glucose levels in diabetics by neural network predictor*. In: IEEE Conf. Engineering in Medicine and Biology Society (2007) 3216–3219.
- [2] Bellazzi, R., Nucci, G., & Cobelli, C. (2001). *The subcutaneous route to insulin dependent diabetes therapy*. Engineering in Medicine and Biology Magazine, IEEE. No. 20(1) (2001) 54–64.
- [3] Ibbini, M. S., & Masadeh, M. A. (2005). *A fuzzy logic based closed-loop control system for blood glucose level regulation in diabetics*. Journal of medical engineering & technology. No. 29(2) (2005) 64–69.
- [4] Amato, F., Lopez, A., Peña-Méndez, E. M., Vañhara, P., Hampl, A., & Havel, J. (2013). *Artificial neural networks in medical diagnosis*. Journal of applied biomedicine. No. 11(2) (2013) 47–58.
- [5] Zhu, K. Y., Liu, W. D., & Xiao, Y. *Application of Fuzzy Logic Control for Regulation of Glucose Level of Diabetic Patient*. In: Machine Learning in Healthcare Informatics (2014) 47–64. Springer Berlin Heidelberg.

A new consensus measure based on Pearson correlation coefficient

T. González-Arteaga¹, R. de Andrés Calle², F. Chiclana³

¹*PRESAD Reasearch Group and IME,
University of Valladolid, Spain
teresag@eio.uva.es*

²*BORDA Research Unit, PRESAD Research Group and IME,
University of Salamanca, Spain
rocioac@usal.es*

³*Centre for Computational Intelligence, School of Computer Science and Informatics,
De Montfort University, UK
chiclana@dmu.ac.uk*

Obtaining *consensual* solutions is an important issue in decision making processes. It depends on several factors such as experts' opinions, principles, knowledge, experience, etc. In the literature we can find a considerable amount of consensus measurement from different research areas (from a Social Choice perspective: Alcalde-Unzu and Vorsatz [1], Alcantud, de Andrés Calle and Cascón [2] and Bosch [3], among others and from Decision Making Theory: González-Arteaga, Alcantud and de Andrés Calle [4] and González-Pachón [5], Herrera, Herrera-Viedma and Chiclana [7], Herrera-Viedma et al. [6] and Wu et al. [8], among others). Most of them have a common point, they are based on distances or similarity functions.

In the present contribution we propose a new approach based on the use of the Pearson correlation coefficient to measure consensus. Moreover, we suppose a general framework considering experts' opinions modelled by fuzzy preference relation. The new *correlation consensus measurement* takes into account concordance between preferences intensities for pairs of alternatives and it verifies important properties. In addition, we prove that our proposal is a different approach to traditional consensus measures based on distances or similarities.

References

- [1] J. Alcalde-Unzu and M. Vorsatz. Measuring the cohesiveness of preferences: An axiomatic analysis. *Social Choice and Welfare*, 41:965–988, 2013.
- [2] J. C. R. Alcantud, R. de Andrés Calle, and J. M. Cascón. Consensus and the act of voting. *Studies in Microeconomics*, 1(1):1–22, 2013.

Financial support by the Spanish Ministerio de Ciencia e Innovación under Project ECO2012–32178 (R. de Andrés Calle, T. González-Arteaga) is gratefully acknowledged.

- [3] R. Bosch. *Characterizations of Voting Rules and Consensus Measures*. PhD thesis, Tilburg University, 2005.
- [4] T. González-Arteaga, J.C.R. Alcantud, and R. de Andrés Calle. A cardinal dissensus measure based on the Mahalanobis distance. *European Journal of Operational Research*, In press.
- [5] J. González-Pachón and C. Romero. Distance-based consensus methods: a goal programming approach. *Omega*, 27(3):341–347, 1999.
- [6] E. Herrera-Viedma, F. J. Cabrerizo, J. Kacprzyk, and W. Pedrycz. A review of soft consensus models in a fuzzy environment. *Information Fusion*, 17:4–13, 2014.
- [7] E. Herrera-Viedma, F. Herrera, and F. Chiclana. A consensus model for multiperson decision making with different preference structures. *IEEE Transactions on Systems, Man, and Cybernetics - Part A: Systems and Humans*, 32(3):394–402, 2002.
- [8] J. Wu, F. Chiclana, and E. Herrera-Viedma. Trust based consensus model for social network in an incomplete linguistic information context. *Applied Soft Computing*, 35:827–839, 2015.

Una nueva visión de los Z-números de Zadeh

Sebastia Massanet, Juan Vicente Riera, Joan Torrens

*Universitat de les Illes Balears, Dept. Matemàtiques i Informàtica
Ctra. Valldemossa km 7.5, E-07122 Palma, Illes Balears
{s.massanet, jvicente.riera, jts224}@uib.es*

El razonamiento humano se caracteriza por la capacidad de realizar múltiples tareas sin la necesidad implícita de efectuar de manera directa ningún tipo de medida o de cálculo, sino que a partir de las propias percepciones y la utilización de datos imprecisos es capaz de tomar decisiones con una gran exactitud. Esta idea sirvió a L. Zadeh para introducir la computación con palabras (CWW) en [4] entendida como aquella en la que se utilizan palabras, percepciones o incluso ciertas frases del lenguaje natural, diferenciándose de la computación tradicional centrada exclusivamente en la manipulación de números y de símbolos. Además, señala dos razones principales que avalan su estudio. La primera de ellas, CWW es una necesidad cuando la información que se maneja es demasiado imprecisa para justificar el uso de números y la segunda cuando se permite una tolerancia con la imprecisión que puede ser utilizada para lograr una mejor adaptación y descripción de la realidad.

Por otra parte, la incertidumbre es un denominador común en una gran variedad de problemas de toma de decisiones. Esta incertidumbre muchas veces se debe a que la información utilizada por los expertos no solo es de tipo cuantitativo sino de tipo cualitativo, por lo que se hace difícil el manejo de la misma. Este hecho ha motivado que la computación con palabras sea un tema candente de constante investigación en la teoría de toma de decisiones. Por esta razón diferentes modelos lingüísticos se han presentado en la literatura con el propósito de modelar adecuadamente la opinión de los expertos. Así en [3] se presenta un sistemático estudio sobre modelos lingüísticos multigranulares borrosos en el que se consideran seis diferentes categorías, entre las que destaca el modelo multigranular basado en números borrosos discretos [2].

Siguiendo con la idea de buscar modelos que describan adecuadamente el lenguaje natural, L. Zadeh en [5] introduce el concepto de Z-número, interpretado como una pareja ordenada de números borrosos (A, B) . Así, cuando un Z-número es asociado a una variable de incertidumbre real valorada X , el triplete ordenado (X, A, B) es entendido como una Z-valoración, donde la primera componente A es interpretada como una restricción sobre los valores que puede tomar X y la segunda componente B nos indica el grado de certeza (seguridad, confianza, posibilidad, probabilidad...) que se tiene sobre el valor descrito por A . En este sentido, *(Estabilidad política, alta, probablemente)* es un ejemplo de Z-valoración. Desde entonces, muchos investigadores han dedicado su estudio a la teoría de Z-números desde diferentes perspectivas (fundamentos teóricos y posibles aplicaciones) destacando el monográfico [1]. Sin embargo, tal como ya apuntó Zadeh [5] “*Problems involving computation with Z-numbers are easy to state but far from easy to solve*”, enfatizando el hecho de que los procesos de inferencia basados en Z-información tienen una complejidad operacional

Este trabajo has sido subvencionado por el proyecto TIN2013-42795-P.

y computacional muy elevada, debida en parte al enfoque probabilístico asignado originariamente a la segunda componente. Dicha complejidad ha sido mencionada ya en la literatura [1] y por esta razón diferentes aproximaciones se han propuesto con la idea de simplificar el cálculo operacional y por tanto el coste computacional. Una de ellas, es la establecida por R. Aliev [1], que introduce el concepto de Z -número discreto, basado en números borrosos discretos pero manteniendo el mismo enfoque probabilístico establecido por Zadeh. Sin embargo, Zadeh [5] también afirma que la segunda componente B , es una medida de confianza (certeza) de la primera componente que puede ser interpretada desde diferentes puntos de vista: seguridad, confianza, creencia o probabilidad.

Teniendo en cuentas estas ideas, en esta comunicación presentamos una nueva visión del concepto de Z -números basada en números borrosos discretos con valores en una cadena finita L_n . En este enfoque, no se considera la perspectiva probabilística de la segunda componente, sino que el grado de confianza que se tiene sobre la primera se representará mediante una expresión lingüística basada en números borrosos discretos con soporte un intervalo de una cadena finita L_n , o abreviadamente números en $\mathcal{A}_1^{L_n}$. Formalmente,

Definición 1 Sean L_n y L_m dos cadenas finitas. Una pareja ordenada de números borrosos discretos (A, B) con $A \in \mathcal{A}_1^{L_n}$, $B \in \mathcal{A}_1^{L_m}$ la llamaremos (L_n, L_m) – Z -número borroso discreto. Si asociamos un (L_n, L_m) – Z -número borroso discreto a una variable de incerteza, X , la primera componente, A , desempeñará la función de restricción borrosa sobre X , mientras que el número borroso discreto, B describirá una estimación imprecisa de la confianza sobre A . El triplete ordenado (X, A, B) será llamado Z -valoración y se escribirá X es (A, B) .

Desde este nuevo enfoque, será posible utilizar las diferentes operaciones establecidas entre números en $\mathcal{A}_1^{L_n}$, así como definir operaciones de agregación entre Z -valoraciones, reduciendo el coste operacional y computacional propio de los clásicos Z -números, facilitando además los procesos de inferencia basados en Z -información. Un ejemplo ilustrativo en el que se aplica este método de agregación para obtener una valoración global en base a las Z -valoraciones dadas por expertos es presentado.

Referencias

- [1] R. A. Aliev, O. H. Huseynov, R. R. Aliyev, and A. A. Alizadeh. *The Arithmetic of Z-Numbers: Theory and Applications*. World Scientific Publishing, 2015.
- [2] S. Massanet, J. V. Riera, J. Torrens, and E. Herrera-Viedma. A new linguistic computational model based on discrete fuzzy numbers for computing with words. *Information Sciences*, 258:277–290, 2014.
- [3] J. Morente-Molinera, I. Pérez, M. Ureña, and E. Herrera-Viedma. On multi-granular fuzzy linguistic modeling in group decision making problems: A systematic review and future trends. *Knowledge-Based Systems*, 74:49 – 60, 2015.
- [4] L. Zadeh. Fuzzy logic = computing with words. *IEEE Transactions of Fuzzy Systems*, 4:103–111, 1996.
- [5] L. Zadeh. A note on Z -numbers. *Information Sciences*, 9(1):43 – 80, 2011.

The weighted Hamming distance as a divergence measure

Susana Montes¹, Jorge Elorza², Carlos Bejines³, Jean Bragard²

¹ *UNIMODE Research Unit, Universidad de Oviedo*
montes@uniovi.es

² *Dept.de Física y Matemática Aplicada, Facultad de Ciencias, Universidad de Navarra*
jelorza@unav.es, jbragard@unav.es

³ *Facultad de Ciencias, Universidad de Málaga*
cbejines@uma.es

The comparison of two sets is a very important topic, due to its application to several fields, such as pattern recognition or decision making. Thus, if we are working in a fuzzy environment, the estimation of the difference between two fuzzy sets is crucial.

This task is usually performed by applying a suitable comparison measure to a pair of fuzzy sets. This relation has been done with different comparison measures, by considering different points of view. A good study about measures of comparison was given by Bouchon-Meunier et al. [1]. More recently, a review of the measures based on the differences, where the relationships among them were studied in detail, was proposed by Couso et al. [2]. Among these measures, distances, dissimilarities and divergences seem to be very appropriate. In the three cases, they have two common axioms:

(Ax1) Irreflexivity: $D(A,A) = 0$ for any fuzzy set $A \in F(X)$ where $F(X)$ denotes the set formed by all the fuzzy subsets of the referential X .

(Ax2) Symmetry: $D(A,B) = D(B,A)$ for any pair of fuzzy sets $A, B \in F(X)$.

Thus, the different points of view are only taken into account in the third axiom:

- In the case of distances, this is the triangular inequality:

$$D(A,B) + D(B,C) \geq D(A,C), \quad \forall A, B, C \in F(X).$$

- In the case of dissimilarities it is some kind of monotonicity:

$$D(A,B) \leq D(A,C) \text{ and } D(B,C) \leq D(A,C) \text{ for any } A, B, C \in F(X) \text{ with } A \subseteq B \subseteq C.$$

- In the case of divergences it is based on the union and intersection with a third set:

$$D(A \cap C, B \cap C) \leq D(A,B) \text{ and } D(A \cup C, B \cup C) \leq D(A,B), \forall A, B, C \in F(X).$$

In the last case, the concept of union and intersection is involved, so, this definition depends on the t-norm and t-conorm selected to define these operations. Thus, in fact, we should call them as T - S -divergences or just T -divergences when we consider the dual t-conorm (which will be here the case).

The research reported in this paper has been partially supported by the Spanish Government under Projects FIS2011-28820 and TIN2014-59543-P.

It is proven that they are not the same concept (see [3]) and there is not any general relation among them. However, in the particular case of the Zadeh t-norm (minimum t-norm) we know that any divergence is a dissimilarity (see [4]). Moreover, when they have the local property (they can be decomposed point by point), the family of the local dissimilarities and the family of the local T_M -divergences are the same, where T_M denotes the minimum t-norm. A very useful measure in this case is the weighted Hamming distance, defined by:

$$D_{wH}(A, B) = \sum_{x \in X} \alpha_x \cdot |A(x) - B(x)|$$

where $\alpha_x \geq 0$ for any $x \in X$ and $\sum_{x \in X} \alpha_x = 1$. This measure is at the same time a distance, a dissimilarity and a T_M -divergence and it has been used in many practical cases. Thus, we would like to study its general behaviour when any t-norm is considered.

Thus, we have proven that D_{wH} is a T -divergence if and only if the t-norm T has the 1-Lipschitz property or, equivalently, if T is a copula. Therefore, we know that only t-norms between the Łukasiewicz t-norm and the minimum t-norm could define a T -divergence. This is not the case of the drastic t-norm.

Apart from that, we know that a continuous Archimedean t-norm satisfies the 1-Lipschitz property with constant 1 if and only if its additive generator is a convex function. Moreover, we know that the minimum t-norm satisfies this property. Thus, we can use the previous results to characterize the behaviour of some of the most famous families of t-norms. Thus, the indexes λ such that D_{wH} is a T -divergence are:

Family	Indexes
Schweizer-Sklar $\lambda \in [-\infty, \infty]$	$\lambda \in [-\infty, -1]$
Hamacher $\lambda \in [0, \infty]$	$\lambda \in [0, 2]$
Frank $\lambda \in [0, \infty]$	$\lambda \in [0, \infty]$
Yager $\lambda \in [0, \infty]$	$\lambda \in [1, \infty]$
Dombi $\lambda \in [0, \infty]$	$\lambda \in [1, \infty]$
Sugeno-Weber $\lambda \in [-1, \infty]$	$\lambda \in [0, \infty]$
Aczél-Alsina $\lambda \in [0, \infty]$	$\lambda \in [1, \infty]$

References

- [1] B. Bouchon-Meunier, M. Rifqi, S. Bothorel. Towards general measures of comparison of objects, *Fuzzy Sets and Systems*, 84, 143–153, 1996.
- [2] I. Couso, L. Garrido, L. Sánchez, Similarity and dissimilarity measures between fuzzy sets: A formal relational study, *Information Sciences*, 229, 122 – 141, 2013.
- [3] V. Kobza, V. Janiš, S. Montes, Generalized local divergence measures between fuzzy subsets, *IEEE Transactions on Fuzzy Systems* submitted.
- [4] S. Montes, I. Couso, P. Gil, C. Bertoluzza, Divergence measure between fuzzy sets, *International Journal of Approximate Reasoning*, 30, 91–105, 2002.

Aplicación del uso de valoraciones *hesitant* lingüísticas en una red social de economía colaborativa

Rosana Montes¹, Ana M^a Sanchez¹, Pedro Villar¹, Francisco Herrera²

¹ *Dpto. Lenguajes y Sistemas Informáticos. Universidad de Granada*
rosana@ugr.es, amlopez@ugr.es, pvillarc@ugr.es

² *Dpto. Ciencias de la Computación e Inteligencia Artificial. Universidad de Granada*
herrera@decsai.ugr.es

En el ámbito de la toma de decisiones (TD) un problema tipo se caracteriza por una serie de alternativas sobre las que elegir, que son valoradas por expertos en el área. Los expertos son personas capaces de proporcionar valoraciones específicas sobre determinadas propiedades de cada alternativa, de cara a que un proceso de computación determine cual es la mejor de todas. Independientemente de la complejidad computacional, estos sistemas pueden no ser fáciles de aplicar, ya que el experto puede no tener clara su valoración, o bien, encuentra reparos a la hora de proporcionar valoraciones exactas o cuantitativas. Gracias a la incorporación de Modelos Lingüísticos –computing with words (CW) [1]– se ha permitido la manipulación efectiva de información cualitativa, mucho más cercana a la forma de razonar de una persona. Pero además se ha encontrado un enorme potencial en la aplicación de valoraciones *hesitant* lingüísticas [2] (en adelante HFLTS), que incorporan mecanismos la expresión de valoraciones poco precisas, esto es, el experto no solo da una valoración individual sino que puede expresar una respuesta múltiple (imprecisa). Con el Enfoque Lingüístico Difuso (ELD) que usa palabras o frases definidas en un lenguaje natural, nos encontramos con sistemas más flexibles que pueden aplicarse a entornos cotidianos con éxito.

Teranga Go! [3] es una plataforma de economía colaborativa cuyo objetivo es el de compartir gastos entre usuarios registrados en la plataforma que viajen entre España y Senegal. Como cualquier red social, los usuarios cuentan con un perfil, un servicio base y diversas formas de consultar información y comunicarse (mensajería interna, foros, comentarios, me gusta, etc). Como cualquier red social tiene un problema común: construir la confianza entre los usuarios. Puede decirse que la confianza es un valor añadido en una comunidad online, ya que para registrarse en una plataforma online tan solo hace falta un correo electrónico. Nosotros pretendemos que la plataforma aporte una valoración de confianza relativa a los conductores involucrando a los propios usuarios. Tras un largo viaje (los trayectos entre España y Senegal implican varios días de conducción) solicitamos a los participantes de éste que se evalúen mutuamente, aportando datos tanto privados como públicos para el resto de la comunidad. Las valoraciones *hesitant* (una etiqueta o varias etiquetas lingüísticas consecutivas) adquiridas sobre una persona a lo largo de varios viajes y aportadas por diversos individuos (véase la Figura 1.a), son empleadas en un sistema de toma de decisión lingüístico.

En este trabajo aplicamos la representación lingüística 2-tuplas [4] con HFLTS a un problema real de toma de decisiones multiexperto-multicriterio. La aplicación del modelo se

Este trabajo está soportado por el Ministerio de Ciencia e Innovación en el marco del proyecto TIN2014-57251-P y por la Junta de Andalucía en el marco del proyecto P11-TIC-7765.

realiza en el que puede ser el ámbito más cotidiano en la actualidad: internet y las redes sociales. Lo que pretendemos es dar un valor añadido a las comunidades online aportando la implementación que hemos realizado de un ELD en Teranga Go!. No aplicamos el problema tipo para obtener un ranking que determine cuál es la mejor de las alternativas, en nuestro caso valoramos una única alternativa (el usuario que recibe las valoraciones). El problema es multicriterio (se valoran aspectos como la forma de conducción y la seguridad, la limpieza o la facilidad de conversación con la persona) y multiexperto (todo el que haya viajado conmigo puede valorarme). El resultado es una etiqueta lingüística dentro del conjunto { 'muy bajo', 'mejorable', 'novel', 'satisfactorio', 'confiable', 'muy bueno', 'excelente' } que determina el *karma* de la persona. El *karma* es un dato público del perfil del usuario, tal y como se muestra en la Figura 1.b.

El enfoque lingüístico difuso puede aplicarse a dar más servicios de la plataforma, para permitir comparar entre varios viajes similares con un mismo origen-destino-fecha, cuál de ellos se ajusta mejor a mis preferencias de viaje.

(a) Mi valoración personal de:
 María Teresa - Teranga

Seguridad en el viaje
☐ Horrible ☐ Muy mal ☐ Mal ☒ Normal ☒ Bien ☐ Muy bien ☐ Excelente
Impresión general relativa tanto a la forma de conducir de la persona como al grado de cuidado del vehículo (puede seleccionar más de una opción siempre que las marcas sean consecutivas).

Limpieza e higiene
☐ Horrible ☐ Muy mal ☐ Mal ☒ Normal ☐ Bien ☐ Muy bien ☐ Excelente
Impresión general relativa tanto a la persona evaluada como al cuidado del vehículo (puede seleccionar más de una opción siempre que las marcas sean consecutivas).

Conversación y compañía
☐ Horrible ☐ Muy mal ☐ Mal ☒ Normal ☒ Bien ☐ Muy bien ☐ Excelente
Impresión general relativa a la persona evaluada (puede seleccionar más de una opción siempre que las marcas sean consecutivas).

Confort del coche
☐ Horrible ☒ Muy mal ☒ Mal ☒ Normal ☐ Bien ☐ Muy bien ☐ Excelente
Impresión general relativo al coche (puede seleccionar más de una opción siempre que las marcas sean consecutivas).

¿Recomendarías el conductor a otros usuarios?

Esta respuesta aparecerá públicamente al igual que los comentarios y no es obligatoria.

¿El conductor se ha ajustado a lo acordado previamente?

Esta respuesta aparecerá públicamente al igual que los comentarios y no es obligatoria.

(b) Rosana Montes
 Tipo de usuario: Conductor
 > Sobre mí
 > Preferencias en un viaje
 > Mi coche
 > Valoración personal
 Seguridad en el viaje: 100 % importancia.
 Limpieza e higiene: 100 % importancia.
 Conversación y compañía: 49 % importancia.
 Confort del coche: 47 % importancia.

Karma: mejorable
 A partir de 3 valoraciones de usuarios.
 Editar imagen de perfil

Figura 1: (a) Tras un viaje los usuarios se evalúan, permitiendo mostrar valoraciones imprecisas como las de los criterios Seguridad y Confort. (b) Estas valoraciones alteran colaborativamente el valor del *karma* de cada implicado.

Referencias

- [1] F. Herrera, S. Alonso, F. Chiclana, E. Herrera-Viedma: *Computing With Words in Decision Making: Foundations, Trends and Prosp.* FUZZY OPTIM DECIS MA 8:4 (2009).
- [2] R.M. Rodríguez, L. Martínez, F. Herrera: *Hesitant Fuzzy Linguistic Term Sets for Decision Making.* IEEE Transactions on Fuzzy Systems. vol.10(1) (2012) 109–118.
- [3] Rosana Montes (coord.) *Plataforma de economía colaborativa para la movilidad entre España y Senegal.* Proyecto CEI-BioTIC. Ref. TIC11-2015 (2015).
- [4] F. Herrera, L. Martínez: Autor(es): *A 2-tuple fuzzy linguistic representation model for computing with words.* IEEE Transactions on Fuzzy Systems. 8(6) (2000) 746–752.

A Novel Consensus Model for Group Decision Making with Hesitant Fuzzy Linguistic Information

R.M. Rodríguez¹, Luis Martínez²

¹ Dept. of Computer Science and A.I., University of Granada, Granada, Spain

rosam.rodriguez@decsai.ugr.es

² Dept. of Computer Science, University of Jaén, Jaén, Spain

martin@ujaen.es

Group Decision Making (GDM) is a usual process in companies and administration in which a group of decision makers try to obtain a joint solution for a complex decision problem by taking into account their different points of view situation [2]. Notwithstanding, in principle group decisions should be better accepted than decisions made by a single decision maker because they try to include several viewpoints, sometimes the decision processes do not consider the agreement in the solution, therefore such solutions can fail in their goal. To overcome such a problem, a consensus reaching process is added to GDM processes (See Fig 1), in which an iterative discussion process supervised by a human figure so-called *moderator* tries to guide by providing advice to experts in order to obtain agreed solutions by all decision makers participating in the decision problem [7].

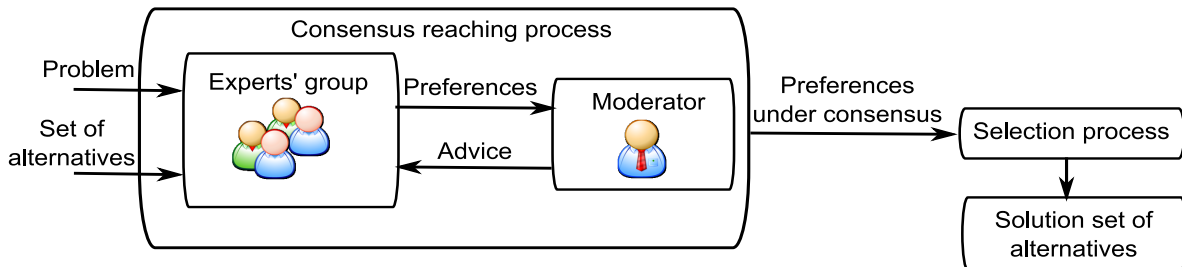


Figura 1: Consensus and GDM scheme

The complexity of GDM problems are often due to the uncertainty related to the imprecision and vagueness of the meaning of the decision situation that is modelled by linguistic descriptors [3]. Different linguistic *consensus* models have successfully dealt with these GDM problems [4, 5]. However, recently it has been pointed out that in GDM problems dealing with linguistic information may be necessary to offer a higher flexibility to experts for eliciting their preferences to manage mainly their hesitancy about linguistic assessments when a single linguistic term does not adjust enough to their knowledge/preference [6]. This contribution provides below a novel *consensus* model for GDM problems dealing with Hesitant Fuzzy Linguistic Term Sets (HFLTS) that have been proposed to deal with hesitancy in linguistic GDM problems.

This work is partially supported by the Research Project TIN2015-66524, Spanish Ministry of Economy and Finance, Postdoctoral Training (FPDI-2013-18193) and ERDF..

1. *Determining Group Decision Problem*: it determines the GDM problem defining the alternatives, experts and domains of expression for experts' preferences.
2. *Gathering Preferences*: Each expert provides her preferences by means of a preference relation P_i , with linguistic terms and linguistic comparative expressions.
3. *Making Information Uniform*: Preferences (fuzzy membership functions representing the linguistic terms and the comparative linguistic expressions) are conducted into a unified expression domain by using fuzzy sets [1].
4. *Computing Consensus Degree*: It computes the degree of agreement or consensus degree, amongst decision makers [2], measured as a value in $[0,1]$ such that the more value the better.
5. *Consensus Control*: Consensus degree is compared with a consensus threshold $\mu \in [0, 1]$, established a priori by the group.
6. *Advice Generation*: When consensus required is not achieved, decision makers are advised to modify their preferences to make them closer to each other and increase the consensus degree in the next round of the Consensus Reaching Process.

This contribution has presented a consensus model that deals with HFLTS information in GDM to achieve agreed solutions in GDM problems defined in qualitative settings in which experts can hesitate when eliciting their preferences.

Referencias

- [1] F. Herrera, L. Martínez, and P.J. Sánchez. Managing non-homogeneous information in group decision making. *European Journal of Operational Research*, 166(1):115–132, 2005.
- [2] J. Kacprzyk. Group decision making with a fuzzy linguistic majority. *Fuzzy Sets and Systems*, 18:105–118, 1986.
- [3] L. Martínez, D. Ruan, F. Herrera, E. Herrera-Viedma, and P.P. Wang. Linguistic decision making: Tools and applications. *Information Sciences*, 179(14):2297–2298, 2009.
- [4] F. Mata, L. Martínez, and E. Herrera-Viedma. An adaptive consensus support model for group decision-making problems in a multigranular fuzzy linguistic context. *IEEE Transactions on Fuzzy Systems*, 17(2):279–290, 2009.
- [5] I. Palomares, J. Liu, Y. Xu, and L. Martínez. Modelling experts' attitudes in group decision making. *Soft Computing*, 16(10):1755–1766, 2012.
- [6] R.M. Rodríguez, L. Martínez, V. Torra, Z.S. Xu, and F. Herrera. Hesitant fuzzy sets: state of the art and future directions. *International Journal of Intelligent Systems*, 29(6):495–524, 2014.
- [7] S. Saint and J.R. Lawson. *Rules for Reaching Consensus. A Modern Approach to Decision Making*. Jossey-Bass, San Francisco, 1994.

Bell-type inequalities and fuzzy probability

Agnieszka Siwiec¹, Susana Díaz², Bernard De Baets³, Susana Montes²

¹ *Instytut Matematyki, University of Gdansk, Poland*
agnieszka.siwiec.zaba@gmail.com

² *UNIMODE research unit, University of Oviedo, Spain*
diazsusana@uniovi.es, montes@uniovi.es

³ *KERMIT, Ghent University, Belgium*
Bernard.DeBaets@UGent.be

A Bell-type inequality is an inequality of the type $0 \leq L \leq 1$, where L is a linear combination with real coefficients of probabilities p_i and joint probabilities $p_{ij}, p_{ijk}, \dots, p_{1,\dots,n}$ corresponding to n events. These inequalities involve t-norms when they are translated to the fuzzy sets context. In this contribution we focus on the continuous t-norms that satisfy the third axiom of Zadeh's definition of fuzzy probability and study which of them satisfy Bell-type inequalities I_2^1 and I_3^2 :

$$\text{Inequality } I_2^1 : \quad 0 \leq x + y - T(x, y) \leq 1 \quad (1)$$

$$\text{Inequality } I_3^2 : \quad 0 \leq x - T(x, y) - T(x, z) + T(y, z). \quad (2)$$

Zadeh's definition of probability for fuzzy events includes the axiom of finite additivity that depends on the chosen t-norm and can be equivalently written as follows:

$$T(1 - x, 1 - y) = 1 - x - y \text{ for every } (x, y) \text{ zero-divisors of } T. \quad (3)$$

Montes et al. [3] characterized the continuous t-norms that satisfy this property.

Theorem 1 *A continuous t-norm satisfies Eq. 3 if and only if it is a strict t-norm, the minimum, the Łukasiewicz operator or an ordinal sum of the following type: $T = (\langle a_i, b_i, T_i \rangle)_{i \in I}$ with nilpotent elements ($Z(T) \neq \emptyset$) such that if $\exists i \in I$ satisfying $a_i = 0$, then:*

- $(a_i, b_i) = (0, i)$, where $i \leq \min(1 - s, \frac{1}{2})$ and
- $T(x, y) = T_L(x, y)$ for every $(x, y) \in \tilde{D}$

where $D = Z(T) = \{(x, y) \in (0, \lambda)^2 \mid T(x, y) = 0\}$, $\tilde{D} = \{(x, y) \mid (1 - x, 1 - y) \in D\}$ and $s = \sup\{x + y \mid (x, y) \in D\}$.

We next explore what t-norms among those considered in Theorem 1 satisfy the Bell-type inequality I_2^1 (Eq. 1).

Observe that since every t-norm is bounded by the minimum, the left inequality holds for any t-norm since $x + y - T(x, y) \geq x + y - \min(x, y) = \max(x, y) + \min(x, y) - \min(x, y) \geq 0$. The right inequality is equivalent to $T(x, y) \geq x + y - 1$, this is, it is equivalent to $T \geq T_L$. Therefore, inequality I_2^1 becomes $T(x, y) \geq T_L(x, y)$. It is clear that not every strict t-norm satisfies the previous inequality. In fact, the following result characterizes those strict t-norms that satisfy it.

The research reported in this paper has been partially supported by the Spanish Ministerio de Economía y Competitividad under Project TIN2014-59543-P.

Proposition 2 A strict t-norm $T = T_P^\varphi$ satisfies Inequality (1) if and only if $\varphi(x+y-1) \leq \varphi(x)\varphi(y)$ for every (x,y) such that $x+y-1 \geq 0$.

And not every strict t-norm satisfies this inequality. As an example, one can consider Hamacher family of t-norms and check that for $\lambda = 30$, the t-norm does not satisfy I_2^1 .

Thus, right inequality I_2^1 is not fulfilled for any strict t-norm. The same situation holds for ordinal sums. Some but not all the ordinal sums considered in Theorem 1 satisfy Ineq. (1). First we show a positive result.

Theorem 3 If $\forall_i T_i \geq T_L$, then $T = \langle T_i \rangle_{i \in I} \geq T_L$.

The condition presented in the previous theorem is sufficient, but not necessary as the following example shows: Consider the ordinal sum: $T = (\langle 0.5, 0.6, T_{T_L}^\varphi \rangle)$, where $T_{T_L}^\varphi = \sqrt{\max(x^2 + y^2 - 1, 0)}$. This ordinal sum is greater than T_L , but $T_{T_L}^\varphi$ is not greater than T_L . Then, condition $T_i \geq T_L$ does not provide a characterization. Nevertheless it is close.

Theorem 4 Let $T = (\langle a_i, b_i, T_i \rangle)_{i \in I}$ be an ordinal sum. Then $T \geq T_L$ if and only if $T_i(x,y) \geq x+y - \frac{1-a_i}{b_i-a_i}$ for all x,y such that $x+y > 1$ and for all $i \in I$ such that $b_i \in]0.5, 1]$.

As commented above, condition $T_i \geq T_L$ is not a necessary condition to assure that $T = (\langle a_i, b_i, T_i \rangle)_{i \in I}$ satisfies $T \geq T_L$. The equivalence only holds if $T = (\langle a_1, b_1, T_1 \rangle)$ and $b_1 = 1$.

Next we study those t-norms that satisfy inequality I_3^2 (Eq. 2). It is easy to check that this condition is stronger than condition I_2^1 and thus, all the necessary conditions for a t-norm to satisfy I_2^1 are also necessary to satisfy inequality I_3^2 .

It was shown in Theorem 6 of [2] that every quasi-copula satisfies inequality I_3^2 (and inequalities I_2^1 and I_4^4). The same proof can be used to prove that every t-norm satisfying the 1-lipschitz property also satisfies inequality I_3^2 . Moreover, Table 3 in [1] shows that for the most important parametric families of t-norms, those parameters for which condition I_3^2 holds are exactly the same parameters for which the corresponding norms are copulas (in particular, quasi-copulas). Then we could derive from that table that there is some connection between condition I_3^2 and the 1-Lipschitz condition. In [2] we can see an example of a commutative conjunctive that satisfies condition I_3^2 but does not satisfy the 1-Lipschitz property. Then, in general, Property I_3^2 is weaker than 1-Lipschitz. However, what happens for t-norms? Are I_3^2 and 1-Lipschitz properties equivalent for t-norms?

References

- [1] S. Janssens, B. De Baets and H. De Meyer. Bell-type inequalities for parametric families of triangular norms. *Kybernetika* 1, 89–106, 2004.
- [2] S. Janssens, B. De Baets and H. De Meyer. Bell-type inequalities for quasi-copulas. *Fuzzy Sets and Systems* 148, 263–278, 2004.
- [3] I. Montes, J. Hernández, D. Martinetti, S. Montes. Characterization of continuous t-norms compatible with Zadeh's probability of fuzzy events. *Fuzzy Sets and Systems* 228, 29–43, 2013.
- [4] L. Zadeh. Probability measures of fuzzy events. *J. Math. Anal. Appl.* 23, 421–427, 1968.

Presentación de una red social docente que integra un sistema de recomendaciones lingüístico difuso para personalizar el acceso a recursos educativos

Alberto Ching-López², Carlos Porcel¹, Enrique Herrera-Viedma²

¹ *Departamento de Informática, Universidad de Jaén, Jaén, 23071, España*
cporcel@ujaen.es

² *Departamento de Ciencias de la Computación e I.A., Universidad de Granada, Granada, 18071, España*
alch@decsai.ugr.es, viedma@decsai.ugr.es

Las nuevas tecnologías enriquecen la formación con la posibilidad no sólo de difundir información de una forma cómoda y eficiente, sino de dotar a los participantes (profesores, alumnos, expertos, etc.) de herramientas para la comunicación personal y grupal que refuercen la acción tutorial y el aprendizaje colaborativo [4]. En este sentido, las **Redes Sociales** son una de las tecnologías Web que posibilitan y fomentan ese trabajo colaborativo. Se definen como servicios Web que permiten a los usuarios construir un perfil público o semi público dentro de un entorno acotado, establecer una lista de usuarios con los que se comparte una conexión y ver y recorrer su lista de conexiones y las establecidas por otros dentro del sistema [1]. Las características de trabajo colaborativo e intercambio social entre los distintos integrantes del sistema educativo, han permitido que el ámbito docente se convierta en un entorno potencial de aplicación con éxito de las redes sociales [2]. Sin embargo el uso cada vez más extendido de Internet en general y de las redes sociales en particular, provoca que cada día se genere más información y contenidos en este tipo de medios, de manera que el número de recursos a los que podemos acceder va creciendo día a día de forma desmesurada.

En este contexto se plantea la necesidad de herramientas automáticas que ayuden a los usuarios en sus procesos de acceso a la información, en la correcta personalización de sus procesos educativos e identificación de posibilidades de colaboración con otros usuarios en similares circunstancias, con el fin de dinamizar y mejorar los procesos educativos. Para hacer frente a dicho problema, planteamos el uso de un **sistema de recomendaciones** [3, 6]. Son herramientas que facilitan a los usuarios el acceso personalizado a información de su interés, distinguiendo entre lo que puede serles relevante y lo que no. Para conseguirlo, estos sistemas explotan comportamientos previos y similitudes entre los usuarios para predecir nuevas necesidades, habiéndose aplicado con éxito en numerosos ámbitos, incluyendo el educativo [5].

Por otro lado nos enfrentamos a otro problema, que es la gran variedad de representaciones y evaluaciones de la información, aún más acusado cuando los usuarios forman parte del proceso, como es el caso que nos ocupa. El uso de la Teoría de Conjuntos Difusos ha dado muy buenos resultados para modelar información cualitativa y se ha probado su utilidad en

Este trabajo ha sido elaborado con la financiación de los proyectos UJA2013/08/41, TIC-5991 y TIN2013-40658-P.

multitud de problemas. Por ello, adoptamos el **modelado lingüístico difuso multi-granular** que nos ayude a representar y gestionar de forma eficiente la información cualitativa presente en los procesos de comunicación [7].

En este trabajo presentamos **SharingNotes**, una red social docente que ayuda a difundir recursos pedagógicos de forma personalizada. Para ello incorporamos un sistema de recomendaciones lingüístico difuso. El éxito o fracaso de una técnica de generación de recomendaciones depende en gran medida del ámbito de aplicación concreto. Una técnica muy usada es el enfoque colaborativo, que permite a los usuarios compartir sus experiencias previas. Un problema de este enfoque es que para empezar a trabajar correctamente, necesita cierto número de valoraciones (difícil de conseguir). Puesto que el esquema basado en contenidos se basa en las características de los recursos y las preferencias y/o necesidades de los usuarios, puede empezar a trabajar con menos información. Por ello, aplicamos un enfoque híbrido entre ambos para solucionar los problemas que presentan por separado, pero aprovechando sus ventajas. Además, al trabajar sobre una red social, se nos abre el abanico a la hora de establecer nuevas técnicas para la generación de recomendaciones. En este sentido, otro enfoque que hemos adoptado para mejorar las recomendaciones generadas, es explotar la confianza derivada de la red de colaboradores o la reputación de un usuario concreto [8]. La versión actual está operativa y es accesible en la siguiente url: <http://sharingnotes.ujaen.es/>

Referencias

- [1] D.M. Boyd, N.B. Ellison: *Social network sites: Definition, history, and scholarship*. Journal of Computer-Mediated Communication. **13(1)** (2007) article 11.
- [2] J.J. DeHaro: *Las redes sociales aplicadas a la práctica docente*. Revista DIM: Didáctica, Innovación y Multimedia. **13** (2009).
- [3] A. Tejeda-Lorente, C. Porcel, J. Bernabé-Moreno, E. Herrera-Viedma: *REFORE: A recommender system for researchers based on bibliometrics*. Applied Soft Computing. **30** (2015) 778–791.
- [4] A. Popovici, C. Mironov: *Students' perception on using eLearning technologies*. Procedia - Social and Behavioral Sciences. **180** (2015) 1514–1519.
- [5] M. Goga, S. Kuyoro, N. Goga: *A recommender for improving the student academic performance*. Procedia - Social and Behavioral Sciences. **180** (2015) 1481–1488.
- [6] R. Burke, A. Felfernig, M.H. Göker: *Recommender systems: An overview*. AI Magazine. **32(3)** (2011) 13–18.
- [7] F. Herrera, L. Martínez: *A model based on linguistic 2-tuples for dealing with multigranularity hierarchical linguistic contexts in multiexpert decision-making*. IEEE Transactions on Systems, Man and Cybernetics. Part B: Cybernetics. **31(2)** (2001) 227–234.
- [8] C. Martinez-Cruz, C. Porcel, J. Bernabé-Moreno, E. Herrera-Viedma: *A Model to Represent Users Trust in Recommender Systems using Ontologies and Fuzzy Linguistic Modeling*. Information Sciences. **311** (2015) 102–118.

Estudio del aprovechamiento de oportunidades de aprendizaje matemático en un aula de 3º de la ESO

Miquel Ferrer¹, Josep M^a Fortuny¹, Itziar García-Honrado²

¹ *Universidad Autónoma de Barcelona*

Miquel.Ferrer.Puigdellivol@uab.cat,

JosepMaria.Fortuny@uab.cat

² *Universidad de Oviedo*

garciaaitziar@uniovi.es

En esta comunicación mostraremos la evolución del aprendizaje seguido por los alumnos a lo largo de una discusión en gran grupo. En trabajos anteriores nos centramos en identificar las oportunidades de aprendizaje matemático surgidas en la discusión en gran grupo [1], motivo por el cual en el presente trabajo es de interés estudiar cómo los alumnos aprovechan estas oportunidades de aprendizaje. Para ello, analizaremos una discusión en gran grupo de un problema de semejanza resuelto por alumnos de tercer curso de la Educación Secundaria Obligatoria (ESO) implementado en el aula utilizando GeoGebra, un software de geometría dinámica. La resolución de este problema requiere la realización de un giro y una homotecia, no importando el orden en el que se efectúen ambas transformaciones.

En la discusión en gran grupo de este problema, guiada por la profesora Sara, se ha identificado, entre otras, la oportunidad de aprendizaje conceptual: Identificar las transformaciones geométricas con GeoGebra que resultan de componer un giro con una homotecia.

Para analizar el aprovechamiento de esta oportunidad de aprendizaje matemático, nos centraremos en la evolución del aprendizaje de una alumna en concreto, Alba, aunque la forma de proceder es análoga con los demás estudiantes. Distinguimos tres fases en el desarrollo de la actividad: la primera fase en la que los alumnos entran en contacto con el problema llevando a cabo una resolución por parejas; la segunda fase, que se corresponde con la propia discusión; y una tercera fase en la que los alumnos hacen una reflexión escrita e individual de forma concluyente.

En todas las fases se aplica una rúbrica construida ex profeso para comprobar el aprovechamiento de la oportunidad de aprendizaje indicada anteriormente. Se desarrolla el hecho de identificar las transformaciones con averiguar los elementos que las caracterizan, en el caso del giro: centro y ángulo, y en el caso de la homotecia: centro y razón.

Se hace una graduación de los niveles en los que se haga esta averiguación:

- B=Bajo= lo menciona.
- M=Medio= hace una explicación sin reflexión.
- A= Alto= hace una explicación propia basada en pruebas empíricas o deducción formal.

Se permite asignar un grado entre 0 y 1 a cada uno de los niveles. Con esto se consigue una evaluación a través de un conjunto borroso discreto en cada una de las fases.

En la primera fase, Alba consigue resolver el problema de una forma intuitiva, ayudándose de traslaciones y evitando un cálculo reflexivo de los centros de cada transformación.

En la segunda fase obtenemos del diálogo transcurrido en la propia discusión en gran grupo que Alba argumenta la necesidad de utilizar al menos dos transformaciones, haciendo hincapié en la homotecia. El hecho de hablar de “al menos dos” indica la posibilidad de eliminar la traslación. Además, Alba introduce la posibilidad de añadir un deslizador en GeoGebra para calcular el valor del ángulo de giro e identifica la equivalencia entre una homotecia de razón -1, un giro de 180° y una simetría central. Por este motivo consigue una mayor profundidad en su conocimiento de los elementos característicos de las transformaciones.

En la tercera fase Alba elimina de su construcción la traslación, indica el centro de la homotecia, calcula el centro de giro razonadamente y utiliza GeoGebra para demostrar de forma empírica tanto el ángulo de giro como la razón de homotecia. Lo único a lo que no hace alusión en su protocolo escrito es a la necesidad de utilizar el giro, aunque evidentemente lo emplea. Recogemos los resultados obtenidos en cada una de estas fases en la Tabla 1.

Tabla 1: Valoraciones de la resolución de Alba.

	Bajo			Medio			Alto		
Identifica...	Fase1	Fase2	Fase3	Fase1	Fase2	Fase3	Fase1	Fase2	Fase3
- Giro	0.7	0.7	0.7	0.3	0.3	0.3			
- Homotecia	0.5			0.5	0.5	0.2		0.5	0.8
- Centro de giro	1	0.7			0.3				1
- Centro de homotecia	1	0.7			0.3	1			
- Ángulo de giro				0.5	0.2		0.5	0.8	1
- Razón de homotecia				0.2	0.2		0.8	0.8	1

Con las representaciones de los tres conjuntos borrosos de cada una de las fases:

- 3.2/6 de B, 1.5/6 de M y 1.3/6 de A, o equivalentemente $0.53/1 + 0.25/2 + 0.22/3$
- 2.1/6 de B, 1.8/6 de M y 2.1/6 de A, o equivalentemente $0.35/1 + 0.3/2 + 0.35/3$
- 0.7/6 de B, 1.5/6 de M y 3.8/6 de A, o equivalentemente $0.11/1 + 0.25/2 + 0.64/3$

Concluimos que Alba ha hecho un aprovechamiento creciente de esta oportunidad de aprendizaje llegando a un nivel medio/alto, pudiendo representar los centros de gravedad de cada conjunto borroso (1.68, 2, 2.52) para ver gráficamente su evolución (véase la Figura 1).

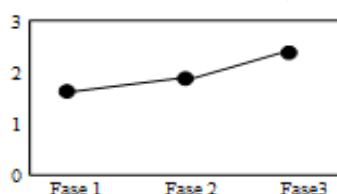


Figura 1: Evolución de Alba

Tras introducir las percepciones de la resolución de los alumnos, en futuros trabajos las agregaremos siguiendo el esquema GLMP [2]. Este esquema nos permitirá conseguir un resumen lingüístico del aprovechamiento de las oportunidades de aprendizaje matemático.

Referencias

- [1] Miquel Ferrer, Josep M^a Fortuny, Laura Morera: *Efectos de la actuación docente en la generación de oportunidades de aprendizaje matemático*. Enseñanza de las Ciencias: Revista de Investigación y Experiencias Didácticas. **32-3** (2014) 385–405.
- [2] Gracián Triviño, Michio Sugeno: *Towards linguistic description of phenomena*. International Journal of Approximate Reasoning. **50- 1** (2013) 22–34.

Propuesta interdisciplinar en Educación Secundaria: creación de un modelo con lógica fuzzy para la hormonal del ciclo ovárico.

Edith Herrera¹, Itziar García-Honrado²

¹ Universidad Autónoma de Barcelona

Departamento de Didáctica de las Matemáticas y las Ciencias Experimentales

Edithherrera@uab.cat

Eherrerera@ubiobio.cl

Departamento de Estadística, I.O. y Didáctica de las Matemáticas.

Universidad de Oviedo.

garciaaitziar@uniovi.es

El proceso de regulación hormonal es un contenido de gran importancia en la Educación secundaria, especialmente relevante resulta el funcionamiento hormonal en el ciclo óvarico, por sus implicaciones en el proceso de fecundación.

En este artículo proponemos un modelo que ayuda a la comprensión de la temática mediante el trabajo colaborativo entre Biología, Matemáticas y Lógica (perteneciente, en la actualidad, a la asignatura de Filosofía) y Tecnología. Hacer conexiones interdisciplinarias [1] incide en los lazos entre las disciplinas y facilita un aprendizaje integral y profundo. Al efectuarlas, los estudiantes refuerzan y expanden sus conocimientos y acceden a información y a diversos puntos de vista, con una visión integral de estos como un “todo”.

Desde el punto de vista de la biología nos interesa que el alumno comprenda cómo funciona el mecanismo del eje hipotálamo-Hipófisis, en el estudio del ciclo menstrual de la mujer, ya que esta regulación se realiza gracias al control que ejercen dos hormonas hipofisiarias, al comienzo la FSH en la etapa preovulatoria y posteriormente la hormona LH en la etapa post ovulatoria.

El mecanismo de control (fuzzy) lo realiza la hipófisis cuando detecta en la sangre aumento o disminución en los niveles de estrógenos y progesterona circulando. Si bajan los niveles de ambas hormonas el ciclo se inicia y niveles altos de ellas desencadenan la salida del óvulo, la ovulación. Niveles medios se producen antes y después de la ovulación, si no hay fecundación del óvulo los niveles vuelven a bajar y el ciclo ovárico se vuelve a iniciar.

La relación con Matemáticas, Lógica y Tecnología viene a través de creación de un modelo de un prototipo que permita, una vez conocidos los niveles de estrógenos y progesterona en sangre, saber el momento del ciclo ovárico: inicio del ciclo, ovulación o post ovulación. Para ello, se darán indicaciones de lo que es un controlador fuzzy (véase Figura 1).

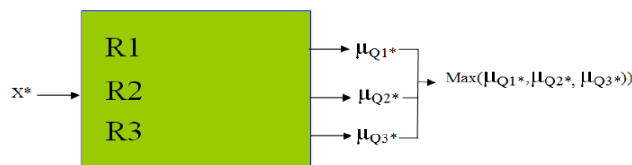


Figura 1: Esquema de controlador fuzzy

El estudio del modelo biológico nos proporciona amplios rangos de concentraciones de hormonas, los cuales dependen de muchas variables en el cuerpo de una mujer, por lo que no funciona con datos precisos, de ahí la necesidad de modelar estas concentraciones a través de conjuntos fuzzy. Además, los alumnos de Educación Secundaria han de conocer el modelo biológico en términos de la existencia de un número bajo/medio/alto de hormonas.

Por lo tanto, para la construcción del modelo indicamos llevar a cabo estas tareas:

1. Crear de reglas del tipo Si/entonces, en las que el antecedente haga alusión a los niveles de estrógenos y progesterona y el consecuente se refiera al momento del ciclo ovárico.
2. Modelar mediante conjuntos fuzzy lineales a través de las tablas de concentraciones los predicados bajo/medio/ alto para los niveles de la hormonas involucradas.
3. Explicar la Regla Composicional de Inferencia de Zadeh [2] para lo que se hará especial énfasis en el uso del Modus Ponens. Así mismo se trabaja la traducción en términos matemáticos de la doble implicación por la función Producto/Mínimo y la traducción de la intersección por la función mínimo, lo que corresponde con el método de Larsen y Mandani, respectivamente.
4. Computar el modelo para una entrada concreta, calculando la forma analítica de los conjuntos fuzzy para posteriormente realizar el cálculo de mínimos, producto y supremo.
5. Analizar las diferencias entre los modelos de Mamdani y Larsen y los distintos diseños de conjuntos borrosos que se den en el aula.

La construcción de un modelo matemático, les hará concebir las matemáticas como una ciencia que puede aplicarse en casos prácticos, así como detenerse en el diseño de los componentes del modelo, lo que les ayudará a entender el comportamiento de las funciones mínimo y producto.

Además a través de este modelo Fuzzy de aprendizaje, el alumno puede profundizar sus conocimientos sobre el mecanismo regulación hormonal y comprender su relación con los anticonceptivos orales., test de embarazos, los cuales están basados en esta lógica de funcionamiento. Así la Matemática nos ofrece un campo riquísimo para proponer situaciones desafiantes en las que no se objetiva exclusivamente la fijación de contenidos, sino que se incentiva la creatividad y la originalidad paralelamente a la observación, análisis y razonamiento lógico.

Referencias

[1] Bridi, J. H., de Fraga Sant, M., Geller, M., & da Silva, J.: *El uso de la actividad de laboratorio de biología para la enseñanza de matemáticas en los años iniciales: Una estrategia interdisciplinaria de enseñanza y aprendizaje* Ensaio Pesquisa em Educação em Ciências. (2010) **12(3)**, 131-150.

[2] Enric Trillas e Itziar García Honrado: *La Regla Composicional de Zadeh, una lección para principiantes*. Actas del XIV Congreso Español sobre Tecnologías y Lógica Fuzzy. (2008) 417–422.

Proyecto Construcción de conocimiento matemático escolar: discurso del profesor y actividad de enseñanza EDU2015-65378-P.

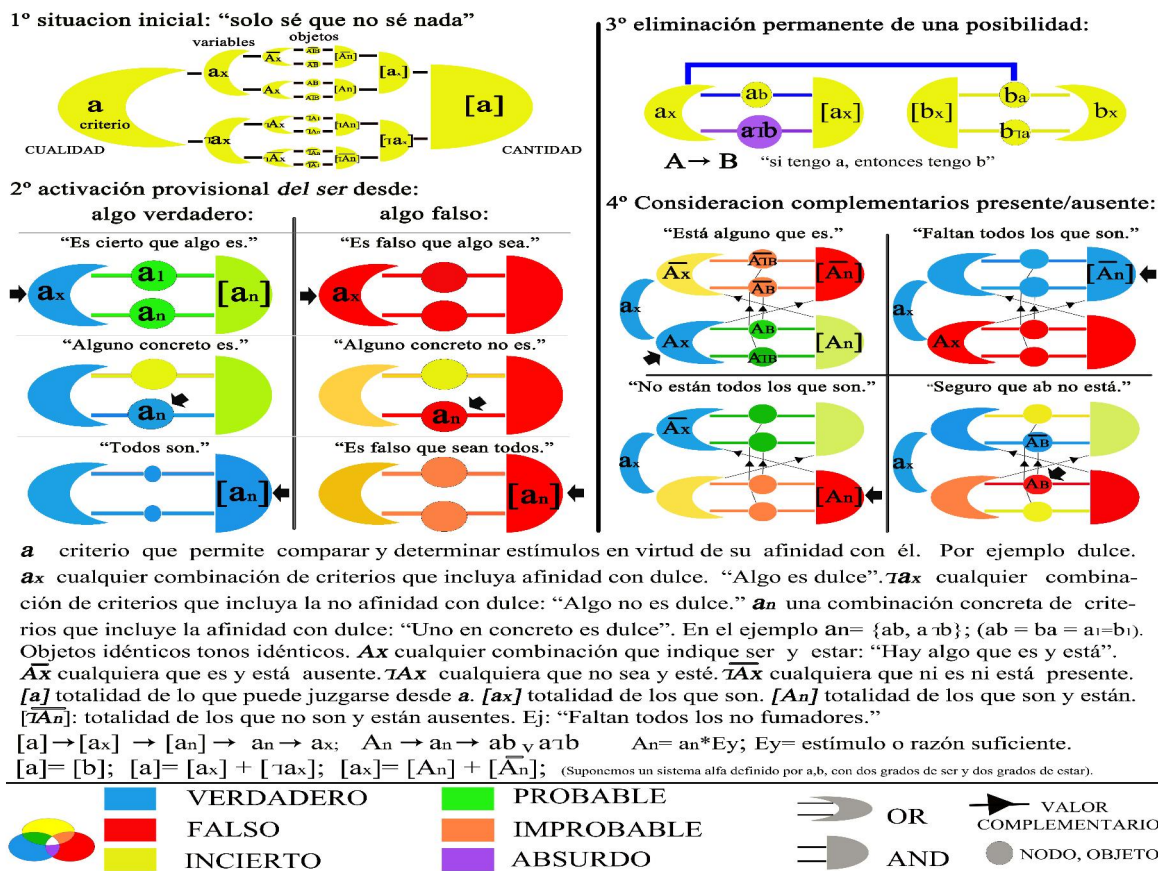
Estructura formal de los sistemas de creencias desde el Diagrama de Marlo

Marcos Bautista López Aznar

C/Fuenteheridos, 25 2d - 21002 Huelva. España

pensamosdistintorazonamosigual@gmail.com

Considerar la lógica como la estructura formal con la que la subjetividad organiza sus creencias a fin de mantener expectativas adaptadas y comunicables facilita su didáctica. Así, para un sujeto, una afirmación puede ser en distintos grados verdadera, falsa, probable, incierta, o absurda, siendo la incertidumbre un estado habitual de los sistemas cognitivos. Estos, integran información de subsistemas capaces de discriminar qué grado de afinidad mantienen los estímulos con criterios relevantes para la vida. Cada subsistema constituye una estructura formal de nodos (objetos) y conexiones (expectativas) que propaga la activación provocada por estímulos. Así, además de grados de afinidad a un criterio (ser), debemos considerar grados de activación o disponibilidad de los objetos (estar), que serán reales en la medida que sean capaces de modificar los estados del sistema. Las variables usadas para discriminar grados de semejanza deben permitir decisiones rentables. Cualquier combinación coherente y conveniente de variables supone un objeto válido en teoría.



Las conectivas, especialmente el bicondicional y el condicional, eliminan posibilidades de forma permanentemente, determinando la propagación de la activación por los nodos permitidos. Y aunque sea útil reducir así la incertidumbre estableciendo lo que es razonable esperar, los sistemas se rebelan ante cualquier intento de imponer una expectativa como verdad absoluta.

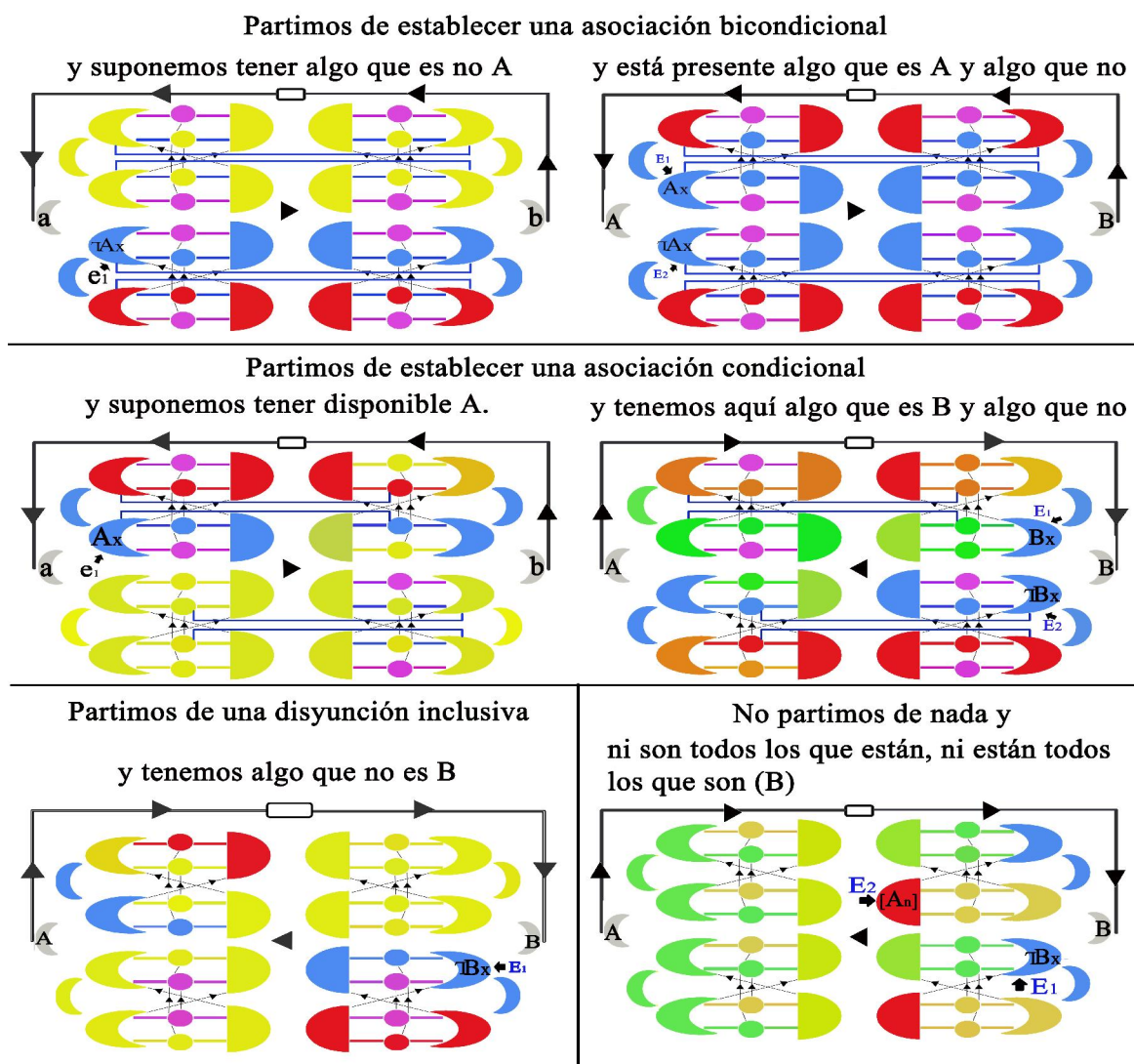


Figura 2. Ejemplos de proposiciones con distinto grado de incertidumbre.

References

- Marcos B. López Aznar: Adiós a $bArbArA$ y VENN. Lógica de predicados en el diagrama. Paideia, No 102 (2015) 35-52.
- Marcos B. López Aznar: Cálculo lógico de modelos proposicionales: la revolución del silogismo en el Diagrama de Marlo. Círculo Rojo. 2014

Extensión borrosa del método de aprendizaje de Leitner

José Otero¹, Rosario Suárez¹, Luís Junco¹, Luciano Sánchez¹

¹ *Edificio Departamental Oeste, Campus de Viesques, Gijón, Asturias.*
 jotero@uniovi.es, mrsuarez@uniovi.es, lajunco@uniovi.es,
 luciano@uniovi.es

Las tarjetas didácticas son un recurso educativo común en la enseñanza de ciertas disciplinas o partes de disciplinas. Un mazo de tarjetas didácticas consiste en una serie de tarjetas que contienen información que debe de ser aprendida. Normalmente contienen una pregunta y una respuesta, en una forma libre: puede tratarse de una pregunta típica, completar una palabra que falta en una frase, reconocer un determinado objeto, una fórmula, etc. El estudiante lee una de las caras de la tarjeta e intenta recordar la respuesta (recuerdo activo). Tras expresar verbalmente o mentalmente la respuesta, la coteja con el dorso de la tarjeta (repaso pasivo). Dependiendo de si la respuesta ha sido correcta o no, el estudiante podría separar el mazo de tarjetas en dos grupos, las que ha respondido correctamente y las que no. Posteriormente, optimizaría su tiempo dedicándose en exclusiva a repasar las tarjetas cuyas preguntas no ha recordado y de nuevo separarlas en los dos grupos mencionados. Este esquema puede perfeccionarse usando más grupos e intercalando las tarjetas de una forma cuidadosa, en lo que se conoce como el método o sistema Leitner [1]. Trabajos posteriores han profundizado en esta técnica dando lugar a lo que se conoce como estudio mediante repeticiones espaciadas [2] [3], que asigna intervalos de repetición más cortos a las preguntas que se fallan más a menudo.

La aparición de aplicaciones informáticas (móviles, de escritorio y web) han hecho muy manejable este sistema en su vertiente digital, añadiendo prestaciones físicamente imposibles, como tarjetas de más de dos caras o multimedia.

Este método de aprendizaje ha demostrado ser efectivo en disciplinas que se apoyan en conocimiento de tipo explícito y se ha aplicado a ciertas áreas de la informática [4] sin embargo presenta ciertos inconvenientes cuando se aplica a disciplinas que manejan un conocimiento de tipo procedural. Un ejemplo de este tipo de disciplinas es la programación usando un determinado lenguaje [5]. Entre los materiales que se usan habitualmente en la enseñanza de esta disciplina se encuentran las baterías de fragmentos de código que el estudiante debe analizar, con el fin de contestar a una serie de preguntas sobre la salida del programa, su corrección, etc. El método de Leitner no es adecuado en este caso porque el estudiante tendería a memorizar la solución del ejemplo concreto que se le muestra, en lugar de analizar (cada vez que se le presenta) el código con el fin de averiguar la respuesta. Una posible solución a este problema sería proveer al estudiante de un gran número variaciones del mismo ejemplo, siendo entonces una aproximación similar a la del testeo frecuente [6]. El inconveniente de trasladar directamente esta metodología al método de Leitner consiste en que la variación de

Este trabajo ha sido subvencionado mediante los proyectos MINECO-15-TIN2014-56967-R y FC-15-GRUPIN14-073.

cada ejemplo transformándolo en n distintos multiplica también por ese mismo factor el tiempo de estudio y, además, podría seguir dándose el caso de que el estudiante memorice cada ejemplo por separado. Por otra parte, sería posible atribuir períodos de repetición distintos a variaciones de la misma pregunta, lo que hace difícil evaluar el grado de conocimiento del concepto relacionado. Con el fin de solventar estos dos problemas, en esta ponencia proponemos una metodología según la cual:

- Cada concepto de programación que se desea dominar se representa por un esqueleto de código o patrón parametrizable, de modo que se puedan obtener automáticamente variaciones del mismo, por ejemplo con distintos valores de constantes, nombres de variables o distintos operadores en las expresiones.
- El grado de dominio de cada concepto se representa mediante un conjunto borroso, definido por el intervalo de repetición asociado a cada variación obtenida a partir del patrón correspondiente.
- Si en una sesión de estudio está programada una variación concreta de un ejemplo, ésta se sustituye por una nueva instancia generada mediante el patrón correspondiente, con el fin de obligar al estudiante a analizar el código de forma efectiva e impidiendo memorizar la respuesta de una instancia concreta.

Se presentan resultados preliminares en los que se comparan los tiempos de estudio y los resultados de aprendizaje obtenidos usando el método de Leitner convencional y la variante basada en el uso de patrones asociados a conceptos y representación borrosa del grado de conocimiento. En los experimentos se ha usado un modelo simplificado de estudiante con componentes de memorización y adquisición de conocimiento procedural parametrizables.

Referencias

- [1] Sebastian Leitner. So lernt man lernen. Der Weg zum Erfolg. Freiburg i. Br, 1972.
- [2] Lisa K. Son y Dominic A. Simon. Distributed Learning: Data, Metacognition, and Educational Implications. *Educational Psychology Review*, 24(3), 379-399, Septiembre 2012.
- [3] A. W. Melton. The situation with respect to the spacing of repetitions and memory. *Journal of Verbal Learning and Verbal Behavior*, 9, 596-606, 1970.
- [4] W. Abramowicz y P.A Wozniak. The educational technology of repetition spacing in information systems development. Wojkowski et al. *Systems development Methods for the Next Century*. Plenum Press, New York, 1997, 341-344.
- [5] A. Robins, J. Rountree y N. Rountree. Learning and teaching programming: A review and discussion. *Computer Science Education*, 13(2):137-172, 2003.
- [6] Roediger, H. L., Putnam, A. L. y Smith, M. A. Ten benefits of testing and their applications to educational practice. In J. Mestre y B. Ross (Eds.), *Psychology of learning and motivation: Cognition in education*, 1-36. Oxford: Elsevier, 2011.

EvalGRAPHS: herramienta para implementar evaluadores inteligentes. Metodología de estudio del comportamiento de los evaluadores

**M. Gloria Sanchez-Torrubia¹, Carmen Torres-Blanc¹,
Silvia García Guerrero¹, Ruben Gonzalez-Martin¹**

¹ *Universidad Politécnica de Madrid*
gsanchez@fi.upm.es, ctorres@fi.upm.es,
silvis.garcia@gmail.com, rubengonzmartin@gmail.com

El Modelo Granular Lingüístico de la Evaluación del aprendizaje humano [1] ha sido desarrollado a partir del paradigma *GLMP* [1], basado en la computación con palabras y percepciones, con el objetivo de diseñar un modelo teórico de representación del aprendizaje que pueda aplicarse a la creación de sistemas (*GLMPs*) que emulen la evaluación realizada por un profesor, en base a sus criterios, y expresen dicha evaluación mediante un informe en lenguaje natural. Un *GLMP* es una estructura jerárquica que organiza y procesa datos mediante inferencia borrosa y modela un fenómeno mediante percepciones o unidades de conocimiento. Estos sistemas utilizan como datos de entrada los generados automáticamente por un entorno de aprendizaje y producen como salida un informe de evaluación que incluye una calificación numérica. De esta manera se puede compaginar la evaluación formativa, comparando logros y objetivos de aprendizaje, que realiza un profesor, con la rapidez y eficacia de un sistema automatizado. Como aplicación del modelo general, hemos diseñado varios sistemas que evalúan, en base a criterios de un profesor, unos los *quizzes* de *Moodle* y otros las simulaciones visuales de algoritmos de grafos hechas en el entorno de aprendizaje *GRAPHS* [2].

Relacionado con el entorno *GRAPHS* hemos desarrollado una herramienta, a la que hemos llamado *EvalGRAPHS*, que ayuda en el diseño y lleva a cabo la ejecución de estos últimos sistemas. La aplicación está compuesta por dos partes: la primera dedicada a la configuración, la salvaguarda, la recuperación y la modificación de los evaluadores que serán implementados por el sistema, y la segunda dedicada a la evaluación, mediante uno de los evaluadores previamente configurados, a partir de los *log* de interacción generados por *GRAPHS* al realizar la simulación de un algoritmo [2]. Dicha evaluación se expresa mediante un informe de evaluación formativa, en lenguaje natural, que describe el nivel de logro del alumno en cada uno de los objetivos de aprendizaje fijados por el profesor y, además, la calificación numérica de dicha simulación.

Cuando se crea una nueva configuración, la herramienta genera un archivo *XML* que almacena todos los datos que la describen: las etiquetas que corresponden a cada uno de los datos de entrada, la estructura de percepciones del *GLMP*, sus etiquetas lingüísticas, las plantillas de generación de texto, las reglas que forman cada uno de los motores de inferencia que obtienen la información de una percepción en función de la de sus subordinadas, la *t-norma* y la *t-conorma* que se utilizan en las funciones de agregación, etc.

Cuando, una vez cargada una configuración concreta para utilizar la herramienta como evaluador, la aplicación pide el *log* de interacción que se desea evaluar, extrae de él los datos de entrada al sistema y genera el informe correspondiente a la simulación evaluada. Además, la aplicación también permite la introducción de datos de forma manual, con el fin de probar las configuraciones que están siendo implementadas.

En este trabajo presentamos una metodología para comprobar la coherencia, el cumplimiento de los objetivos de diseño y el funcionamiento local de un *GLMP*: el estudio de las superficies de comportamiento [3]. Estas superficies representan la salida desborrosificada de una percepción en función de las entradas, también desborrosificadas, a dos de sus percepciones subordinadas. A continuación se describen de forma esquemática los detalles de diseño que se pueden detectar mediante el estudio de estas superficies.

- Influencia de la forma de las etiquetas lingüísticas sobre el funcionamiento del sistema.
 - El conjunto de etiquetas lingüísticas de cada percepción debe ser ortogonal. En caso contrario aparecerán picos y zonas no definidas en las superficies.
 - El uso de etiquetas con forma más *crisp* produce escalones.
- Influencia de la asignación de consecuentes.
 - Como nuestros sistemas se diseñan para la evaluación del aprendizaje, las superficies de comportamiento siempre deben ser crecientes o decrecientes.
 - Crecientes, cuando se agregan aciertos para medir el nivel de acierto o se agregan errores para medir el nivel de error.
 - Decrecientes, cuando se agregan errores para medir nivel de acierto.
 - Las reglas deben ser coherentes, en caso contrario aparecerán zonas de no convexidad.
 - Aunque en nuestros sistemas no pueden existir cambios de comportamiento, en otros *GLMPs* pueden observarse cambios de comportamiento debidos al modelo.
- Validación de que se han definido todas las reglas y usado todas las etiquetas.
 - Si faltan reglas por definir, en la superficie de comportamiento aparecerá una zona no definida (agujero).
 - Si alguna de las etiquetas de la percepción de salida no ha sido asignada en ninguno de los consecuentes de las reglas, aparecerá un escalón en la superficie.

Los sistemas descritos han sido desarrollados, probados según la metodología descrita y mediante la comparación, sobre casos reales, de los resultados generados por el sistema con los facilitados por el profesor cuyos criterios emulan. Además han sido utilizados en el aula con resultados muy satisfactorios.

Referencias

- [1] Sanchez-Torrubia, M. Gloria, Torres-Blanc, Carmen, y Trivino, Gracian: *An approach to automatic learning assessment based on the computational theory of perceptions*. Expert Syst. Appl. **39**, 15 (2012) 12177–12191.
- [2] Sanchez-Torrubia, M. Gloria, Torres-Blanc, Carmen, y Escribano-Blanco, Sonia: *GRAPHS: a learning environment for graph algorithm simulation primed for automatic fuzzy assessment*. En Proc. 10th Koli Calling. (2010) 62–67.
- [3] Sanchez-Torrubia, M. Gloria: *Especificaciones eMathTeacher, creación del Modelo Granular Lingüístico de la evaluación del aprendizaje humano, y diseño de sistemas que implementan estos conceptos*. UPM, Madrid. (2016).

Este trabajo ha sido financiado en parte por UPM.

Un software de modelado de color basado en lógica difusa como herramienta docente para aprender el concepto de imprecisión del color de manera visual

Jose Manuel Soto-Hidalgo¹, Pedro Manuel Martínez-Jiménez², Jesús Chamorro-Martínez², Daniel Sánchez-Fernández²

¹ Dpto. Arquitectura de Computadores, Electrónica y Tecnología Electrónica. Universidad de Córdoba

jmsoto@uco.es

² Dpto. Ciencias de la Computación e Inteligencia Artificial. Universidad de Granada
{pedromartinez, jesus, daniel}@decsai.ugr.es

El color es una característica fundamental y representativa del contenido visual en imágenes. Computacionalmente, el color se representa mediante un sistema vectorial, también conocido como espacio de color. Sin embargo, es sabido que el ser humano usa *categorías* de color a la hora de describirlo, si bien no existe una correspondencia directa entre la representación del color en un ordenador y los términos que el ser humano utiliza para identificarlos (problema conocido como “hueco semántico”, del término inglés “semantic gap”).

En este contexto, el **modelado del color** es un problema importante a la vez que complejo debido, entre otros motivos, a que el color es, en general, **impreciso** (no es posible establecer una frontera clara que delimite unos colores de otros), **subjetivo** (no todas las personas distinguen o nombran de la misma forma los colores) y **dependiente del contexto** (un mismo color puede tener distintos significados en distintos ámbitos). En la Figura 1, formada por un degradado lineal entre dos colores, se puede observar la vaguedad asociada al color, donde no es inmediato establecer un límite preciso que separe los dos colores, sino que la frontera entre ambos es borrosa. Igualmente, hay colores en el degradado, especialmente en la zona central, para los cuales un ser humano no sabría indicar a cuál de los dos nombres de color (asociados a los colores de los extremos) correspondería (posiblemente a más de uno, con un cierto grado en cada caso). Es más, si hubiese que nombrar dichos colores, la respuesta podría ser diferente en función del sujeto y el contexto (por ejemplo, un enólogo podría nombrarlos como “bermellón” y “púrpura”, mientras que un frutero podría usar términos como “fresa” y “berenjena”).



Figura 1: Imagen de un degradado entre dos colores

Por un lado, en [1] abordamos la imprecisión, la subjetividad y la dependencia del contexto asociada al color introduciendo los conceptos de *color difuso* y *espacio de color difuso*. También proponemos una metodología para obtener automáticamente conjuntos difusos que definen colores difusos. Cada color difuso comprende una porción del espacio de color con fronteras borrosas, de manera que cada subconjunto difuso de colores corresponda a uno de los colores que usamos. Por otro lado, en [2] proponemos un software de código abierto, llamado **JFCS** (Java Fuzzy Color Space), que proporciona una solución adecuada a los problemas de imprecisión, subjetividad y dependencia del contexto asociados al color mediante el modelado del color en base a colores difusos. *JFCS* es portable y puede ejecutarse en cualquier máquina independientemente del sistema operativo.

JFCS, aparte de proponer una solución al modelado de color, puede ser utilizado como herramienta didáctica y docente para entender el problema de la imprecisión del color, además de comprender cómo la lógica difusa es adecuada para resolver este problema. *JFCS* integra herramientas gráficas que permiten describir, de manera simple y visual, los colores presentes en imágenes, interactuar con ellos, así como visualizar y navegar por los modelos difusos de color en tres dimensiones. En este sentido, estas herramientas favorecen la comprensión de los conceptos de núcleo y soporte de un conjunto difuso tridimensional (color difuso) gracias a la visualización y navegación en 3D a través de los colores difusos de una manera sencilla e intuitiva. Además, se puede interactuar y visualizar las distintas superficies que delimitan tanto el núcleo, como el soporte, como el alfa-corte 0.5 con un simple click de ratón, facilitando así el aprendizaje de éstos conceptos. Asimismo, con *JFCS* se pueden obtener y visualizar los grados de pertenencia de colores precisos a colores difusos de forma amigable mediante la selección de colores precisos en imágenes o seleccionados de una paleta de colores. También, con *JFCS* se pueden calcular imágenes mapeadas que ilustran visualmente las correspondencias de color entre píxeles y categorías de color modeladas por colores difusos. Gracias a esta utilidad, se puede entender de forma visual el concepto de frontera borrosa entre colores así como degradados entre colores en imágenes, tal y como se mostró en la Figura 1.

Finalmente, *JFCS* tiene asociada una web en <http://www.uco.es/semanticolor/JFCS> con descargas, material complementario y videos demostrativos de la utilidad que *JFCS* proporciona, tanto para el modelado de color como para un enfoque didáctico de comprensión de la lógica difusa como herramienta para modelar la imprecisión del color.

Referencias

- [1] J. Chamorro-Martínez, J. Soto-Hidalgo, P. Martínez-Jiménez, and D. Sánchez, “Fuzzy color spaces: A conceptual approach to color vision,” *IEEE Transactions on Fuzzy Systems*, (submitted).
- [2] J. Soto-Hidalgo, P. Martínez-Jiménez, J. Chamorro-Martínez, and D. Sánchez, “JFCS: A color modeling java software based on fuzzy color spaces,” *IEEE Computational Intelligence Magazine*, (In press) 2016.

¿Qué debe saberse para emplear la lógica fuzzy?

Enric Trillas¹

¹ etrillasetrillas@gmail.com

1. Introducción

Cuanto sigue proviene de la experiencia adquirida tras muchos años de investigar en el campo de la lógica borrosa, de introducirla a estudiantes universitarios de ingeniería y de observar cómo suele abordarse la representación práctica de los sistemas borrosos. Por resumir la respuesta al título: Los sistemas borrosos no son únicamente una materia de grado, como suele decir Lotfi A. Zadeh, sino y de forma muy especial, un asunto de diseño que exige un gran cuidado.

La razón de lo anterior es el hecho de que la función de pertenencia a un conjunto borroso de etiqueta lingüística P y con uso impreciso en un universo del discurso X , no es sino una aproximación a una medida del significado de P en X . Un significado que, en todo caso, depende del contexto en el cual se inscriba el problema y, a la vez, está guiado por el propósito con el cual P se use en él. El diseño de un sistema borroso no sólo comprende especificar las funciones de pertenencia de los términos lingüísticos implicados, sino también las funciones que puedan representar los conectivos involucrados; en el campo de lo impreciso casi no hay especificaciones universales, como sucede con los términos precisos especificados mediante conjuntos clásicos.

La lógica borrosa ayuda, básicamente, a representar en términos fuzzy bien a sistemas dinámicos descritos lingüísticamente, o bien a enunciados en lenguaje natural reflejando razonamientos no necesariamente deductivos; casi siempre es el razonamiento de sentido común, expresado en un lenguaje natural, lo que debe tenerse en cuenta. No se dispone, en el campo impreciso, de un ‘axioma de especificación’ como el del campo preciso y que autorice representar los predicados imprecisos por un único conjunto borroso; es bien conocido que un mismo predicado admite muchas especificaciones que, si mantienen algunas características comunes son, realmente, funciones de pertenencia distintas.

2. Función de pertenencia y colectivos

2.1. Los únicos predicados P representables matemáticamente son los medibles; aquellos de los que se conoce tanto su significado (comparativo) primario, como una medida del mismo. Sin embargo y usualmente, el diseñador no llega a conocer por completo el primero sino y a lo más, una parte del mismo; con ello, no puede tener la seguridad de que la función de pertenencia que defina lo sea realmente y debe suponer que aproxima suficientemente a una medida. Además, tampoco conoce todos los pares de elementos relacionados en la forma ‘menos P que’ y debe emplear, en su lugar, la relación que la función de pertenencia permite establecer; en conclusión, al linearizarlo el diseñador modifica el significado del predicado. Ello le exige un gran cuidado en la obtención de la función de pertenencia que, fácilmente, puede especificar un significado distinto del

predicado, es decir, el de otro predicado; de ahí el gran cuidado que debe poner al darla por buena.

2.2. Por lo que respecta a los conectivos, se suelen representar por funciones bien conocidas matemáticamente (por ejemplo, t-normas, negaciones fuertes, funciones de agregación, etc.) que, por su definición, implican propiedades que deben suponerse de los conectivos lingüísticos originales (como son, por ejemplo, la asociatividad, la conmutatividad, estar entre el mínimo y el máximo, etc.) y que, con frecuencia son por lo menos dudosas en el lenguaje. Por ejemplo, la propiedad conmutativa de la conjunción deja de valer en cuanto el tiempo interviene y la asociatividad no siempre cabe suponerla. De nuevo, el diseñador debe elegir las funciones que representen los conectivos de acuerdo con propiedades realmente existentes en el lenguaje o que, por lo menos, pueda suponerlas razonablemente; a la vez, debe tener en cuenta que, una vez elegidas, esas funciones fuerzan propiedades que pueden alterar la realidad.

3. Conclusión

No deben confundirse ni un conjunto borroso con una de sus funciones de pertenencia, ni la lógica fuzzy con el estudio matemático de las funciones que ‘puedan eventualmente’ representar los conectivos ‘and’, ‘or’, ‘if/then’, etc. Los conjuntos borrosos y la lógica fuzzy se refieren, básica y respectivamente, a la especificación de predicados imprecios y a la representación simbólica del razonamiento de sentido común, del cual el deductivo formal no es sino un caso particular y, además y de ordinario, poco frecuente. Tanto los conjuntos borrosos como la llamada lógica fuzzy no son, sin más, una continuación de la teoría de conjuntos y la lógica; su principal diferencia con ellos es la falta de universalidad de las correspondientes representaciones y, por consiguiente, el hecho de que estas deban diseñarse. No debe confundirse, sin más, a la lógica borrosa con una lógica. El ‘practitioner’ de las metodologías borrosas debe saberlo y, para ello, ha de estudiar previamente algún texto que vaya más allá de una descripción matemática de operadores y atienda al fondo de las cuestiones relativas a la imprecisión sin, por otro lado, confundirla con la incertidumbre.

References

- [1] E. Trillas, L. Eciolaza: *Fuzzy Logic: An Introductory Course for Engineering Students*. Springer. 2015.
- [2] E. Trillas: *Razonamiento; significado, incertidumbre y borrosidad*. Ed. UPNA. 2015.
- [3] H.T. Nguyen, E.A. Walker: *A First Course in Fuzzy Logic*. CRC Press. 2005.
- [4] R. López de Mántaras: *Approximate Reasoning Models*. Prentice Hall PTR. 1990.
- [5] H.J. Zimmermann: *Fuzzy Set Theory and Its Applications*. Springer. 1992.
- [6] L.A. Zadeh: *Computing with Words. Principal Concepts and Ideas*. Springer. 2012.

Hipercontextos L-fuzzy

Cristina Alcalde¹, Ana Burusco²

¹ Dept. Matemática Aplicada. Escuela de Ingeniería de Gipuzkoa. UPV/EHU.
Plaza de Europa, 1. 20018 San Sebastian. España
c.alcalde@ehu.es

² Dept. Automática y Computación. Instituto de Smart Cities.
Universidad Pública de Navarra. Campus de Arrosadía. 31006 Pamplona. España
burusco@unavarra.es

La Teoría de Conceptos L -fuzzy [5, 6] analiza la información existente en un contexto L -fuzzy mediante el uso de los conceptos L -fuzzy. Estos contextos L -fuzzy son tuplas (L, X, Y, R) , con L retículo completo, X e Y conjuntos de objetos y atributos respectivamente y $R \in L^{X \times Y}$ una relación L -fuzzy entre los objetos y atributos. Podemos decir que esta teoría nace como una extensión al caso fuzzy de la Teoría Formal de Conceptos de R. Wille ([8]) cuando queremos estudiar relaciones entre los objetos y los atributos que toman valores en un retículo L .

Existen otras generalizaciones importantes de la Teoría Formal de Conceptos utilizando operadores de implicación residuados (R. Belohlavek [3, 4]) y otras extensiones como los retículos de conceptos property-oriented y multi-adjoint ([7]).

En algunas ocasiones necesitamos extender estos contextos L -fuzzy estableciendo marcos de trabajo que permitan representar cada $R(x, y)$, con $x \in X$ e $y \in Y$, como una colección de valores que tenga estructura de contexto con un conjunto de objetos Q_x asociados al objeto x y un conjunto de atributos S_y asociados al atributo y .

Para hacerlo, definiremos tuplas $(L, X, Y, (Q_x)_{x \in X}, (S_y)_{y \in Y}, R)$, con L retículo completo, X e Y conjuntos de objetos y atributos respectivamente, $(Q_x)_{x \in X}$ familia de conjuntos asociada a los objetos de X a la que llamaremos de objetos asociados, $(S_y)_{y \in Y}$ que se llamará familia de atributos asociados y R tal que $R(x, y)$ es a su vez otra relación $R_{xy} \in L^{Q_x \times S_y}$. En algunos casos nos interesará hacer un estudio únicamente en función de los elementos de X o de Y , pero en otros casos también estaremos interesados en analizar los resultados en función de las familias $(Q_x)_{x \in X}$ y $(S_y)_{y \in Y}$.

Estas tuplas, que guardan relación con los contextos L -fuzzy, serán llamadas hipercontextos L -fuzzy. Cada una de estas relaciones R_{xy} representará un contexto L -fuzzy (L, Q_x, S_y, R_{xy}) de forma que la R inicial será una relación de relaciones. En este trabajo estudiaremos los hipercontextos L -fuzzy cuando $L = [0, 1]$.

Nuestros hipercontextos serán extensiones al caso fuzzy de ciertos multicontextos nítidos de Wille [9] que guarden unas determinadas condiciones respecto de sus conjuntos de objetos y de atributos. Por otra parte, el enfoque dado en este trabajo utilizando operadores OWA [10] será también distinto al allí tratado.

Este trabajo ha sido parcialmente subvencionado por el Grupo de investigación Sistemas Inteligentes y Energía (SI+E) de la UPV/EHU (Proyecto IT677-13) y por el Grupo de investigación Inteligencia Artificial y Razonamiento Aproximado de la UPNa (Proyecto TIN2013-40765-P).

El primer problema con el que nos encontramos consiste en que estas nuevas estructuras no tienen forma de contexto L -fuzzy. Por eso, la manera de abordar el estudio de estos hipercontextos L -fuzzy consistirá en transformarlos en contextos L -fuzzy ordinarios realizando distintos tipos de agregaciones mediante operadores OWA. La elección del operador dependerá de cuáles sean nuestros intereses de estudio. Se pueden además realizar estudios más completos cuando una de las dos familias $((Q_x)_{x \in X})$ o $((S_y)_{y \in Y})$ sea fija.

A continuación podemos obtener información mediante el estudio de los conceptos L -fuzzy de los contextos resultantes. Para ello partiremos de un conjunto de partida que represente lo que queramos analizar siguiendo las pautas definidas en [1, 2] y calcularemos el concepto L -fuzzy asociado. Del estudio de este concepto obtendremos información relevante para nuestros intereses.

Referencias

- [1] C. Alcalde, A. Burusco, R. Fuentes-González: *L-Fuzzy Concepts and linguistic variables in knowledge acquisition processes*, in: Proc. of CLA'10, Sevilla (2010) 38–49.
- [2] C. Alcalde, A. Burusco, R. Fuentes-González, I. Zubia: The use of linguistic variables and fuzzy propositions in the L -Fuzzy Concept Theory, *Computers and Mathematics with Applications* **62** (2011) 3111–3122.
- [3] R. Bělohlávek, *Fuzzy relational Systems*, IFSR International Series on Systems Science and Engineering **20** (2002).
- [4] R. Bělohlávek, *Fuzzy Galois connections and fuzzy concept lattices: from binary relations to conceptual structures*, in: V. Novak, I. Perfilova (Eds.), *Discovering the World with Fuzzy Logic*, Physica-Verlag (2000) 462–494.
- [5] A. Burusco, R. Fuentes-González: *The Study of the L-Fuzzy Concept Lattice*, *Mathware and Soft Computing* **1** (3) (1994) 209–218.
- [6] A. Burusco, R. Fuentes-González: *Construction of the L-Fuzzy Concept Lattice*, *Fuzzy Sets and Systems* **97** (1)(1998) 109–114.
- [7] J. Medina, M. Ojeda-Aciego, J Ruiz-Calviño: *Formal concept analysis via multi- adjoint concept lattices*, *Fuzzy Sets and Systems* **160** (2) (2009) 130–144.
- [8] R. Wille: *Restructuring lattice theory: an approach based on hierarchies of concepts*, in: Rival I. (Ed.), *Ordered Sets*, Reidel, Dordrecht-Boston, (1982) 445–470.
- [9] R. Wille: *Conceptual structures of multicontexts* In: Eklund, P.W., Ellis, G., Mann, G. (eds.) *Proceedings of the 4th ICCS 1996*, Sidney, *Lecture Notes in Computer Science* **1115** (1996) 23–39.
- [10] R.R. Yager: *On ordered weighted averaging aggregation operators in multi-criteria decision making*, *IEEE Transactions on Systems, Man and Cybernetics* **18** (1988) 183–190.

Fuzzy adjunctions revisited

I.P. Cabrera¹, P. Cordero¹, B. De Baets², F. García-Pardo¹, M. Ojeda-Aciego¹

¹ *Universidad de Málaga, Depto. Matemática Aplicada, Spain
Blv. Louis Pasteur 35, 29071 Málaga, Spain*

{ipcabrera,cordero,fgarciap,aciego}@uma.es

² *KERMIT. Department of Mathematical Modelling, Statistics and Bioinformatics. Ghent University, Coupure links 653, 9000 Gent, Belgium
Bernard.DeBaets@UGent.be*

Extended abstract

In this work we are aiming at obtaining the weakest notion of adjunction between fuzzy structures. This topic continues our research line on the study and construction of adjunctions [2, 5, 6].

We will focus on the notion of fuzzy relation which, in some sense, can be interpreted as a fuzzy function. A number of papers dealing with this problem can be found in the literature [3, 4, 7] and among all these approaches, we will work with the notion of *functional* fuzzy relation, which is equivalent to what Ćirić calls partial fuzzy function (but actually does not coincide with the original notion of partial fuzzy function given by Demirci).

Definition 1 A fuzzy relation $\mu \in L^{A \times B}$ is said to be functional if for each $a_1, a_2 \in A$ and $b_1, b_2 \in B$ the following inequality holds

$$\mu(a_1, b_1) \otimes \mu(a_2, b_1) \otimes \mu(a_1, b_2) \leq \mu(a_2, b_2) \quad (1)$$

A functional fuzzy relation μ induces two fuzzy equivalence relations, $\approx_\mu$ on B and $\approx_{\mu^{-1}}$ on A , defined, respectively by

$$(b_1 \approx_\mu b_2) = (b_1 = b_2) \vee \bigvee_{a \in A} (\mu(a, b_1) \otimes \mu(a, b_2))$$

$$(a_1 \approx_{\mu^{-1}} a_2) = (a_1 = a_2) \vee \bigvee_{b \in B} (\mu(a_1, b) \otimes \mu(a_2, b))$$

for all $b_1, b_2 \in B$ and $a_1, a_2 \in A$. The two fuzzy equivalences above allow to consider A and B as fuzzy structures $\langle A, \approx_A \rangle$ and $\langle B, \approx_B \rangle$ where the fuzzy equivalences $\approx_A$ and $\approx_B$ are the induced by μ .

In order to consider an adequate notion of fuzzy adjunction between fuzzy structures, it is necessary to fix fuzzy preorders on both of them. A fuzzy relation $\rho_A: A \times A \rightarrow L$ is said to be a *fuzzy preorder compatible with $\langle A, \approx_A \rangle$* if and only if the two following conditions hold

$$(a_1 \approx_A a_2) \leq \rho_A(a_1, a_2) \quad \text{and} \quad \rho_A(a_1, a_2) \otimes \rho_A(a_2, a_3) \leq \rho_A(a_1, a_3)$$

In this abstract setting, it makes sense to study suitable notions of isotonicity for a functional fuzzy relation which generalize the well-known definitions of isotonicity for compatible crisp mappings and for perfect fuzzy functions. A possible definition could be

$$\mu(a_1, b_1) \otimes \mu(a_2, b_2) \otimes \rho_A(a_1, a_2) \leq \rho_B(b_1, b_2) \quad (2)$$

for all $b_1, b_2 \in B$ and $a_1, a_2 \in A$.

Similarly, another related notion which has to be extended is the inflationary property, whose definition for an internal functional fuzzy relation $\varphi: A \times A \rightarrow L$ in this setting could be $\varphi(a_1, a_2) \leq \rho_A(a_1, a_2)$.

What would be a desirable notion of adjunction in this extended framework? An adjunction between fuzzy preorders $\langle A, \rho_A \rangle$ and $\langle B, \rho_B \rangle$ can be defined as a pair (μ, ν) of functional fuzzy relations. In principle, we can find the three possibilities below:

- $\mu(a_1, b_1) \otimes \nu(b_2, a_2) \leq \rho_B(b_1, b_2) \leftrightarrow \rho_A(a_1, a_2)$.
- $\nu(b_2, a_2) \otimes \rho_A(a_2, a_2) = \mu(a_1, b_1) \otimes \rho_B(b_1, b_2)$.
- μ and ν are isotone, $\nu \circ \mu$ inflationary and $\mu \circ \nu$ deflationary.

It has to be studied which of the definitions above behaves as expected in relation to closure operators and closure systems.

References

- [1] R. Belohlavek. *Fuzzy Relational Systems: Foundations and Principles*. Kluwer Academic Publishers, Norwell, MA, USA. 2002.
- [2] I.P. Cabrera, P. Cordero, F. Garca, M. Ojeda-Aciego, B. De Baets. *On the construction of adjunctions between a fuzzy preposet and an unstructured set*. Submitted for publication, 2016.
- [3] M. Ćirić, J. Ignjatović, S. Bogdanović: Uniform fuzzy relations and fuzzy functions. *Fuzzy Sets and Systems* 160(8):1054–1081, 2009.
- [4] M. Demirci. *Fuzzy functions and their applications*. J. Mathematical Analysis Applications, 252:495–517, 2000.
- [5] F. Garca-Pardo, I.P. Cabrera, P. Cordero, and M. Ojeda-Aciego. *On the construction of fuzzy Galois connections*. In XVII Spanish Conference on Fuzzy Logic and Technology, ESTYLF, pages 99-102, 2014.
- [6] F. Garca-Pardo, I.P. Cabrera, P. Cordero, M. Ojeda-Aciego, and F.J. Rodríguez. *On the definition of suitable orderings to generate adjunctions over an unstructured codomain*. Information Sciences, 286:173-187, 2014.
- [7] S. Gottwald. *Fuzzy Sets and Fuzzy Logic. Foundations of Application—from a Mathematical Point of View*. Vieweg, Braunschweig, Wiesbaden 1993.

Towards attribute implications in multi-adjoint context

M. E. Cornejo¹, J. Medina², E. Ramírez-Poussa², R. Rodríguez-Galván²

¹ *Department of Statistic and O.R., University of Cádiz. Spain*

mariaeugenia.cornejo@uca.es

² *Department of Mathematics, University of Cádiz. Spain*

{jesus.medina,eloisa.ramirez,rafael.rodriguez}@uca.es

Extracting information from databases has become a key task in different areas such as stock market prediction, disease diagnosis, census data analysis, etc. Often, the obtained information is represented as a set of rules which summarizes the knowledge contained in databases. These sets of rules are known as logic programs. However, traditional logic program languages are not able to handle uncertainty and approximated reasoning. In order to solve this trouble, different fuzzy logic programming systems have been proposed such as Probabilistic Deductive Databases, Hybrid Probabilistic Logic Programs, Annotated Logic Programming, Residuated Logic Programming, etc.

Multi-adjoint Logic Programming [8] is a generalization of previous logic programming frameworks whose semantic structure is the multi-adjoint lattice. From a multi-adjoint lattice, a set of weighted rules and facts of a given language is defined, which is called multi-adjoint logic program. This work is focused on the use of a formal tool which allow us to obtain multi-adjoint logic programs from databases. Specifically, we are interested in employing Formal Concept Analysis.

Formal Concept Analysis is a mathematical technique to extract information from relational databases that contain a set of objects and a set of attributes. The pieces of the extracted information are called concepts and they are hierarchized to obtain concept lattices. A deep survey about attribute dependencies in data, which are known as attribute implications, is presented in [1]. Nevertheless, powerful advantages provided by multi-adjoint environment are not considered.

Our contribution complements the study presented in [6], related to attribute implications in multi-adjoint concept lattices [7], where the algebraic structure considered is a multi-adjoint algebra [2] instead of the residuated lattice. For example, in this framework, operators do not need to be neither monotone nor associative.

From now on, we will fix a multi-adjoint frame $(L_1, L_2, P, \&_1, \swarrow^1, \lhd_1, \dots, \&_n, \swarrow^n, \lhd_n)$ where $(L_1, \preceq_1)$ and $(L_2, \preceq_2)$ are complete lattices, $(P, \leq)$ is a poset and $(\&_i, \swarrow^i, \lhd_i)$ is an adjoint triple with respect to L_1, L_2, P , for all $i \in \{1, \dots, n\}$. In addition, we will consider a multi-adjoint context (A, B, R, σ) such that A and B are non-empty sets of attributes and objects, respectively, R is a P -fuzzy relation $R: A \times B \rightarrow P$ and $\sigma: A \times B \rightarrow \{1, \dots, n\}$ is a mapping which associates any element in $A \times B$ with a particular adjoint triple in the frame [3].

First of all, we introduce the notion of fuzzy implication in the multi-adjoint framework, which is a direct generalization of the classical one, and the degree in which this fuzzy implication is valid in a context. In this definition we will consider truth-stressing hedge, these operators were studied in [5].

Definition 1. Given three fuzzy subsets of attributes $f_1, f_2, f_3 \in L_1^A$, the degree in which a fuzzy implication $f_2 \Leftarrow f_1$ is valid in f_3 is

$$\|f_2 \Leftarrow f_1\|_{f_3} = S(f_2, f_3) \frown S(f_1, f_3)^*$$

where $S(f_i, f_j) = \inf\{f_i(a) \frown f_j(a) \mid a \in A\}$ for all $i, j \in \{1, 2, 3\}$ and $*$ is a truth-stressing hedge.

This notion is extrapolated to a set of fuzzy subsets of attributes $\mathcal{F} \subset L_1^A$, defining the degree in which the fuzzy implication $f_2 \Leftarrow f_1$ is valid in $\mathcal{F}$ as

$$\|f_2 \Leftarrow f_1\|_{\mathcal{F}} = \bigwedge_{f \in \mathcal{F}} S(f_2, f) \frown S(f_1, f)^*$$

In order to know the degree in which a fuzzy implication is valid in a multi-adjoint concept lattice, it is not necessary to calculate the degree of validity in all the concepts of the concept lattice, that is, on the whole set of intensions of the concepts of the lattice. We only need to compute the degree in which a fuzzy implication is valid on the set of join-irreducible elements of the lattice [4].

Theorem 1. Let $f_1, f_2 \in L_1^A$ be fuzzy subsets of attributes. Then,

$$\|f_2 \Leftarrow f_1\|_{J_F(A, B, R, \sigma)} = \|f_2 \Leftarrow f_1\|_{Int(A, B^*, R, \sigma)}$$

where $J_F(A, B, R, \sigma)$ denotes the set of join-irreducible elements of the lattice associated with (A, B, R, σ) and $Int(A, B^*, R, \sigma)$ denotes the set of intensions of the concepts of the lattice associated with (A, B^*, R, σ) .

References

- [1] R. Belohlavek and V. Vychodil. Attribute dependencies for data with grades. *CoRR*, 2014.
- [2] M. E. Cornejo, J. Medina, and E. Ramírez-Poussa. Multi-adjoint algebras versus non-commutative residuated structures. *International Journal of Approximate Reasoning*, 66:119–138, 2015.
- [3] M. E. Cornejo, J. Medina, and E. Ramírez-Poussa. A comparative study of adjoint triples. *Fuzzy Sets and Systems*, 211:1–14, 2013.
- [4] M. E. Cornejo, J. Medina, and E. Ramírez-Poussa. Attribute reduction in multi-adjoint concept lattices. *Information Sciences*, 294:41–56, 2015.
- [5] P. Hájek. On very true. *Fuzzy Sets Systems* 124, 329–333, 2001.
- [6] V. Liñeiro-Barea, J. Medina and I. Medina-Bulo. Towards generating fuzzy rules via Fuzzy Formal Concept Analysis. *7th European Symposium on Computational Intelligence and Mathematics (ESCIM 2015)*, I:60–65, 2015.
- [7] J. Medina, M. Ojeda-Aciego, and J. Ruiz-Calviño. Formal concept analysis via multi-adjoint concept lattices. *Fuzzy Sets and Systems*, 160(2):130–144, 2009.
- [8] J. Medina, M. Ojeda-Aciego, and P. Vojtáš. Multi-adjoint logic programming with continuous semantics. *Proc of Logic Programming and Non-Monotonic Reasoning, LPNMR'01*, Lecture Notes in Computer Science, 2173:351–364, Springer, 2001.

On the size of reducts in multi-adjoint contexts

M. Eugenia Cornejo¹, Jesús Medina², Eloísa Ramírez-Poussa²

¹ *Department of Statistic and O.R., University of Cádiz. Spain*
mariaeugenia.cornejo@uca.es

² *Department of Mathematics, University of Cádiz. Spain*
{jesus.medina,eloisa.ramirez}@uca.es

Formal Concept Analysis has become one of the most forceful mathematical tools to extract information from databases that contain two non-empty sets A and B , usually interpreted as attributes and objects, together with a relation between them. The units of extracted information are called concepts and these concepts can be ordered, obtaining the algebraic structure of a concept lattice. Due to the fact that the complexity to build concept lattices is exponential in the number of attributes and objects, it is important to find mechanisms to reduce the number of attributes [1, 2].

In the fuzzy framework of multi-adjoint concept lattices [3, 4], a suitable attribute reduction method and several results were introduced in order to classify the set of attributes in [5]. This classification provides a procedure to know whether an attribute should be considered or not. Consequently, it can be used to extract reducts, that is, minimal sets of attributes preserving the original information of the dataset.

In addition, when the attribute classification verifies that the set of relatively necessary attributes is not empty several reducts can be obtained. From this fact two interesting questions arise: How can the reducts be computed? Should all the reducts have the same cardinality? Our contribution is an initial study which aims to answer these questions. The most complex part in the construction process of reducts is the selection of relatively necessary attributes. In order to simplify this selection we need to define the following set:

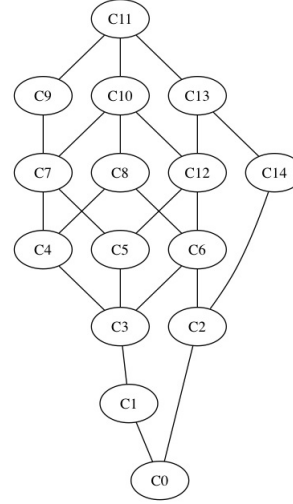
Definition 1. *Given a multi-adjoint frame $(L_1, L_2, P, \&_1, \dots, \&_n)$ and a context (A, B, R, σ) associated with the concept lattice $(\mathcal{M}, \leq)$. Let C be a concept of $(\mathcal{M}, \leq)$, we define the set of attributes generating C , as the set:*

$$Atg(C) = \{a_i \in A \mid \text{there exists } \phi_{a_i, x} \in \Phi \text{ such that } \langle \phi_{a_i, x}^\downarrow, \phi_{a_i, x}^\uparrow \rangle = C\}$$

In order to show an illustrative example, we will consider a multi-adjoint framework $(L_1 = [0, 1]_{10}, L_2 = [0, 1]_4, L_3 = [0, 1]_5, \preceq, \&_G^*)$ where $\&_G^*$ the discretization Gödel conjunc-tor defined on $L_1 \times L_2$, and a context (A, B, R, σ) where $A = \{a_1, a_2, a_3, a_4, a_5, a_6, a_7\}$, $B = \{b_1, b_2, b_3\}$, $R: A \times B \rightarrow L_3$ is given below and σ is constantly $\&_G^*$.

From this framework and this context, according to the attribute classification theorems presented in [5], we obtain $C_f = \{a_1, a_2\}$, $K_f = \{a_4, a_5, a_6, a_7\}$ and $I_f = \{a_3\}$. Consequently, the attributes a_1 and a_2 must belong to all reducts and a_3 should be removed. To select the relatively necessary attributes, we analyze the attributes generating each meet-irreducible concept by means of the above definition. In this case, we find two different meet-irreducible

R	a_1	a_2	a_3	a_4	a_5	a_6	a_7
b_1	0.6	0.2	0.2	1	0.6	0.2	0
b_2	0.8	0.4	0.4	1	0.8	0.6	0.6
b_3	0.6	0.6	0.2	0	0	0.2	0



elements C_1 and C_9 such that $Atg(C_1) = \{a_6, a_7\}$ and $Atg(C_9) = \{a_4, a_5, a_6\}$. The rest of meet-irreducible elements are generated from the attributes a_1 and a_2 .

Therefore, three reducts $Y_1 = \{a_1, a_2, a_4, a_7\}$, $Y_2 = \{a_1, a_2, a_5, a_7\}$ and $Y_3 = \{a_1, a_2, a_6\}$ can be built. We can observe that the different reducts do not have the same cardinality, which reveals that we can find reducts with different number of attributes. Therefore, the degree of difficulty in the construction process of the concept lattice can be reduced if we consider the proper starting reduct. Finally, we show a result that states under what conditions all reducts of a multi-adjoint context have the same cardinality.

Proposition 1. *Given a multi-adjoint frame $(L_1, L_2, P, \&_1, \dots, \&_n)$ and a context (A, B, R, σ) with the associated concept lattice $(\mathcal{M}, \leq)$. If the set $\{Atg(C) \mid C \in M_F(A) \text{ and } Atg(C) \cap K_f \neq \emptyset\}$ is a partition of K_f , then all reducts $Y \subseteq A$ have the same cardinality and, that is:*

$$card(Y) = card(C_f) + card(\{Atg(C) \mid C \in M_F(A) \text{ and } Atg(C) \cap K_f \neq \emptyset\})$$

More results can be given in order to fix a proper reduct. These properties and new efficient reduction procedures will be introduced in an extended version.

References

- [1] J. Medina. Relating attribute reduction in formal, object-oriented and property-oriented concept lattices. *Computers & Mathematics with Applications*, 64:1992–2002, 2012.
- [2] L. Wei and J. J. Qi. Relation between concept lattice reduction and rough set reduction. *Knowledge-Based Systems*, 23:934–938, 2010.
- [3] J. Medina and M. Ojeda-Aciego. Multi-adjoint t-concept lattices. *Information Sciences*, 180:712–725, 2010.
- [4] J. Medina, M. Ojeda-Aciego and J. Ruiz-Calviño. Formal concept analysis via multi-adjoint concept lattices. *Fuzzy Sets and Systems*, 160:130–144, 2009.
- [5] M. E. Cornejo, J. Medina, and E. Ramírez-Poussa. Attribute reduction in multi-adjoint concept lattices. *Information Sciences*, 294:41–56, 2015.

Similarity-based reasoning using prototypes and counterexamples

Soma Dutta¹, Francesc Esteva², Lluís Godo³

¹ *Vistula University, Warsaw & University of Warsaw, Poland*
somadutta9@gmail.com

² *IIIA - CSIC, Bellaterra, Spain*
{esteva,godo}@iiia.csic.es

Vague properties, in the sense of gradualness, are characterized by the existence of borderline cases; that is, objects or situations for which the property only partially applies. The aim of this paper is to investigate how a logic for vague concepts can be defined assuming that a vague concept α is given by a set of prototypical situations $[\alpha]^+ \subseteq \Omega$ where α definitely applies, as well as a set of counterexamples $[\alpha]^- \subseteq \Omega$ where α does not apply for sure. In this paper we will assume that this information is complete, that is, $[\alpha]^+$ is the whole set of prototypes and $[\alpha]^-$ is the whole set of counter-examples. This also means that the remaining set of situations $\Omega \setminus ([\alpha]^+ \cup [\alpha]^-)$ are those where α only partially applies. Of course, to be in a consistent scenario, we will require that $[\alpha]^+ \cap [\alpha]^- = \emptyset$. In such a case, one might think of a three-valued framework, where for each situation $w \in \Omega$ the degree to which α applies at w is defined as follows:

$$app(w, \alpha) = \begin{cases} 1, & \text{if } w \in [\alpha]^+ \\ 0, & \text{if } w \in [\alpha]^- \\ 1/2, & \text{otherwise} \end{cases}$$

This 3-valued vagueness model, where third value $1/2$ does not represent unknown but borderline (see [1] for a discussion on this topic), is indeed very rough. A more refined model can be introduced by assuming the availability of a (fuzzy) similarity relation $S : \Omega \times \Omega \rightarrow [0, 1]$ among situations. In such a case, for $w \in \Omega \setminus ([\alpha]^+ \cup [\alpha]^-)$ one can measure how close w is to some prototype of α , and on the one hand how close w is to some of its counterexamples.

$$\begin{aligned} S(w, [\alpha]^+) &= \sup\{S(w, w') : w' \in [\alpha]^+\} \\ S(w, [\alpha]^-) &= \sup\{S(w, w') : w' \in [\alpha]^-\} \end{aligned}$$

Finally, to aggregate how much α applies to situation w , considering both the values, one can implement a commonsense rule like this one:

“The closer w is to some prototype and the farther is to any of the counterexamples,
the **more α applies** to w ”

Of course, one can think of different models of formalization following this rule; in principle one can think of a suitable aggregation operator $\otimes$ and define:

$$app^*(w, \alpha) = S(w, [\alpha]^+) \otimes (1 - S(w, [\alpha]^-))$$

This may in principle be appropriate as soon as $\otimes$ properly extends the above three-valued model in the sense that if $S(w, [\alpha]^+) = 1$ then $app^*(w, \alpha) = 1$ and if $S(w, [\alpha]^-) = 1$ then $app^*(w, \alpha) = 0$, and otherwise $0 < app^*(w, \alpha) < 1$. Assuming the similarity is strict, i.e. such that $S(w, w') = 1$ iff $w = w'$, a relevant example of such an aggregation operator, given in [4], is:

$$x \otimes y = \frac{y}{1 - x + y},$$

but other operators may be suitable as well. Note that the mapping $\mu_\alpha : \Omega \rightarrow [0, 1]$, defined as $\mu_\alpha(w) = app^*(w, \alpha)$, specifies a fuzzy set which can smoothly incorporate the finer distinctions of the apparent 3-valued nature of α . Or equivalently, one can also interpret $app^*(w, \alpha)$ as the degree to which α is satisfied by an interpretation, model or situation $w \in \Omega$.

Following the latter logical interpretation, the aim of this paper is to extend the approach used in [2, 3] (where only the values of $S(w, [\alpha]^+)$'s were considered) to define a logical framework to reason with fuzzy concepts given by a set of prototypes and counterexamples in the line of [4], but with some differences. The main difference is that we consider here as a working assumption that the base logic for our model of vague concepts based on prototypes and counter-examples is 3-valued Łukasiewicz logic $\mathbb{L}_3$. Then, to refine such a logic, we will extend the language of $\mathbb{L}_3$ with a modality $\Diamond$, and we will consider a Kripke-style semantics given by models $M = (W, e, S)$, where W is a set of worlds, $S : W \times W \rightarrow [0, 1]$ is a similarity relation on worlds, and $e(w, \cdot) : V \rightarrow \{0, 1/2, 1\}$ is a $\mathbb{L}_3$ -valuation of variables. The evaluation will be extended to (non-nested) modal formulas by stipulating $e(w, \Diamond\alpha) = app^*(w, \alpha)$, where $[\alpha]^+ = \{w \in W \mid w(\alpha) = 1\}$ and $[\alpha]^- = \{w \in W \mid w(\alpha) = 0\}$. In this framework we plan to study different notions of graded entailment, as well as, to explore a Hilbert-style axiomatization and a proof system for the logical system.

References

- [1] D. Ciucci, D. Dubois, J. Lawry. Borderline vs. unknown: comparing three-valued representations of imperfect information. *Int. J. Approx. Reasoning* **55**(9): 1866-1889 (2014)
- [2] F. Esteva, P. Garcia, L. Godo, R. O. Rodríguez, A modal account of similarity-based reasoning, *Int. J. Approx. Reasoning* **16** (1997), 235 - 260.
- [3] F. Esteva, L. Godo, R. O. Rodríguez, T. Vetterlein., Logics for approximate and strong entailments, *Fuzzy Sets Syst.* **197** (2012), 59 - 70.
- [4] T. Vetterlein. Logic of prototypes and counterexamples: possibilities and limits. In Proc. of IFSA-EUSFLAT-15, J.M. Alonso et al. (eds.), Atlantis Press, pp. 697-704, 2015.

Extensión a la Lógica Borrosa de algunas nociones asociadas a Relaciones y a Grafos ordinarios utilizando Retículos de Conceptos Borrosos.

Ramón Fuentes-González¹
 Universidad Pública de Navarra
 rfuentes@unavarra.es

Si $(L, \leq)$ es un retículo completo con un operador tipo implicación $\rightsquigarrow: L \times L \rightarrow L$ y si X e Y son referenciales ordinarios no vacíos, consideremos los retículos $(L^X, \leq)$, $(L^Y, \leq)$, $(L^{X \times Y}, \leq)$ de subconjuntos y relaciones L -borrosos con sus órdenes habituales obtenidos como extensiones del orden en el retículo L . El operador $\triangleleft: L^X \times L^{X \times Y} \rightarrow L^Y$ [1] está asociado a la implicación $\rightsquigarrow$ y es tal que para $M \in L^X$ y $S \in L^{X \times Y}$, el L -borroso $(M \triangleleft S) \in L^Y$ es:

$$(M \triangleleft S)(y) = \inf \{ M(x) \rightsquigarrow S(x, y) \mid x \in X \} \quad \forall y \in Y.$$

Si se interpretan los referenciales X e Y respectivamente como conjuntos de objetos y de atributos y si $R: X \times Y \rightarrow L$ es una L -relación de incidencia de X en Y , entonces un L -concepto del L -contexto (L, X, Y, R) es un par $(A, B) \in L^X \times L^Y$ de subconjuntos L -borrosos tales que $B = (A \triangleleft R)$ y $A = (B \triangleleft R^{op})$, siendo $R^{op}: Y \times X \rightarrow L$ la opuesta de R . Como es conocido [2], el conjunto $\mathcal{L}(L, X, Y, R)$ de esos L -conceptos (A, B) , $(C, D), \dots$ es un retículo completo con la relación de orden $\preceq$ tal que: $(A, B) \preceq (C, D) \Leftrightarrow A \leq C$ (equivalente a $B \geq D$). El par $(\mathcal{L}(L, X, Y, R), \preceq)$ es el *Retículo de L -Conceptos Borrosos* del L -contexto.

En este trabajo analizaremos casos particulares interesantes: los que aparecen al considerar L -conceptos (A, B) en L -contextos del tipo (L, X, X, R) . Es decir, en los que el conjunto de objetos coincide con el conjunto de atributos ($X = Y$) [3]. En esta situación, la relación de incidencia R pertenece a $L^{X \times X}$ por lo que puede ser reflexiva, simétrica, etc.

Así, un L -contexto (L, X, X, R) proporciona una nueva perspectiva del sistema relacional L -borroso (X, R) , en el que R puede ser un orden o un semi-orden L -borroso, o una semejanza L -borrosa, etc. Parece pues oportuno investigar si esa nueva visión constituye una herramienta útil en el estudio de propiedades del sistema (X, R) , así como en la traducción al marco borroso de ideas y nociones asociadas a relaciones sobre conjuntos nítidos.

En esta línea, en este trabajo analizamos esas dos situaciones. Concretamente:

1. Los L -conceptos $(A, B) \in \mathcal{L}(L, X, X, R)$ están relacionados con ciertos elementos o subconjunto L -borrosos del sistema (X, R) que sean distinguidos precisamente por propiedades que caracterizan a la L -relación R en X .
2. En esa interpretación de (X, R) como contexto (L, X, X, R) , y utilizando la clase de los L -conceptos pertenecientes a $\mathcal{L}(L, X, X, R)$, se puede extender, al campo de la lógica borrosa, nociones importantes asociadas a relaciones ordinarias (nítidas), de manera que esas extensiones resultan coherentes, es decir, si el sistema relacional (X, R) es nítido, la noción extendida recupera la noción de la que se parte.

¹ Miembro honorario del grupo de investigación de la Universidad Pública de Navarra: “Adquisición de conocimiento y minería de datos, funciones especiales y métodos numéricos avanzados”.

En lo que respecta a la situación 1., ya existe en la literatura algunos resultados en los que se encuentran vínculos entre L -contextos $(L, \mathbb{R}^n, \mathbb{R}^n, R')$ y $(L, \mathbb{Z}^n, \mathbb{Z}^n, R')$ y los sistemas algebraicos $(\mathbb{R}^n, R)$ y $(\mathbb{Z}^n, R)$ que aparecen en el tratamiento de señales e imágenes mediante técnicas de la Morfología Matemática [4,5]. Aquí, el retículo L tiene además una t -norma y una negación y R es una relación L -borrosa asociada a imágenes consideradas como elementos estructurantes [4]. R' representa su negación.

En este trabajo se incluye un nuevo ejemplo de esta situación 1 de carácter distinto: utilizando un sistema referencial (X, R) en el que $R \in L^{X \times X}$ es un orden L -borroso en X , justificamos que, para todo subconjunto L -borroso $M \in L^X$, el operador $\triangleleft$ proporciona una conexión entre las nociones: $MAX_R(M)$, $MIN_R(M)$, $SUP_R(M)$, $INF_R(M)$, [6,7,8] y $MAXIMAL_R(M)$, $MINIMAL_R(M)$ [7], (todos en L^X), y los L -conceptos de $\mathcal{L}(L, X, X, R)$.

En lo que se refiere a la situación 2., dada una L -relación $R: X \times X \rightarrow L$ reflexiva y simétrica (semejanza o tolerancia L -borrosa), e identificando el sistema relacional L -borroso (X, R) con un grafo L -borroso, (conjunto de vértices con aristas etiquetadas con elementos del retículo L), en este trabajo se utiliza el retículo de L -conceptos $(A, B) \in \mathcal{L}(L, X, X, R)$ del L -contexto asociado al grafo (X, R) para dar definiciones en el grafo L -borroso (X, R) de clique L -borroso $C \in L^X$ y de subconjunto L -borroso independiente $H \in L^X$. Además, se justifica que si la implicación $\rightsquigarrow: L \times L \rightarrow L$ es residuo de una t -norma continua a la izquierda, entonces el conjunto de cliques L -borrosos no es vacío y que todo clique L -borroso C verifica la ecuación relacional $(C \triangleleft R) = C$ [3]. También se prueba que estas definiciones son coherentes, en el sentido de que recuperan las nociones usuales de clique y de conjunto independiente en grafos ordinarios.

Palabras clave Retículos de conceptos borrosos, órdenes borrosos, cotas y extremos borrosos, grafos borrosos, cliques, conjunto independiente en grafos.

Referencias

- [1] L.J. Kohout, W. Bandler: *Use of fuzzy relations in knowledge representation, acquisition and processing*. In Fuzzy logic for the management of uncertainty. John Wiley & Sons.(1992) 415–435.
- [2] A. Burusco, R. Fuentes-González: *Concept Lattices defined from Implication Operators*. Fuzzy Sets and Systems **Vol 114 No 3** (2000) 431-436
- [3] C. Alcalde, A. Burusco, R. Fuentes-González : *Analysis of certain L-Fuzzy Relational Equations and the study of its Solutions by means of the L-Fuzzy Concept Theory* . International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems **Vol 20 No 1** (2012) 21–40.
- [4] C. Alcalde, A. Burusco, R. Fuentes-González: *An interpretation of the L-Fuzzy Concept Analysis as a tool for the Morphological Image and Signal Processing* In CLA 2012. (2012) 81–92.
- [5] C. Alcalde, A. Burusco, J.C. Díaz, R. Fuentes-González, J. Medina-Moreno : *Fuzzy Property-Oriented Concept Lattices in Morphological Image and Signal Processing*. In Advances in Computational Intelligence **Vol 7903** Springer (2013) 246–253.
- [6] R. Fuentes-González: *Cotas y extremos asociados a implicaciones difusas y a t-normas*. En Memorias del I Simposium de Inteligencia Artificial. CIMAF'97. La Habana, Cuba.(1997) 99-103.
- [7] P. Burillo, R. Fuentes-González, L. González: *On Completeness and Direction in Fuzzy Relational Systems* . Mathware & Soft Computing. **No 5** (1998) 243–251.
- [8] L. Běhounek: *Maxima and Minima in fuzzified linear orderings* . Fuzzy Sets and Systems **No 289** (2016) 82–93.

Nuevos órdenes sobre los números y cantidades borrosas

Pablo Hernández Varela², Carmen Torres-Blanc^{1,2},
Susana Cubillo Villanueva¹, José Atilio Guerrero²

¹ Universidad Politécnica de Madrid

scubillo@fi.upm.es, ctorres@fi.upm.es,

² Universidad Nacional Experimental del Táchira, Venezuela

phernandezv@unet.edu.ve, jaguerrero4@gmail.com

En los procesos de los sistemas borrosos, comúnmente se debe comparar y elegir entre varios números borrosos para emitir una respuesta o tomar una decisión. En cierto sentido, los números y cantidades borrosas se pueden considerar como "valores reales" con grados de incertidumbre, imprecisión, o vaguedad, y por tanto se representan como conjuntos borrosos sobre la recta real. Es decir, en general, no son simplemente números reales, y no se pueden ordenar con el orden total de los mismos. En lo que sigue, para simplificar las notaciones, identificaremos un conjunto borroso F sobre la recta real con su función de pertenencia f .

Por otra parte, el orden parcial usual de los conjuntos borrosos inducido por el orden de los números reales, es decir

$$f \leq g \text{ si } f(x) \leq g(x), \text{ para todo } x \in \mathbb{R} \quad (1),$$

no extiende el orden de los números reales. En efecto, supongamos que se quiere tratar a dos números reales a y b como números borrosos, por medio de sus funciones características $\bar{a}$ y $\bar{b}$, entonces el orden parcial usual (1), no extiende el orden usual ($\leq$) de los números reales. Por ejemplo, si $a, b \in \mathbb{R}$, con $a < b$, no se cumple que $\bar{a} \leq \bar{b}$, ya que $\bar{b}(a) = 0 < \bar{a}(a) = 1$. Además, tampoco se cumple $\bar{b} \leq \bar{a}$, por lo que no son comparables con dicho orden parcial. Por tanto, el orden parcial (1) no es recomendable para comparar dos números o cantidades borrosas. Por este motivo, en [1] se presenta el orden parcial que se obtiene a partir de las extensiones en forma horizontal del Principio de Extensión de Zadeh

$$(f \bullet g)_\alpha = f_\alpha \bullet g_\alpha = \{a \circ b : a \in f_\alpha, b \in g_\alpha\} \quad (2),$$

tomando como $\circ$ las operaciones $\vee$ y $\wedge$; es decir, el siguiente orden parcial:

$$f \leq g \text{ si } f_\alpha \leq_I g_\alpha, \text{ para todo } \alpha \in (0,1], \quad (3)$$

donde f_α representa el α -corte correspondiente de la función f , y $[a_1, a_2] \leq_I [b_1, b_2]$ si $a_1 \leq b_1$ y $a_2 \leq b_2$. Este orden parcial sí extiende el orden de los números reales, ya que dados dos números reales a y b tales que $a \leq b$, se tiene que $\bar{a} \leq \bar{b}$, al ser $\bar{a}_\alpha = [a, a]$ y $\bar{b}_\alpha = [b, b]$ para todo α , y por consiguiente $\bar{a}_\alpha \leq_I \bar{b}_\alpha$ para todo α . Sin embargo, este orden parcial no está definido para cualesquiera conjuntos borrosos sobre $\mathbb{R}$, y sólo es válido para ordenar parcialmente conjuntos borrosos con funciones de dominio real convexas, semicontinuas superiormente (upper-semicontinuas), con soporte acotado y fuertemente normales (lo que llamaremos números borrosos, FNs), condiciones que aseguran que los α -cortes sean intervalos cerrados y acotados.

En este trabajo, proponemos el uso de dos órdenes parciales, obtenidos a partir de la forma vertical del Principio de Extensión de Zadeh ([4]), de las operaciones $\vee$ y $\wedge$:

$$(f \sqcap^* g)(x) = \sup \{f(y) \wedge g(z) : y \wedge z = x\} \text{ y } (f \sqcup^* g)(x) = \sup \{f(y) \wedge g(z) : y \vee z = x\} \quad (4)$$

Teniendo en cuenta estas operaciones, Walker et al. ([3]) obtuvieron en el conjunto de las funciones de $[0,1]$ en $[0,1]$ dos órdenes parciales, estudiando sus propiedades. En este trabajo, proponemos una extensión de estos órdenes al conjunto M^* de las funciones de $\mathbb{R}$ en $[0,1]$. Así consideramos los órdenes parciales

$$f \sqsubseteq^* g \text{ si } f \sqcap^* g = f \text{ y } f \leq^* g \text{ si } f \sqcup^* g = g, \text{ siendo } f, g \text{ en } M^* \quad (5).$$

Al ser aplicados a todo M^* son más generales que el orden parcial dado en (3).

Los órdenes parciales dados en (5), también extienden el orden de los números reales, puesto que se verifica $\bar{a} \sqsubseteq^* \bar{b}$ y $\bar{a} \leq^* \bar{b}$ si $a \leq b$, lo que se demostrará utilizando una caracterización de estos órdenes en el caso de funciones convexas y con la misma altura (de hecho, para estas funciones los dos órdenes dados en (5) coinciden).

Utilizando estos órdenes parciales y sus caracterizaciones comparamos números borrosos cuyas funciones de pertenencia no sean necesariamente semicontinuas superiormente o no tengan soporte acotado, casos en los que no se puede emplear la representación horizontal del orden de los números borrosos. Además, basados en las caracterizaciones de estos órdenes parciales, se propondrá un nuevo método que permitirá comparar totalmente (linealmente) números borrosos. En dicho método se obtendrá un pre-orden total, ya que la relación no será antisimétrica, en el conjunto de los números borrosos generalizados extendidos, es decir, el conjunto de los FSs sobre $\mathbb{R}$ convexas. También se obtendrán otras propiedades deseables en los métodos de ordenación, según se propone en [2].

La idea en que nos basamos, es determinar, en general, qué caracterización del orden parcial “se cumple más”, si la que resulta de $f \sqsubseteq^* g$, o la que resulta de $g \sqsubseteq^* f$, calculando ciertas áreas determinadas por las funciones de pertenencia de los conjuntos a comparar.

Cabe destacar que, en trabajos previos no se ha encontrado evidencia, al menos directamente, sobre el empleo de la definición y caracterización de los órdenes parciales obtenidos a partir de la forma vertical del principio de extensión de Zadeh, para comparar números o cantidades borrosas.

Referencias

- [1] Kerre, E., Mares, M. y Mesiar, R.: *Generate fuzzy quantities and their orderings*. Information, Uncertainty and Fusion, edited by Bernadette Bouchon-Meunier, Ronald R. Yager, Lotfi Zadeh, Springer Science and Business Media, (2012) 119–130.
- [2] Wang, X y E. Kerre, E.: *Reasonable properties for the ordering of fuzzy quantities (I)*. Fuzzy Sets and Systems, **118**, (2001) 375–385.
- [3] Walker, C. y Walker, E.: *The algebra of fuzzy truth values*. Fuzzy Sets and Systems, 149, (2005) 309–347.
- [4] Zadeh, Lotfi: *The concept of a linguistic variable and its application to approximate reasoning-I*. Inf. Sci. **8**, (1975) 199–249.

Este trabajo ha sido financiado parcialmente por la UPM (España)

Application of F -transforms to Reduction of Relational Databases

Petra Hodáková¹, Nicolás Madrid², Pavel Rusnok¹

¹ *University of Ostrava, Centre of Excellence IT4Innovations,
Institute for Research and Applications of Fuzzy Modeling,
30. dubna 22, 701 03 Ostrava 1, Czech Republic
{petra.hodakova, pavel.rusnok}@osu.cz*

² *Departamento de Matemática Aplicada. Universidad de Málaga
Doctor Pedro Ortiz Ramos, s/n, CP:29071, Málaga, Spain
nicolas.madrid@uma.es*

The development and applications of fuzzy techniques in databases has had a great interest in the last three decades. There are mainly two reasons for such an interest. On the one hand, there is a necessity of performing fuzzy queries in databases (non necessarily fuzzified); for instance, one user can be interested in “sport and small cars around 20,000\$”. A family of approaches dealing with this kind of tasks is commonly denoted as Flexible Query Answering, see e.g. [1]. On the other hand, there is a necessity of representing (and working with) fuzzy data in databases. For instance data that is subjective, usually obtained by surveys and enriched with uncertainty. A database of sportmen with parameters representing speed, technique, resistance, etc is a good example of such kind of data. There are many ways to carry out this kind of tasks (usually depending on the kind of data represented) such as Possibilistic databases [2], Fuzzy relational databases [3] or Fuzzy object-oriented databases [4].

Among the fuzzy techniques defined on relational databases we point out those related to Fuzzy Association Rules [5] and Fuzzy Concept Lattices [6]. The goal of these two approaches is to represent the knowledge appearing in relational databases by *inference rules* and by a *lattice of concepts*, respectively. In both kind of approaches (i.e. in the creation of Fuzzy Association Rules and Fuzzy Concept Lattices), it is usually necessary to do a reduction of the information retrieved. There are mainly two ways to perform such a reduction, either to reduce the information from the original database (also called context) or to reduce the information obtained from the fuzzy technique (either from the inference system or from the concept lattice). In this talk, we focus on the first kind of reduction i.e., on reducing the size of the relational database by keeping the significant information contained in it.

In general, the relational databases carry the information about several attributes of plenty of objects. The idea underlying in the technique described in this talk is to create a new set of “fuzzy” objects. These new objects are representative among those satisfying a fuzzy attribute, for instance a “quite big house” or a “barely power computer”. These new objects are defined by using direct F -transforms [7]. The F -transform is generally an approximation technique. Its advantage is that by using the direct F -transform we compress and reduce the information about the original objects, and by using the respective inverse F -transform we obtain back the approximation of the relations associated to the original objects.

Finally, in the talk we also show how the technique is potentially applicable, firstly to the restoration of imperfect information, secondly to the reduction of Fuzzy Concept Analysis and thirdly to the achievement of Fuzzy Association Rules. Note that to obtain fuzzy inference rules from a relational database it is needed a fuzzification in some step. In addition, in fuzzy concept lattice the step of fuzzification is missing in the most of approaches (perhaps because the goal of the most papers in the literature is to present the technique, not its real application). So the technique presented here can be considered also a fuzzification procedure for relational databases in both kind of approaches.

References

- [1] T. Andreasen, H. Christiansen, J. Kacprzyk, H. L. Larsen, G. Pasi, O. Pivert, G. De Tré, M. A. Vila, A. Yazici and S. Zadrozny: *Flexible Query Answering Systems 2015. Proceedings of the 11th International Conference FQAS 2015*. Advances in Intelligent Systems and Computing. **No. 400** (2016). Springer.
- [2] P. Bosc and O. Pivert: *About projection-selection-join queries addressed to possibilistic relational databases*. IEEE Transaction on Fuzzy Systems. **No. 1** Vol. 13 (2005) 124 – 139.
- [3] S. Shenoit and A. Melton: *An extended version of the fuzzy relational database model*. Information Sciences. **No. 1** Vol. 52 (1990) 35–52.
- [4] F. Berzal, N. Marín, O. Pons and M. A. Vila: *Managing fuzziness on conventional object-oriented platforms*. International Journal of Intelligent Systems. **No. 7** Vol. 22 (2007) 781–803.
- [5] R. Agrawal, T. Imielinski, A. Swami: *Mining Association Rules between Sets of Items in Large Databases*. Proceedings of the International Conference on Management of Data (ACM SIGMOD), 1993.
- [6] J. Medina, M. Ojeda-Aciego and J. Ruiz-Calviño: *Formal concept analysis via multi-adjoint concept lattices*. Fuzzy Sets and Systems. **No. 2** Vol. 160 (2009) 130–144.
- [7] I. Perfilieva: *Fuzzy transforms: Theory and applications*. Fuzzy Sets and Systems. **No. 8** Vol. 157 (2006) 993–1023.

Removing redundancy for attribute implications in data with grades

E. Rodríguez-Lorenzo, Pablo Cordero, Manuel Enciso, Ángel Mora

¹ Universidad de Málaga. Andalucía Tech. Spain

{estrellarodlor, amora}@ctima.uma.es, {pcordero, enciso}@uma.es

Extended abstract

Reasoning with if-then rules –in particular, with those taking from of implications between conjunctions of attributes– is crucial in many disciplines ranging from theoretical computer science to applications. One of the most important problems regarding the rules is to remove redundancies in order to obtain equivalent implicational sets with lower size. This problem, which has been widely studied in the classical setting, is partially addressed in [4] for the case of attribute implications in data with grades. In this work, we tackle this problem using the so-called *Fuzzy Attribute Simplification Logic*, FASL, which has been introduced in [1]. This logic leads to an automatic reasoning method for implications in data with grades.

FASL manages formulas $A \Rightarrow B$ where A and B are fuzzy sets over an alphabet Ω whose elements are named attributes. Interpretations are fuzzy formal contexts and, informally, an implication such as $\{a,^{0.5}/b\} \Rightarrow \{^{0.9}/c\}$ means every object that has attribute a to degree 1 (i.e. fully possesses a), and attribute b to degree 0.5, has attribute c to degree at least 0.9.

Specifically, truthfulness structures in FASL are tuples $\langle L, \vee, \wedge, \otimes, \rightarrow, \searrow, *, 0, 1 \rangle$ where $\langle L, \vee, \wedge, \otimes, \rightarrow, 0, 1 \rangle$ is a complete residuated lattice, $*$ is an hedge (a “very true” function [2]) and $\searrow$ is a binary operation satisfying the following adjointness property: $a \searrow b \leq c$ if and only if $a \leq b \vee c$ for all $a, b, c \in L$. As a consequence, $a \searrow b = \bigwedge \{c \in L \mid a \leq b \vee c\}$. In particular, when the lattice is linearly ordered, we have

$$a \searrow b = \begin{cases} a, & \text{if } a > b, \\ 0, & \text{otherwise.} \end{cases} \quad (1)$$

These operations are pointwise extended to fuzzy sets.

For a fuzzy formal context $\mathbb{K} = \langle X, Y, I \rangle$ such that $Y = \Omega$, the degree in which $A \Rightarrow B$ is true in $\mathbb{K}$ is defined as follows:

$$\|A \Rightarrow B\|_{\mathbb{K}} = \bigwedge_{x \in X} \left(\left(\bigwedge_{y \in Y} (A(y) \rightarrow I(x, y)) \right)^* \rightarrow \bigwedge_{y \in Y} (B(y) \rightarrow I(x, y)) \right) \quad (2)$$

The axiomatic system in FASL is defined as follows: for all $A, B, C, D \in L^{\Omega}$ and $c \in L$,

$$[\text{Ax}] \text{ infer } A \cup B \Rightarrow A \quad (\text{Axiom})$$

Supported by project TIN2014-59471-P of the Science and Innovation Ministry of Spain, co-funded by the European Regional Development Fund (ERDF).

[Mul] from $A \Rightarrow B$ infer $c^* \otimes A \Rightarrow c^* \otimes B$ (Multiplication)

[Sim] from $A \Rightarrow B$ and $C \Rightarrow D$ infer $A \cup (C \setminus B) \Rightarrow D$ (Simplification)

The soundness and completeness are ensured when we assume that both L and Ω are finite. In addition, for any $A, B, C, D \in L^Y$, the following equivalences hold true:

(DeEq) $\{A \Rightarrow B\} \equiv \{A \Rightarrow B \setminus A\};$

(UnEq) $\{A \Rightarrow B, A \Rightarrow C\} \equiv \{A \Rightarrow B \cup C\};$

(SiEq) If $A \subseteq C$ then $\{A \Rightarrow B, C \Rightarrow D\} \equiv \{A \Rightarrow B, A \cup (C \setminus B) \Rightarrow D \setminus B\}.$

In [4], V. Vychodil consider *globalization* as hedge and provides a polynomial algorithm for computing a non-redundant equivalent set for a given implicational set. An implicational set T is said to be non-redundant if no implications in T can be inferred from others implications. The main goal here is in the same direction, but considering an arbitrary hedge and using FASL. As a first stage in this line, we propose the use of **(DeEq)**, **(UnEq)** and **(SiEq)** for removing redundant information in an implicational set. Specifically, the method here proposed crosses the sets of implications (quadratic complexity) taking pairs of implications and applies the above mentioned equivalences.

For instance, consider $\langle \{0, 0.5, 1\}, \vee, \wedge, \otimes, \rightarrow, \setminus, *, 0, 1 \rangle$ being the three-element equidistant subchain of the standard Łukasiewicz algebra with $*$ being the identity and $\setminus$ as described in (1). For the implicational set

$$T_1 = \{ \{0.5/x, y, 0.5/z\} \Rightarrow \{x, 0.5/y, 0.5/z\}, \{0.5/y\} \Rightarrow \{0.5/x, 0.5/y, z\}, \\ \{0.5/x, 0.5/y\} \Rightarrow \{y, z\}, \{x, z\} \Rightarrow \{0.5/x, y, 0.5/z\}, \{z\} \Rightarrow \{0.5/x\} \}$$

in [4], the author obtains

$$T_2 = \{ \{0.5/y\} \Rightarrow \{x, y, z\}, \{z\} \Rightarrow \{0.5/x\} \{x, z\} \Rightarrow \{0.5/x, y, 0.5/z\} \}$$

However, our method, which has been implemented in PROLOG, renders the following equivalent set:

$$T_3 = \{ \{0.5/x\} \Rightarrow \{x\}, \{0.5/y\} \Rightarrow \{0.5/x, y, z\}, \{z\} \Rightarrow \{y\} \}$$

References

- [1] Radim Belohlavek, Pablo Cordero, Manuel Enciso, Ángel Mora, Vilem Vychodil: *Automated prover for attribute dependencies in data with grades*. International Journal of Approximate Reasoning. Vol. 70 (2016) 51–67.
- [2] Radim Belohlavek, Vilem Vychodil. Attribute implications in a fuzzy setting. Lecture Notes in Computer Science. Vol. 3874, 45–60.
- [3] Radim Belohlavek, Vilem Vychodil. Attribute implications in a fuzzy setting. Lecture Notes in Computer Science. Vol. 3874, 45–60. Vilem Vychodil. On minimal sets of graded attribute implications, Information Sciences. Vol. 294, 478–488.
- [4] Vilem Vychodil: *On minimal sets of graded attribute implications*. Information Sciences. Vol. 294 (2015) 478–488.

Un método para obtener una representación difusa de series de tiempo

Antonio Moreno-Garcia¹, Juan Moreno-Garcia¹

¹ *Departamento de Tecnologías y Sistemas de Información, Uni. de Castilla-La Mancha*
antmorgarcia@gmail.com, juan.moreno@uclm.es

Las series de tiempo (TS) son utilizadas en infinidad de problemas y contienen información importante del ámbito de donde han sido capturadas. Uno de sus problemas consiste en que tratarlas suele ser un proceso complejo ya que están formadas por una gran cantidad de puntos. Por otro lado, la lógica difusa permite modelar la incertidumbre del proceso de captura y trabajar con los datos de una forma cualitativa mediante sus operadores.

Este trabajo presenta una representación difusa de una TS y cómo obtenerla. El proceso consta de dos fases:

1. Se calcula una representación basada en rectas de la TS mediante un nuevo método.
2. Transformación de dicha representación en una estructura difusa que puede ser utilizada en diferentes problemas.

1. Transformación de la TS en un conjunto ordenado de segmentos S

La TS será representada mediante un conjunto ordenado de segmentos, similar a las propuestas de Keogh [1], y su expresión se muestra en la Ecuación 1, donde cada segmento $s_{f_i, l_i} = (m_{f_i, l_i} \times x) + c_{f_i, l_i}$ es válido en el intervalo de tiempo definido entre los instantes f_i y l_i . Un segmento $s_{f, l}$ es formalmente representado como $\{m_{f, l}, c_{f, l}\}$.

$$S = \{s_{f_1, l_1}, s_{f_2, l_2}, \dots, s_{f_m, l_m}\} \quad (1)$$

Nuestro método es novedoso y está basado en el uso de ventanas deslizantes y regresión lineal. Las ventanas deslizantes permiten modelar cada segmento de la TS que se representan utilizando la ecuación de la recta entre dos instantes f_i y l_i . La Figura 1 muestra un ejemplo de una TS obtenida por nuestro método.

2. Conversión del conjunto de segmentos en una estructura difusa T

Kacprzyk et al. presentan algunas ideas en este sentido en [2]. El modelado del conjunto de tendencias se realizará mediante la tupla $T = \{t_0, LT\}$, donde t_0 es un conjunto difuso que indica el instante temporal donde empieza la lista de tendencias, y $LT = \langle TEND_1, \dots, TEND_n \rangle$ es la lista de tendencias propiamente dicha. Cada elemento de LT se modela mediante $TEND_i = \{type_i, t_i, v_i, power_i\}$ donde:

Este trabajo ha sido financiado bajo el proyecto de investigación TIN2015-64776-C3-3-R del Ministerio de Economía y Competitividad.

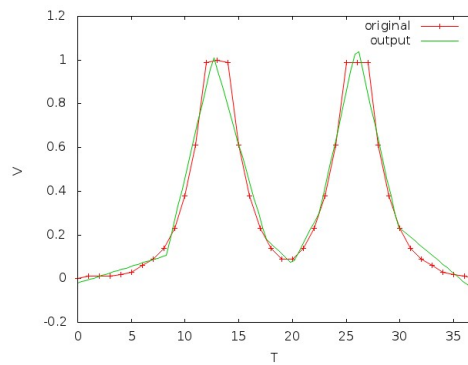


Figura 1: Ejemplo de TS obtenida.

- $type_i$ indica el tipo de tendencia tomando uno de los valores *ZERO*, *INC* o *DEC*, se calculará en base a las pendientes de los segmentos.
- t_i es un número difuso triangular que define el instante temporal donde acaba el intervalo que representa el segmento. El origen del segmento es t_{i-1} de $TEND_{i-1}$.
- v_i es un número difuso que define el valor final de salida del último valor del intervalo. El valor del primer punto del segmento se puede obtener de v_{i-1} de $TEND_{i-1}$.
- $power_i \in SLL$ es una etiqueta lingüística que modela la potencia de incremento o decremento. El conjunto ordenado de etiquetas lingüísticas *SLL* se define apriori. Se utiliza la pendiente de la recta para calcular esta etiqueta.

El soporte de los números difusos t_i y v_i se calcula utilizando una medida del error de la recta obtenida respecto a los puntos originales.

3. Conclusiones

Se proponen dos métodos de representación de TS. La primera está basada en segmentos obtenidos con regresión lineal que posteriormente se modela mediante un Conjunto Ordenado de Segmentos, también se describe brevemente el algoritmo para obtenerlas representaciones a partir de la serie. La segunda está basada en la lógica difusa y es obtenida desde el conjunto ordenado de segmentos.

Referencias

- [1] E. Keogh: *A fast and robust method for pattern matching in time series databases*. Proceedings of the Ninth IEEE International Conference on Tools with Artificial Intelligence. (1997) 578–584.
- [2] J. Kacprzyk, A. Wilbik, S. Zadrozny: *Linguistic summarization of time series using a fuzzy quantifier driven aggregation*. Fuzzy Sets and Systems. **159(12)** (2008) 1485–1499.

Aprendizaje de Reglas de Asociación Difusas en Flujo de Datos para el Análisis del Electroencefalograma en Psicofisiología

Elena Ruiz ¹, Pandelis Perakakis ², Jorge Casillas ¹

¹ Dept. Ciencias de la Computación e Inteligencia Artificial, Universidad de Granada
eruiz@decsai.ugr.es, casillas@decsai.ugr.es

² Centro de Investigación Mente, Cerebro y Comportamiento, Universidad de Granada
peraka@ugr.es

Hoy en día prácticamente la totalidad de la actividad humana es registrada (o susceptible de serlo). A veces, esta información es empleada para obtener modelos (minería de datos), generalmente con el objetivo de predecir comportamientos o tendencias. Sin embargo, otras veces el interés principal reside en supervisar el sistema. Es cada vez más habitual que estos datos provengan de fuentes de datos que, de forma continua, generan información ordenada cronológicamente y que superan las capacidades de almacenamiento y procesamiento habituales.

Para abordar este problema, se pueden manejar flujos de datos (*data streams*), secuencias infinitas de registros estructurados que llegan continuamente –es decir, no son datos transitorios sino flujos de información duradera en el tiempo– [1, 2]. La característica clave de estos sistemas es que los datos de estos flujos no se almacenan de forma permanente sino que se procesan “al vuelo”, esto es, cada dato es analizado, procesado y finalmente olvidado. Las aplicaciones reales de este problema crecen cada día, siendo frecuente el uso en telecomunicaciones, tráfico de vehículos, consumo energético, comercio, finanzas, fisiología, robótica o redes sociales, entre otras.

La mayoría de la literatura especializada en este campo se centra en clasificación (y *concept drift*), es decir, se asume la existencia de un flujo de datos etiquetado con una clase. Este tipo de tarea supervisada presenta una falta de aplicabilidad en problemas reales al ser necesario disponer de un experto (máquina, persona o comunidad) que etiquete cada ejemplo que llega desde el flujo de datos. Esto hace que sea habitual encontrarnos con experimentación basada solo en *benchmarks* y datos sintéticos –por ejemplo, SEA [3]-. Disponer de datos estructurados y etiquetados en todo momento no resulta habitual en problemas reales, y cuando se dispone de ellos (por ejemplo, actividad web), la habilidad predictiva del algoritmo resulta irrelevante. Un enfoque de aprendizaje descriptivo parece tener más sentido en flujo de datos. En este caso, el modelo se adapta a los datos para explicar qué está ocurriendo en cada instante. Se utiliza, por tanto, para monitorizar el sistema, no para predecir salidas. Se requiere modelar relaciones complejas entre las variables de una forma fácilmente interpretable.

Esto lleva, dentro del campo de flujo de datos, a lo que se ha dado a conocer como *association stream mining* [9], donde el objetivo no está centrado en una determinada variable dependiente sino en el descubrimiento de patrones de relaciones entre todas ellas. Al contrario de los conocidos enfoques para encontrar ítems frecuentes (*frequent items*) en flujos de datos [4, 5], el campo de *association stream mining* consiste en generar de forma continua modelos causales en dominios complejos. Su objetivo es extraer asociaciones de interés entre los atributos de los datos de modo no supervisado y adaptándose al entorno. La principal limitación del actual estado del arte en aprendizaje de patrones a partir de flujo de datos reside en la dificultad para extraer reglas que definan relaciones de causalidad al no

poder aplicar el enfoque tradicional en dos fases usado en datos estáticos, donde primero se obtienen objetos frecuentes y luego se extraen reglas, algo imposible cuando los datos llegan en modo de flujo. No obstante existen unas pocas propuestas basadas en mantener reglas de asociación y actualizarlas conforme se reciben datos [6, 7, 8] pero resultan muy ineficientes e impracticables en problemas con una afluencia de datos relativamente alta. Además, la mayoría de las soluciones tratan solo variables nominales, lo cual limita su aplicación en problemas reales (por ejemplo, mediciones a través de sensores).

En este artículo proponemos varias mejoras al algoritmo Fuzzy-CSar [9] para extraer reglas de asociación difusas en flujo de datos. Se proponen dos ejes principales: por un lado, el uso de una nueva representación y operadores genéticos para los atributos que forman el flujo de datos que permite manejar particiones difusas de distinta granularidad, -pudiéndose adaptar así a las necesidades de precisión de cada variable en cada regla. Por otro lado, en problemas reales con flujos continuos de información es muy posible que desconozcamos cuál va a ser el dominio utilizado por ciertos atributos. Por esto, en esta nueva versión se incorpora al algoritmo un método por el cual no es necesario conocer el rango en el que se moverán las variables, los valores mínimos y máximos para cada atributo de entrada se van actualizando en tiempo real a medida que los datos van siendo procesados.- Todas las modificaciones se han realizado con el objetivo de dotar al sistema de una mayor flexibilidad que le permita adaptarse mejor a las peculiaridades de cada problema y a las condiciones de los problemas reales de flujo de datos.

Los resultados del nuevo algoritmo se validan en un problema real del ámbito de la Psicofisiología donde se analizan asociaciones entre las distintas señales de electrodos del encefalograma en sujetos sometidos a diferentes estímulos.

References

- [1] Gama, J.: *Knowledge Discovery from Data Streams*. Chapman and Hall/CRC. (2010).
- [2] Bifet, A., Holmes, G., Kirkby, R., Pfahringer, B.: *MOA: Massive Online Analysis*. Journal of Machine Learning Research. 11 (2010) 1601-1604.
- [3] Street, W.N., Kim, Y.: *A streaming ensemble algorithm (SEA) for large-scale classification*. Proceedings of the 7th ACM SIGKDD international conference on Knowledge discovery and data mining. (2001) 377-382.
- [4] Karp, R.M., Shenker, S., Papadimitriou, C.H.: *A simple algorithm for finding frequent elements in streams and bags*. ACM Trans. Database Syst. 28 (2003) 51-55.
- [5] Cormode, G., Korn, F.: *Finding hierarchical heavy hitters in streaming data*. ACM Transactions on Knowledge Discovery from Data. 1:4, article no. 16 (2008) 1-48.
- [6] Fan, W., Watanabe, T., Asakura, K.: *Ratio rules mining in concept drifting data streams*. Proceedings of the World Congress on Engineering and Computer Science 2009 Vol II, WCECS 2009, San Francisco, USA. (2009) 809-814.
- [7] Chen, P., Su, H., Guo, L., Qu, Y.: *Mining fuzzy association rules in data streams*. 2nd International Conference on Computer Engineering and Technology, ICCET. (2010) 153-158.
- [8] Tan, J., Bu, Y., Zhao, H.: *Incremental maintenance of association rules over data streams*. 2nd International Conference on Networking and Digital Society (ICNDS). (2010) 444-447.
- [9] Sancho-Asensio, A., Orriols-Puig A., Casillas J.: *Evolving Association Streams*. Information Sciences 334-335 (2016) 250-272.

An Approach to Multilabel Classification using $\mathcal{K}$ -Formal Concept Analysis

Francisco J. Valverde-Albacete, Carmen Peláez-Moreno

*Departamento de Teoría de la Señal y de las Comunicaciones.
Universidad Carlos III de Madrid, 28911 Leganés, Spain
{fva,carmen}@tsc.uc3m.es*

Motivation. Multi-label classification (MLC) is a relatively-recently formalized task in Machine Learning [1], encompassing tasks like *text categorization*, *gene expression analysis*, etc. Consider a feature space $\mathcal{X} \equiv \mathcal{K}^F$. A $\mathcal{K}$ -valued F -vector $x \in \mathcal{X}$ is an *instance* to be tagged with any combination of a set of *labels* $\mathcal{L}$. Label associations or *labelsets* are represented by means of their characteristic vectors $y \in \mathcal{Y} \equiv 2^{\mathcal{L}}$, and pairs of instances and labelsets are called *examples or observations*, $(x, y) \in \mathcal{X} \times \mathcal{Y}$.

Given a set of *training observations* $\mathcal{D}_T = \{(x^j, y^j)\}_{j=1}^{n_T}$ and a set of *test observations* $\mathcal{D}_E = \{(x^k, y^k)\}_{k=1}^{n_E}$ the purpose of the *multi-label classification task* is to find a (*multi-label*) *classifier* function $f: X \rightarrow Y$ using the information in the training examples $\mathcal{D}_T$ that generates (*multi-label*) *predictions* $\hat{y} = f(x)$ on the test *instances* optimizing a certain performance measure $d: Y \times Y \rightarrow P$ where P is any set of figure-of-merit values (such as precision, recall, F-measure, Hamming loss, etc.) Table 1 shows a selection of low- to middle-complexity datasets used as test banks for the MLC task.

Name	$n_E + n_T$	F	L	distinct
emotions	391+202	72	6	27
scene	1 211+1 196	294	6	15
yeast	1 500+917	103	14	198

Table 1: A selection of multi-label classification databases (data from [2]). “Distinct” refers to distinct label sets occurring in the training decision data.

Several issues have to be dealt with when trying to solve the task, such as a) Classifier and performance measure selection, b) Data preparation, c) Classifier construction, and d) Classifier evaluation.

Since the MLC task can be considered a strict generalization of the *binary and multi-class classification* tasks in that instances may have more than one label (class) assigned to them, most of the techniques for classifier selection and construction have been imported therefrom—albeit with a limited success.

On the other hand, performance measure selection, data preparation and classifier evaluation have required extensions to cater for the peculiarities of MLC: since the theory of

CPM & FVA have been partially supported by the Spanish Government-MinECo projects TEC2014-53390-P and TEC2014-61729-EXP.

machine learning is firmly-grounded only for the binary, mutually-exclusive labelling cases, dealing with label sets poses a challenge usually solved by means of *problem transformation*. Two extreme cases of these transformations are [2] are the **Label Powerset (LP)** method, and the **Binary relevance (BR)** method.

Results. In this paper we propose an entirely new approach to multi-label classification based on $\mathcal{K}$ -Formal Concept Analysis, an extension to FCA, where the incidences take value in an idempotent semifield [3].

Consider a relation $R \in \mathcal{K}^{L \times F}$. A Galois connection between $\mathcal{X} \equiv \mathcal{K}^L$ and $\mathcal{Y} \equiv \mathcal{K}^F$ is created from inequation $x^\top \otimes R \otimes y \leq \sigma$, where $x \in \mathcal{X}$, $y \in \mathcal{Y}$ and $\sigma \in \mathcal{K} \setminus \{\perp, \top\}$ by using residuation to give closed formulae for the polars of the connection:

$$x_{R,\sigma}^\uparrow = R^{\circledast} \dot{\otimes} x^{-1} \dot{\otimes} \sigma \quad y_{R,\sigma}^\downarrow = R^{-1} \dot{\otimes} y^{-1} \dot{\otimes} \sigma \quad (1)$$

Then it is clear that for (u, v) being a σ -formal concepts of $\underline{\mathfrak{B}}^\sigma(L, F, R)_{\mathcal{K}}$ we may write,

$$R \leq u^{-1} \dot{\otimes} \sigma \dot{\otimes} v^{\circledast} \quad (2)$$

This proves that σ -concepts are part of a decomposition of R whence we may use (2) to synthesize $R \leq \bigvee \{u^{-1} \dot{\otimes} \sigma \dot{\otimes} v^{\circledast} \mid (u, v) \in \underline{\mathfrak{B}}^\sigma(L, F, R)_{\mathcal{K}}\}$. To solve this recursive equation we propose a set of concepts that generate the required equation, and this sets the inductive bias of the learning method:

$$R \approx Y_T^{-1} \dot{\otimes} \Sigma \dot{\otimes} X_T^{\circledast} \quad (3)$$

where a diagonal matrix $\Sigma \in \mathcal{K}^{n_T \times n_T}$ scales the individual concepts. It can be proven that this formula conforms the training observation intents to the training labelset extents and vice-versa.

In particular, for $\sigma = I_{n_T}$ we have $\hat{Y}_T = X_T^\downarrow$ and $\hat{Y}_E = X_E^\downarrow$ which are the predictions for the training and test sets, respectively. By optimizing a threshold with respect to the performance measures $\hat{\phi}$ we can obtain a measure of the fitness of the approach on the test set. Note that since the Σ is actually a matrix of singular values for decomposition (2) it pays to try to find those singular values that better approximate R .

References

- [1] J. Read. *Scalable Multi-label Classification*. PhD thesis, The University of Waikato, Australia, 2010.
- [2] G. Tsoumakas, I. Katakis, and I. Vlahavas. Mining multi-label data. In O. Maimon and L. Rokach, editors, *Data Mining and Knowledge Discovery Handbook*, pages 667–685. Data Mining and Knowledge Discovery, 2010.
- [3] F. J. Valverde-Albacete and C. Peláez-Moreno. The Linear Algebra in Formal Concept Analysis over Idempotent Semifields. In *Formal Concept Analysis*, number 9113 in LNAI, pages 97–113. Springer Berling Heidelberg, Heidelberg, Apr. 2015.

On the Application of Data Mining Techniques to Graded Ontology Building

Fernando Bobillo¹, M. Dolores Ruiz², Juan Gómez-Romero², Daniel Sánchez²

¹ Aragon Institute of Engineering Research (I3A), University of Zaragoza, Spain

² Dpt. of Computer Science and A. I., CITIC-UGR, University of Granada, Spain

Email: fbobillo@unizar.es, {mdruiz, jgomez, daniel}@decsai.ugr.es

In recent years, we are witnessing an increase in the development of *graded* (including *fuzzy*, *probabilistic*, and *possibilistic*) ontology applications, which are more appropriate to several real world domains with uncertain, imperfect, incomplete, imprecise, or ambiguous knowledge [3]. In graded ontologies, axioms might have a degree associated. The semantics of such degrees depends on the formalism: in fuzzy ontologies it indicates to what extent the axiom is satisfied; in probabilistic/possibilistic ontologies it estimates to what extent we are confident that it is true.

Despite the existence of some applications and some methodologies, building graded ontologies is still a hard problem. Since data mining techniques have proved to be very useful to discover meaningful and interesting knowledge for the user (even in the presence of uncertainty) [1], the objective of this paper is to discuss how data mining could help in the development of graded ontologies. In particular, we will discuss how to learn graded ontology axioms from association rules.

We propose to study some of the many statistical measures for association rules in the literature [1, 2]. Depending on the type of axiom that we want to build, we may use different measures for learning the degrees. Let $x_1 \Rightarrow y_1$ be an association rule expressing the relation between the antecedent x_1 and the consequent y_1 in a database D . Let a denote the number of transactions satisfying both the antecedent and the consequent, b the number of transactions satisfying the antecedent but not the consequent, c the number of transactions satisfying the consequent but not the antecedent, and d the number of transactions satisfying neither the antecedent nor the consequent. For example, in these terms the confidence-descriptive measure is defined as $\frac{a}{a+b}$.

Let us assume that we have a database where the transactions include information about some attributes $A_1, A_2, \dots, A_n$. Optionally, we may also have that the transactions are classified, i.e., we have information about some class that every transaction belongs to. This scenario is illustrated in Figure 1 (a).

Conventional data mining algorithms can obtain association rules of the form $x_1 \Rightarrow y_1$. Such rules can be seen as graded ontology axioms of the form $\langle \exists A_1.\{x_1\} \sqsubseteq \exists A_2.\{y_1\}, \alpha \rangle$, for some degree α [3]. To learn such axioms, we have to search for those

We were partly funded by the CICYT project TIN2013-46238-C4-4-R and University of Granada (Programa de Proyectos para la incorporación de jóvenes doctores a nuevas líneas de investigación).

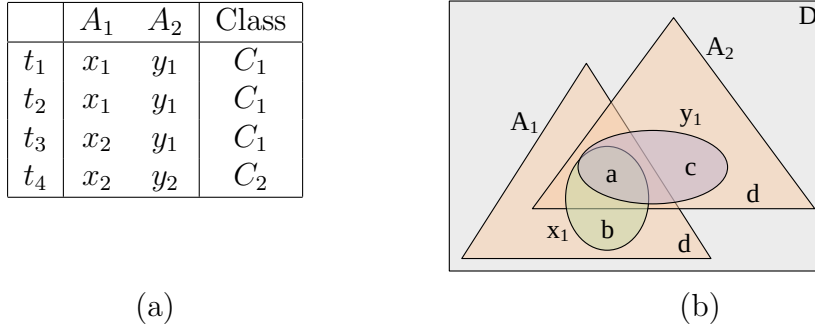


Figure 1: (a) A database example; (b) Venn diagram for the axiom $\langle \exists A_1.\{x_1\} \sqsubseteq \exists A_2.\{y_1\}, \alpha \rangle$ (obtained from a different database).

rules involving individuals (values) x_1 and y_1 belonging to the columns A_1 and A_2 , respectively (see Figure 1 (a)), and the degree α of fulfilment of this axiom. Figure 1 (b) shows that the higher the value of a , the higher the reliability of that relation between two individuals. Additionally, if x_1 would be completely included in y_1 we would have that b is near to zero. Summarizing these two ideas, when $\frac{a}{a+b}$ is near to 1 we will have more confidence in the satisfiability of the previous axiom. This value corresponds to the confidence-descriptive of the rule $x_1 \Rightarrow y_1$.

In more general cases, we can obtain axioms of the form $\langle \exists A_1.\{x_1\} \sqcap \exists A_2.\{x_2\} \sqcap \dots \sqcap \exists A_n.\{x_n\} \sqsubseteq \exists A_{n+1}.\{x_{n+1}\}, \alpha \rangle$. The difference here is that a, b, c , and d now involve several attributes in the antecedent. For example, b is now the number of tuples such that the value of A_1 is x_1 , the value of A_2 is x_2 , etc., but the value of A_{n+1} is not x_{n+1} .

The use of data mining to learn graded ontology axioms is not restricted to this single case. For example, one can consider association rules of the form $y_1 \Rightarrow C_1$ to produce graded ontology axioms of the form $\langle \exists A_2.\{y_1\} \sqsubseteq C_1, \beta \rangle$ or, in more general cases, $\langle \exists A_1.\{x_1\} \sqcap \exists A_2.\{x_2\} \sqcap \dots \sqcap \exists A_n.\{x_n\} \sqsubseteq C_1, \beta \rangle$.

As a final remark, it is worth to stress that the obtained axioms should be reviewed by a human expert before actually incorporating them to a graded ontology, because data mining techniques discover relationships that are only potentially interesting. In that regard, we are currently developing further evaluation tests of our approach.

References

- [1] M. Delgado, M. D. Ruiz, and D. Sánchez. Studying interest measures for association rules through a logical model. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 18(1):87–106, 2010.
- [2] Y. Kodratoff. Comparing machine learning and knowledge discovery in databases: An application to knowledge discovery in texts. In *Machine Learning and Its Applications, Advanced Lectures*, volume 2049 of *LNCS*, pages 1–21. Springer, 2001.
- [3] T. Lukasiewicz and U. Straccia. Managing uncertainty and vagueness in Description Logics for the Semantic Web. *Journal of Web Semantics*, 6(4):291–308, 2008.

Reglas de asociacion difusas en Big Data

Carlos Fernandez-Basso, M. Dolores Ruiz, Maria J. Martin-Bautista

**CITIC-UGR, University of Granada (Spain), Department of Computer Science and
A.I., University of Granada (Spain)**

karloos@correo.ugr.es, {mdruiz,mbautis}@decsai.ugr.es

La gran cantidad de información almacenada por las empresas y el aumento de las redes sociales y el Internet de las cosas están produciendo un crecimiento exponencial en la cantidad de datos que se producen. Por lo tanto, las técnicas de análisis de datos deben mejorarse para permitir que toda esta información pueda ser procesada. Una de las técnicas más utilizadas para la extracción de información en el campo de minería de datos es el de las reglas de asociación, que permite representar con precisión la frecuente co-ocurrencia de los elementos en un conjunto de datos. En particular reglas de asociación difusas han demostrado ser muy útiles en algunos dominios para hacer frente a la imprecisión en bases de datos y al mismo tiempo, ofrecen una buena representación del conocimiento encontrado.

El Big Data es un fenómeno que no para de crecer. Para las empresas y los usuarios es de gran interés poder analizar esta información y conseguir extraer de ella conocimiento con el fin de obtener beneficios, ya sean económicos o para la mejora y el interés de la sociedad en general.

Las técnicas de minería de datos, y en particular, las reglas de asociación [2] tienen como objetivo identificar conocimiento útil. Las reglas de asociación tradicionales se basan en describir las relaciones entre la presencia conjunta de elementos. Para ello, las reglas de asociación de tipo crisp, sólo nos dicen si un elemento se encuentra o no en una transacción. Sin embargo, la naturaleza de los datos puede ser diversa y puede venir descrita en forma numérica, categórica, imprecisa, etc. En el caso de elementos numéricos, una primera aproximación podría ser categorizarlos de forma que, por ejemplo, la altura de una persona pueda venir dada por un intervalo al que pertenece, por su valor o incluso por una etiqueta lingüística que además proporcione semántica. Por ejemplo, si la altura es 1'80 m, esta podría representarse por un intervalo ya que dicho valor se encuentra en el intervalo [1'70, 2], o bien por una etiqueta semántica que nos dice que la persona es "ALTA". Es por esto que necesitamos otras formas de representar este tipo de conocimiento que sean más intuitivas para el usuario. Una de las técnicas más utilizadas para representar el conocimiento impreciso es mediante la Teoría de Subconjuntos Difusos propuesta por L.A. Zadeh [5].

Vamos a continuación a definir cómo esta teoría se aplica a las reglas de asociación anteriormente explicadas. Por lo tanto primero vamos a definir que es una transacción difusa [1]. Sea I un conjunto de ítems, una transacción difusa t será un subconjunto difuso no vacío de I en el que el grado de pertenencia a t de un ítem $i \in I$, será un número en el intervalo $[0,1]$ y lo notaremos por $t(i)$. Para un ítemset, $A \subseteq I$ podremos calcular su grado de pertenencia a una transacción difusa t como sigue

$$t(A) = \min_{i \in A} t(i) \quad (1)$$

Según esta definición una transacción es un caso particular de transacción difusa. Nosotros representamos un conjunto de transacciones difusas por medio de una tabla donde las columnas identifican los ítems y las filas las transacciones. La celda del ítem i_k y la transacción t_j contiene un valor en el intervalo $[0,1]$ que sería el grado de pertenencia de i_k en t_j ($t_j(i_k)$).

Aunque se han propuesto varios métodos para la extracción de reglas de asociación difusas, estos métodos no funcionan bien en bases de datos con cantidades masivas de transacciones debido a los altos costes computacionales y los problemas de memoria. En este artículo abordamos estos problemas mediante el estudio de las tecnologías actuales para el procesamiento de grandes volúmenes de datos con el fin de proponer una paralelización del proceso de minería de reglas de asociación para bases de datos transaccionales difusas utilizando tecnologías Big Data. Además proponemos un algoritmo eficiente capaz de manejar grandes cantidades de datos y lo comparamos con algoritmos tradicionales de minería de reglas de asociación difusas.

El algoritmo utilizado para la extracción de los conjuntos frecuentes es el algoritmo Apriori [3]. Este algoritmo lo hemos adaptado para mejorar la eficiencia del algoritmo de extracción de reglas de asociación mediante el uso de tecnologías BigData creando además un algoritmo escalable. Entre las herramientas utilizadas podemos destacar Hadoop [4] para la distribución y gestión de los archivos en entornos distribuidos y Spark para el desarrollo del algoritmo en distribuido. Para todo ello se ha realizado un algoritmo que utiliza MapReduce con el propósito de poder realizar un procesamiento distribuido dentro de un cluster de servidores. Para ello el algoritmo se divide en dos fases, en la primera realizamos un recuento de el número de apariciones de cada ítem en cada α – corte mediante MapReduce y en la segunda fase se realiza el mismo recuento pero en este caso se busca cada uno de los itemset frecuentes calculados en la fase 1 o en una ejecución anterior de la fase 2.

Mediante los experimentos realizados con el algoritmo propuesto hemos comprobado que se mejora la eficiencia del algoritmo, en cuanto a tiempo y memoria, cuando el número de transacciones y ítems es muy elevado.

Agradecimientos

Este trabajo ha sido parcialmente sufragado por la Junta de Andalucía con el proyecto P11-TIC-7460 y la Unión Europea con el proyecto Energy IN TIME con Grant agreement no. 608981.

Referencias

- [1] M. Delgado et al. Fuzzy association rules: general model and applications. *Fuzzy Systems, IEEE Transactions on*, 11(2):214–225, 2003.
- [2] R. Agrawal et al. Mining associations between sets of items in large databases. In *ACM-SIGMOD Int. Conf. on Data*, pages 207–216, 1993.
- [3] R. Agrawal et al. Fast algorithms for mining association rules in large databases. In *Proc of the 20 Int. Conf. on Very Large Databases*, pages 487–499, Chile, 1994.
- [4] Tom White. *Hadoop: The Definitive Guide. Fourth edition*. O'Reilly, 2015.
- [5] L.A. Zadeh. Fuzzy sets. *Information and Control*, 8:338–353, 1965.

Open Linked Data Mining for Environmental Scanning*

Juan Gómez-Romero¹, Fernando Bobillo², M. Dolores Ruiz¹, Miguel Delgado¹

¹ *Dep. of Computer Science and Artificial Intelligence, CITIC-UGR, Univ. of Granada*
 {jgomez, mdruiz, mdelgado}@decsai.ugr.es

² *Aragon Institute of Engineering Research (I3A), Univ. of Zaragoza*
 fbobillo@unizar.es

Governmental institutions are increasing their efforts to offer bigger amounts of data as *open data*. By definition, open data can not only be freely accessed in electronic format in the Web, but also reused and redistributed for any purpose. Publishing open data has a two-fold impact to the society. First, it promotes transparency, when referred to public bodies' activities, which is essential in a democratic society. Second, it has social and commercial value, since it can be the enabler of further studies (e.g. independent academics can derive new knowledge) and commercial applications (e.g. socio-economic data can be used in private market analysis). Traditionally, publicly funded statistical offices, such as the Spanish Statistical Office and EUROSTAT, have been the most important open data providers. Nowadays, the definition of open data is widening, and arguably, other sources like social networks and open-licensed digital newspapers can be as well included.

Environmental scanning is defined as the process of acquisition and use of information about events, trends and relationships in an organization's external environment to support decision-making and planning [1]. The wide availability of open data has renewed the interest in this topic, in which Data Mining plays a key role regarding automatic discovery of non-trivial, new and valid knowledge from data. We identify here three challenges of open data mining for environmental scanning:

Data heterogeneity. Using different information sources is essential to provide a wider context to understand the whole scenario. Data combination implies incompatible formats, which must be expressed in a common language, and more importantly, semantic heterogeneity, which requires the alignment of diverse conceptual schemas based on information meaning rather than data format. This is particularly difficult if unstructured sources are to be considered; e.g. texts in natural language.

Less structured process. The typical data mining process encompasses four steps: data collection and cleaning (pre-processing), filtering (selection), exploration and model building (analysis), and visualization (for data description and prediction). To some extent, open data mining invert this process, because it implies the collection of data before determining an explicit purpose. Therefore, we may be more interested in exploratory analysis than model building, which would not require a thorough pre-processing. Conversely, it may happen that the pre-processing stage requires more effort before starting a detailed analysis aimed at decision-making.

Imprecision and vagueness. Open data makes it necessary to reinforce the capabilities of mining algorithms to discern between relevant and irrelevant information, and to be

* This work is partly funded by the Univ. of Granada (Programa de Proyectos para la incorporación de jóvenes doctores a nuevas líneas de investigación), the Spanish Ministry of Economy and Competitiveness (CICYT project TIN2013-46238-C4-4-R), and the Andalusian Government (P11-TIC-7460).

resilient to noise, uncertainty, and outliers. In this regard, Fuzzy Logic methods have proved to be more robust than conventional approaches [2]. In non-fuzzy methods, small changes in the inputs result in major changes in the outputs, whereas in fuzzy methods, changes are more gradual.

Semantic Web technologies, based on ontological and graph-based knowledge representations, provide an infrastructure for publishing, storing, retrieving, reusing, integrating, and analyzing open and linked data [3]. The standard Semantic Web Languages (RDF, RDFS, OWL) allow us to express data and metadata in the form of triples (subject, property, object) with formal properties, and more importantly, to easily link pieces of information. The Linked Data initiative purposely focus on connecting related information on the Web [4]. For example, DBPedia (a semi-structured version of Wikipedia) is linked to GeoNames (a geographical knowledge base). Specialized vocabularies based on RDF and OWL have been proposed to characterize, find and use open datasets: VoID, to express metadata about datasets; PROV-O, to define data provenance; Data Cube Vocabulary, for statistical data and metadata exchange; and so forth. Queries can be issued to distributed RDF sources, published in triplestore repositories, by using the SPARQL standard query language. All these technologies favor the creation of rich information contexts, with wider scope than the initial target data, at a low cost.

To illustrate the potential of open data mining, we have developed an example based on EUROSTAT, DBPedia and GeoNames open data. From EUROSTAT, we have identified a dataset about unemployment in European regions, encoded with the Data Cube vocabulary and linked to GeoNames and DBPedia, and we have retrieved unemployment rates for “all the European regions where German is an official language in 2000-2012”. Let us notice that the regions are defined in GeoNames, the official languages in DBPedia, and the unemployment rates in EUROSTAT. We have applied a fuzzy clustering algorithm to characterize groups of regions with similar evolution, which allows the use of imprecise cluster boundaries. The results suggest that the unemployment trends mostly depend on the initial rates and the changes in the first two years of the economical crisis, regardless of the country. We could incorporate other sources to the process, and refine the conditions to explore specific scenarios with small effort.

We consider that the use of ontological descriptions for linked data is one of the most promising directions for future work. Both the input data and the results –the clusters, in our example– could be semantically described with meaningful ontological expressions. Moreover, we could use a fuzzy ontology, in which concepts and relations can be imprecise. Another area of interest is the utilization of automatic ontology alignment algorithms in the pre-processing stage, in order to locate and link similar datasets.

References

- [1] C.W. Choo: *The art of scanning the environment*. Bulletin of the American Society for Information Science and Technology 25(3) (1999), 21-24.
- [2] E. Hüllermeier: *Fuzzy methods in machine learning and data mining: Status and prospects*. Fuzzy Sets and Systems 156(3) (2005), 387-406.
- [3] K. Janowicz, F. van Harmelen, J. Hendler, P. Hitzler: *Why the data train needs semantic rails*. AI Magazine 36 (1) (2014).
- [4] C. Bizer, T. Heath, T. Berners-Lee: *Linked Data - The story so far*. International Journal on Semantic Web and Information Systems 5(3) (2009), 1-22.

Análisis de Tópicos en Redes Sociales mediante Técnicas de Clustering Difuso

**Karel Gutierrez-Batista, Jesús R. Campaña,
M. Amparo Vila, María J. Martín-Bautista**

*Department of Computer Science and Artificial Intelligence, ETSIIT
University of Granada, 18071, Granada, Spain*

karel.gutierrez1984@gmail.com, [jesuscg|vila|mbautis]@decsai.ugr.es

Hoy día prolifera la acumulación en Internet y las organizaciones, de inmensas cantidades de datos textuales no estructurados en repositorios provenientes de documentos de trabajo, correos electrónicos, sitios web de opinión, encuestas, opiniones en redes sociales, etc. En estos datos se encuentran grandes cantidades de información muy relevantes. Sin embargo, el gran cúmulo y falta de estructura de éstos, hace que sea prácticamente imposible su procesamiento y análisis automático de forma masiva. En este entorno, resulta particularmente útil, detectar automáticamente los principales tópicos que son abordados y que constituyen información relevante. Desde el ámbito de la Minería de Textos, una solución propuesta para la detección de tópicos, es la aplicación de métodos clásicos de agrupamiento. Dichos algoritmos permiten agrupar un conjunto de textos y etiquetarlos con los principales tópicos tratados.

Si tenemos en cuenta que la gran mayoría de los textos de redes sociales están restringidos en cuanto a cantidad de caracteres y además un usuario puede realizar comentarios y expresar su opinión sobre cualquier tema, es de esperar que un simple documento pueda abordar varios temas, por lo que el principal problema que se pretende resolver es la detección de tópicos en redes sociales teniendo en cuenta que un mismo documento puede exponer varias temáticas. Motivado por la problemática anterior, una posible solución es el uso de algoritmos de agrupamiento difuso, los cuales permitirían identificar el grado de pertenencia de cada documento a cada cluster, y a su vez al tópico que representa dicho cluster.

En este trabajo proponemos un nuevo enfoque para detectar automáticamente los principales tópicos de datos textuales en redes sociales, utilizando técnicas de agrupamiento difuso y una base de conocimiento multilingüe, y de esta forma, facilitar el análisis de los textos en redes sociales. Nuestro acercamiento tiene la capacidad de agrupar los textos a partir de un conjunto de etiquetas pertenecientes a una taxonomía, las cuales sustituyen a los términos originales del texto tras ser procesados sintácticamente y de pasar por un proceso de desambiguación. Al sustituir las palabras por las etiquetas, se reduce en gran medida la dimensionalidad del problema, ya que independientemente del número de términos diferentes en los textos, se trabajará con el conjunto de etiquetas mencionado anteriormente, ya que con ellas se logra representar un gran número de los principales dominios del conocimiento. Las principales ventajas de nuestra propuesta y que lo diferencia de otras, es que permite analizar

The research reported in this paper was partially supported by the Andalusian Government (Junta de Andalucía) under projects P11-TIC-7460 and P10-TIC-6109.

los datos textuales de redes sociales en cualquier intervalo de tiempo, no es necesario conocer de antemano los temas tratados, y es independiente del idioma.

Para lograr esta tarea, se propone una metodología basada en el uso del recurso léxico Multilingual Central Repository 3.0 (MCR 3.0) [1], y de las herramientas Stanford Part-of-Speech (POS) Tagger [4] y Stanford Named Entity Recognition (NER) [2] combinados con algoritmos de agrupamiento difuso. El último paso, consiste en el etiquetado de los grupos obtenidos. Dicha etapa se encargará de seleccionar etiquetas descriptivas para cada grupo. Las etiquetas de los grupos pertenecen al conjunto de etiquetas definidas en el recurso semántico WordNet Domains [3] incluido en la base de conocimiento MCR 3.0. El uso de algoritmos de Agrupamiento Difuso brinda la posibilidad de conocer en qué medida se aborda cierto tópico en cada documento. Para demostrar la viabilidad de la propuesta, la experimentación se realizó con conjuntos de datos reales pertenecen a las redes sociales Twitter y DreamCatchers, en inglés y español respectivamente.

Con la aportación del presente trabajo, es posible manejar datos textuales integrados con información bien estructurada e independiente del idioma, especialmente en procesos de Minería de Datos y Data Warehousing, facilitando el desarrollo de herramientas robustas capaces de tratar datos heterogéneos e independientes del idioma. Esto le permitirá a los decisores realizar análisis con un alto nivel de detalle sobre el contenido de las redes sociales.

Referencias

- [1] Gonzalez-Agirre Aitor, Laparra Egoitz, and Rigau German. Multilingual central repository version 3.0. In *Proceedings of the Eight International Conference on Language Resources and Evaluation (LREC'12)*, Istanbul, Turkey, may 2012. European Language Resources Association (ELRA).
- [2] Jenny Rose Finkel, Trond Grenager, and Christopher Manning. Incorporating non-local information into information extraction systems by gibbs sampling. In *Proceedings of the 43rd Annual Meeting on Association for Computational Linguistics*, ACL '05, pages 363–370, Stroudsburg, PA, USA, 2005. Association for Computational Linguistics.
- [3] Bernardo Magnini and Gabriela Cavaglia. Integrating subject field codes into wordnet. In *LREC*. European Language Resources Association, 2000.
- [4] Kristina Toutanova, Dan Klein, Christopher D. Manning, and Yoram Singer. Feature-rich part-of-speech tagging with a cyclic dependency network. In *Proceedings of the 2003 Conference of the North American Chapter of the Association for Computational Linguistics on Human Language Technology - Volume 1*, NAACL '03, pages 173–180, Stroudsburg, PA, USA, 2003. Association for Computational Linguistics.

Fuzzy approximate dependencies in terms of representation by levels

C. Molina¹, M.D. Ruiz², D. Sánchez¹, J.M. Serrano²

¹ Dept. of Computer Science, University of Jaén (Spain)
 {jschica, carlosmo}@ujaen.es

² CITIC-UGR, Dept. Computer Science and A.I., University of Granada (Spain)
 {mdruiz, daniel}@decsai.ugr.es

Let D be a database over a relational schema $R = \{X_1, X_2, \dots, X_m\}$, and let X, Y be two disjoint, not empty, subsets of R . A functional dependency (FD) $X \rightarrow Y$ holds in D , iif

$$\forall t, s \in D : \text{if } t[X] = s[X] \text{ then } t[Y] = s[Y] \quad (1)$$

FDs still can be useful even when they do not totally hold. Works like [3][7][8], define approximate dependencies (AD), informally, as FDs with exceptions. In [5][10], ADs are defined in terms of association rules (AR), allowing the application of usual AR measures (support, confidence) on ADs. Furthermore, this same approach is followed in [1], together with the definition for fuzzy association rule (FAR) proposed in [4], to introduce fuzzy approximate dependencies (FAD). Restrictions imposed by equation 1 are relaxed. Accuracy of the dependency is measured by support and confidence; plus, a value may belong to a tuple with a certain membership degree; and finally a similarity, instead of equality, relation between values is defined. Later, [2] argues that similarity relations are excessively restrictive, and instead, resemblance relations are applied in FAR mining.

Representations by levels (RL) [9] has been proven as an valid alternative to fuzzy sets, and an effective approach to mine FARs has been proposed in [6]. Hence, it has sense to apply this same approach in the discovery of FADs. Let $\tilde{D}$ be a fuzzy database over R . For each attribute $X_i \in R$, a set of linguistic labels $\{E_1^i, \dots, E_n^i\}$ is defined. For each tuple $\tilde{t} \in \tilde{D}$, the intersection with X_i is a fuzzy set,

$$\tilde{t}[X_i] = \mu_{\tilde{t}}(E_1^i)/E_1^i + \dots + \mu_{\tilde{t}}(E_n^i)/E_n^i$$

where $\mu_{\tilde{t}}(E_1^i) \in [0, 1]$ is the membership degree of E_1^i in $\tilde{t}$, and so on.

Let $I_R = \{it_{X_i} | X_i \in R\}$ be the set of items associated with the attributes in R . Let $\tilde{T}$ be a set of fuzzy transactions, where for each $\tilde{t}, \tilde{s}$ in $\tilde{D} \times \tilde{D}$, we have $\tilde{t}\tilde{s} \in \tilde{T}$. The membership degree of an item it_{X_i} to the fuzzy transaction $\tilde{t}\tilde{s}$ is established by means of a resemblance relation R_i , defined as

$$\tilde{t}\tilde{s}[it_{X_i}] = R_i(\tilde{t}[X_i], \tilde{s}[X_i]) = \begin{cases} 1, & \text{if } \tilde{t} = \tilde{s} \\ \max_{j=1 \dots n} \min(\mu_{\tilde{t}}(E_j^i), \mu_{\tilde{s}}(E_j^i)), & \text{otherwise} \end{cases} \quad (2)$$

The present contribution has been partially supported by the Andalusian Government (Junta de Andalucía) under project P11-TIC-7460.

Once we have the set of fuzzy transactions $\tilde{T}$, mining FADs as FARs by RL is straightforward. Besides, since this approach ensures that crisp operations and measures can be easily extended to RLs, while maintaining all the properties of the crisp case, a possible scenario appears when the user is only interested in crisp ADs, but they can only be found at a given representation level. Crisp ADs can be mined at every RL, and next, without any loss of generality, aggregated into a global FAD. In this work, we aim to study additional properties for FADs by RLs as well as analyze extended application scenarios.

References

- [1] F. Berzal, I. Blanco, D. Sánchez, J.M. Serrano, and M.A. Vila. A definition for fuzzy approximate dependencies. *Fuzzy Sets and Systems*, 149(1):105 – 129, 2005. Fuzzy Sets in Knowledge Discovery.
- [2] F. Berzal, J. C. Cubero, D. Sánchez, M. A. Vila, J. Calero, G. Delgado, and J.M. Serrano. A fuzzy data mining application over user knowledge involving fuzzy resemblance relations. In *Proc Alternative Techniques in Data Mining and Knowledge Discovery Workshop/4th Int Conf on Data Mining (ICDM2004)*, 2004.
- [3] P. Bra and J. Paredaens. *Advances in Data Base Theory: Volume 2*, chapter Horizontal Decompositions for Handling Exceptions to Functional Dependencies, pages 123–141. Springer US, Boston, MA, 1984.
- [4] M. Delgado, N. Marín, D. Sánchez, and M.A. Vila. Fuzzy association rules: General model and applications. *IEEE Transactions on Fuzzy Systems*, 11(2):214–225, 2003.
- [5] M. Delgado, M.J. Martín-Bautista, D. Sánchez, and M.A. Vila. Mining strong approximate dependencies from relational databases. In *Proc. IPMU'2000*, volume 2, pages 1123–1130, 2000.
- [6] M. Delgado, M.D. Ruiz, D. Sánchez, and J.M. Serrano. A formal model for mining fuzzy rules using the RL representation theory. *Information Sciences*, 181(23):5194–5213, 2011.
- [7] J. Kivinen and H. Mannila. Approximate inference of functional dependencies from relations. *Theoretical Computer Science*, 149(1):129 – 149, 1995. Fourth International Conference on Database Theory (ICDT '92).
- [8] B. Pfahringer and S. Kramer. Compression-based evaluation of partial determinations. In *KDD*, pages 234–239, 1995.
- [9] D. Sánchez, M. Delgado, M.A. Vila, and J. Chamorro-Martínez. On a non-nested level-based representation of fuzziness. *Fuzzy Sets and Systems*, 192:159–175, 2012.
- [10] D. Sánchez, J.M. Serrano, I. Blanco, and M.J. Martín-Bautista. Using association rules to mine for strong approximate dependencies. *Data Mining & Knowledge Discovery*, 16(3):313–348, 2008.

Sistema Híbrido AHP-Fuzzy Logic, para Toma de Decisiones en Entornos de Incertidumbre. Aplicación a FMA

Luis Santacreu Rios¹, Alejandro Talavera Ortiz¹, Ricardo Aguasca Colomo¹, Maximo Mendez Babey¹, Blas Galvan Gonzalez¹, Dagoberto Castellano Nieves²

¹ *Instituto Universitario SIANI-CEANI, Universidad de Las Palmas de Gran Canaria
3500, Las Palmas de Gran Canaria, Las Palmas. España
ricardo.aguasca@ulpgc.es*

² *Grupo de Computación Inteligente, Universidad de La Laguna
38207, San Cristóbal de La Laguna, Santa Cruz de Tenerife. España
dcastell@ull.es*

Resumen

Las Comunidades Autónomas en España tratan los Fenómenos Meteorológicos Adversos (FMA) según las especificidades de sus Planes de Protección Civil (instrumentos que utilizan para atender las emergencias) y salvo las que tienen agencia propia para la predicción meteorológica (Agencia Vasca de Meteorología), todas trabajan con la información facilitada por la Agencia Estatal de Meteorología (AEMET) y con conjuntos de expertos. Esto hace que todas traten los FMA con bastantes similitudes y que los problemas que surgen tengan múltiples coincidencias.

Introducción

En general, la toma de decisiones ante un FMA, sigue las siguientes fases:

1. La información facilitada por la AEMET ante la detección de un FMA, recibida en el 1-1-2 Canarias, se pone en conocimiento de la Dirección General de Seguridad y Emergencias de la Comunidad Autónoma (DGSE), como órgano competente en las declaraciones de prealerta, alerta o alerta máxima.
2. En base al análisis de la información recibida y en aplicación al Plan Específico de Emergencias de la Comunidad por Riesgos de Fenómenos Meteorológicos Adversos (PEFMA), la DGSE realiza la declaración de prealerta, alerta o alerta máxima. En el tratamiento de la información influye la cantidad recibida, como ha sido transmitida a través de los diferentes canales y si ha sido influenciada en el proceso, la experiencia previa de los decisores, el tiempo disponible para tomar la decisión, etc.
3. La declaración es notificada al 1-1-2 que realiza los comunicados a las administraciones, organismos, servicios operativos y empresas involucradas en la operatividad del PEFMA, además de los avisos a la población en caso necesario. Cuando se recibe información adicional o datos reales, cuando el FMA está en curso, puede cambiar la declaración y se repite el proceso desde la fase 1.

La toma de decisiones ante los FMA se enfrenta a varias problemáticas: a) El proceso está influenciado por el conjunto de expertos involucrados en la toma de decisiones, que no siempre es el mismo, al tratarse de un servicio que funciona las 24 horas al día, los 365 días al año; b) Se trabaja con imprecisión, ambigüedad, datos erróneos o ausencia de estos y expresiones de los seres humanos. Además, en meteorología se trabaja con conceptos imprecisos y las predicciones meteorológicas se realizan utilizando modelos matemáticos que no siempre se cumplen. Por otro lado, la singularidad de la geografía, lo que se conoce

como efectos locales, no la tienen en cuenta los modelos utilizados pero sin embargo si se encuentra en el conocimiento de los expertos; c) Se maneja mucha información y se toman múltiples decisiones, que afectan a la vida y los bienes de las personas, en un tiempo limitado. La información debe ser tratada en el menor tiempo posible para generar las prealertas, alertas y alertas máximas, así los como comunicados y avisos a la población.

Sistema Experto para Tomar Decisiones Basado en Lógica Difusa y AHP

Las soluciones propuestas a los problemas planteados son:

1. Seleccionar un conjunto de expertos, especialistas en seguridad y emergencias con un alto grado de conocimiento en el PEFMA y FMA. Para la selección se aplica el “Coeficiente de Competencia Experta” o “Coeficiente K” [1].
2. Utilizar el método AHP, propuesto por Saaty [2], para recoger el conocimiento de los expertos, dando consistencia a los juicios a la vez de soporte matemático AHP se puede combinar con lenguaje Fuzzy [3] por lo que es ideal para trabajar con la Lógica Difusa.
3. Mediante La utilización de Lógica Difusa, construimos funciones de pertenencia para poder evaluar las alternativas posibles en cada momento (declaraciones de prealerta, alerta o alerta máxima). El grado de pertenencia a determinada función es lo que el sistema utilizará para proponer la alternativa idónea.
4. Actualizar la tradicional metodología AHP Fuzzy [3] para simplificar los cálculos con los conjuntos difusos y desarrollar un sistema híbrido AHP-Fuzzy.
5. Gestionar la gran cantidad de información en un corto lapso de tiempo, automatizando todo el proceso mediante un sistema experto basado en reglas [4], en el que la base de conocimiento ha sido generada mediante AHP y el motor de inferencia se basa en Lógica Difusa [5].

Conclusiones

Mediante la aplicación del sistema híbrido para la toma de decisiones se ha conseguido sistematizar, por un lado, la extracción de conocimiento para la generación de una base de reglas coherente, y por otro se ha generalizado un método para la construcción de las funciones de pertenencia utilizadas por el motor de inferencia difuso con buenos resultados, contrastados con el histórico de datos proporcionados por CECOES 1-1-2 Canarias.

Referencias

- [1] Cabero-Almenara J, Barrosos-Osuna J. *La Utilización del Juicio de Experto para la Evaluación de TIC: el Coeficiente de Competencia Experta*. Bordón, **65**, 2, (2013) 25-38.
- [2] Saaty TL. *Fundamentals of Decision Making and Priority Theory with the Analytic Hierarchy Process*. 2a edition. Pittsburgh: RWS Publications, (2000).
- [3] Jie Lu, Guangquan Zhang, Da Ruan, Fengjie Wu. *Multi-Objective Group Decision Making. Methods, Software and Applications with Fuzzy Set Techniques*. 1ª edition. London: Imperial College Press, (2007).
- [4] Calvo-Rolle JL, Casteleiro-Roca JL, Vilar-Martinez XM et al. *Expert System Development to Assist on the Verification of “Tacan” System Performance*. DYNA. **89**, 1, January (2014) 112-121.
- [5] García-Fernández N, Puente J, Fernández I et al. *How to Improve the Suppliers Evaluation Process Using Fuzzy Inference Systems*. DYNA. **89**, 4, July (2014) 449-456.

The Team Orienteering Problem with Fuzzy Times

Julio Brito, Airam Expósito and José A. Moreno

¹ *Department of Computer Engineering and Systems, University of La Laguna, Spain*
 {jbrito, aexposim, jamoreno}@ull.edu.es

1 Introduction

Route Planning Problems constitutes one of the most important operational transport problems in many industrial environments. In this context, the present work is focused on the resolution of special cases of vehicle routing problems. A set of potential customers is considered and a fleet of vehicle is used to visit a subset of customers, there is no constraint to visit all customers. These problems are named The Orienteering Problems (OP) [5] and are inspired in outdoor competitive games. In these games, you have to carry out routes through a number of locations where certain scores are obtained and the winner is whoever gets the highest score. The main objective is to determine a path that consists of a subset of nodes and design a route with selected nodes so that the total collected score is maximized, and not exceeding the predefined maximum travel time. In particular, this paper focuses in The Team Orienteering Problem (TOP) that was first introduced by Chao et al. [2]. The Team Orienteering Problem (TOP) is a variant of the Orienteering Problem (OP) that only considers one route whereas TOP provides the ability to plan a certain number of simultaneous tours.

2 Fuzzy TOP

In many routing planning problems of the real world the available information about the problem is often imprecise, vague or uncertain. For instance, travel times depend on the surrounding conditions and the situation of traffic, roads or weather. The available information on these conditions is often sparse and imprecise, and not easily accessible. In addition, decision makers take decisions with some subjectivity and a certain degree of flexibility. Thus it has been thought necessary to propose models and methodologies that support new applications in the route planning, incorporating imprecision and flexibility. Soft Computing includes an appropriate family of models and methods to deal with this characteristics. This work considers a version of TOP where the travel times are considered fuzzy. To solve this TOP variant we propose a methodological approach that uses metaheuristics in a fuzzy approach for optimization problems. In this paper we use the ideas introduced by Bellman and Zadeh (1970) [1] for fuzzy optimization problems and the methods developed by Verdegay and others [7].

This work has been partially funded by the Spanish Ministry of Economy and Competitiveness (project TIN2015-70226-R). Airam Expósito-Márquez would like to thank the Canary Government for the financial support they receive through their post-graduate grants.

3 Solution Approach

This approach provide Fuzzy Linear Programming (FLP) formulations of the problems with a number of methods for solving them, in a direct and easy way, obtaining solutions that are coherent with their fuzzy nature. We formulated the TOP as a linear programming problem with fuzzy inequalities in some constraints. Different FLP models can be considered according to the elements that contain imprecise information that are used as a basis for the classification proposed in [6]. In this context, heuristics and metaheuristics constitute a good referent for efficiently providing good solutions and strategies that are integrated with other Soft-Computing tools to facilitate approximate solutions to more complex real world problems. This paper proposes a solution algorithm based on Greedy Randomized Adaptive Search Procedures (GRASP) [3]. GRASP is a relevant metaheuristics that have shown appropriated for solving real and large instances of routing problems. This multistart two-phase metaheuristic for combinatorial optimization problems basically consists of a construction phase and a local search improvement phase. A feasible solution is obtained in the construction phase. The solution construction mechanism builds a solution step-by-step by adding a random new element from a candidate list to the current partial solution under construction without destroying feasibility. Subsequently the neighborhood of the solution is explored until a local minimum is found in the local search phase. GRASP has been successfully applied to a wide variety of combinatorial optimization problems including route planning [4].

References

- [1] R.E. Bellman and L.A. Zadeh. Decision making in a fuzzy environment. *Management Science*, 17(4):141–164, 1970.
- [2] I-Ming Chao, Bruce L. Golden, and Edward A. Wasil. The team orienteering problem. *European Journal of Operational Research*, 88(3):464 – 474, 1996.
- [3] T. A. Feo and M. G. C. Resende. Greedy randomized adaptive search procedures. *Journals Global Optimization*, 6:109–133, 1995.
- [4] Mauricio G.C. Resende and Celso C. Ribeiro. Greedy randomized adaptive search procedures: Advances, hybridizations, and applications. In Michel Gendreau and Jean-Yves Potvin, editors, *Handbook of Metaheuristics*, volume 146 of *International Series in Operations Research and Management Science*, pages 283–319. Springer US, 2010.
- [5] T. Tsiligrirides. Heuristic methods applied to orienteering. *The Journal of the Operational Research Society*, 35(9):797–809, 1984.
- [6] J. L. Verdegay. Fuzzy optimization: models, methods and perspectives. In *In proceeding 6th IFSA-95 World Congress.*, pages 39–71, 1995.
- [7] J.L. Verdegay. *Fuzzy Information and Decision Processes*, chapter Fuzzy mathematical programming. North-Holland, 1982.

FLMICMAC aplicado a las moratorias turísticas en las Islas Canarias

D. Castellanos Nieves,¹ C. Cruz Corona,² J. Brito Santana¹

¹ *Grupo de Computación Inteligente*

Universidad de La Laguna

38207, San Cristóbal de La Laguna, Santa Cruz de Tenerife. España

dcastell@ull.es, jbrito@ull.es

² *Grupo de Investigación en Modelos de Decisión y Optimización*

Universidad de Granada

E-18071, Granada. España

carloscruz@decsai.ugr.es

Resumen

Las organizaciones logran ventajas competitivas identificando de forma precisa los escenarios a los que se enfrentan. Identificando los posibles futuros se facilita la anticipación para afrontar los cambios en el entorno y su seguimiento. Las organizaciones se apoyan en los sistemas que realizan análisis prospectivos. En este trabajo se presenta la aplicación de la metodología FLMICMAC a una problemática en el dominio del turismo. La metodología FLMICMAC mejora la información en los procesos de toma de decisiones estratégicas, al identificar áreas de actuación prioritarias, ante posibles escenarios de moratorias turísticas.

Una moratoria es una regulación gubernamental que limita, de forma temporal, la expansión de la capacidad productiva de un sector. En el sector turístico, las moratorias han sido utilizadas para favorecer el rejuvenecimiento de destinos maduros. Las moratorias pueden tener efectos contraproducentes, lo cual depende en gran medida de si el instrumento que se utiliza se ajusta a los problemas a solucionar [1]. El impacto y la complejidad del problema de la moratoria turística y su dependencia sobre la sostenibilidad conducen a la necesidad de estudiar y evaluar futuros hipotéticos. Una de las metodologías más utilizadas para el análisis prospectivo es el método de escenarios [2]. Estos escenarios suelen elaborarse en base a información de diferentes expertos que se obtiene a través de encuestas, paneles, etc. Debido al carácter impreciso de mucha de esta información así como a la incertidumbre inherente en los estudios prospectivos, las técnicas utilizadas en el método de escenarios presentan en muchas ocasiones ciertas carencias a la hora de modelar. En este sentido, el uso de técnicas de Soft Computing es de gran interés. La Soft Computing puede verse como una familia de métodos que permiten el manejo de la imprecisión y la incertidumbre. En este trabajo nos hemos centrado en el método de escenarios propuesto por Godet [2], y las mejoras introducidas al análisis estructural por la metodología FLMICMAC [3][4].

El impacto de las moratorias sobre los alojamientos turísticos no ha sido suficientemente estudiado [1]. Los métodos de análisis de escenarios pueden contribuir a la evaluación y contraste de propuestas y sus dependencias con la situación actual. En la metodología FLMICMAC se propone el uso de variables lingüísticas en el modelado de las interrelaciones de los sistemas [3]. Esto permite a los expertos valorar el grado de las

interrelaciones con valoraciones lingüísticas, en lugar de valoraciones numéricas exactas [2]. Además de esto, proporciona unos resultados con información cercana y fácil de interpretar. De esta manera se consigue obtener modelos más robustos y reales que proveen una mejor caracterización de los escenarios y por tanto de los futuros posibles.

Aplicando la metodología FLMICMAC en el dominio de las moratorias turísticas en Canarias: 1) Se determinó el conjunto de etiquetas lingüísticas definidas mediante números difusos triangulares. Este conjunto de etiquetas lingüísticas permiten establecer el grado de influencia y dependencia entre las variables; 2) Se empleó un operador de agregación, para establecer la etiqueta lingüística, que corresponde a cada variable perteneciente al modelo definido. En el caso de las relaciones directas, se emplea el operador suma, y en el caso de las relaciones indirectas el operador potencia; 3) Obtención del ranking de las variables, en las relaciones directas e indirectas, de una manera análoga a la propuesta por Godet [2]. Esto es posible ordenando los resultados con el empleo de los operadores de agregación, lo que permite que los expertos puedan evaluar la influencia entre las variables del sistema en términos lingüísticos (alto, medio-alto,..., etc). El procedimiento facilita el consenso entre las distintas opiniones de los expertos en el ámbito de las moratorias turísticas. Esta información, junto con las etiquetas lingüísticas asociadas a cada variable, permite analizar los resultados de forma absoluta y relativa en el dominio [5].

Como conclusiones se tienen que la metodología FLMICMAC es aplicada al análisis estructural (método de los escenarios) en las moratorias turísticas en las Islas Canarias. El uso de la lógica difusa y las variables lingüísticas, permite un mejor modelado de las opiniones de los expertos; así como proporciona una solución más sólida, fiable e interpretable. Lo cual además, facilita el proceso de decisión y la creación de modelos cualitativos.

Referencias

- [1] Hernández-Martín, R., Álvarez-Albelo, C. D., Padrón-Fumero, N., *The economics and implications of moratoria on tourism accommodation development as a rejuvenation tool in mature tourism destinations*, Journal of Sustainable Tourism, 23, 6 (2015) 881-899.
- [2] Godet M.: *How to be rigorous with scenario planning*. Foresight, 2 (1) (2000) 5-9.
- [3] Villacorta P. J., Masegosa A. D., Castellanos D., Lamata M. T, *A new fuzzy linguistic approach to qualitative Cross Impact Analysis*, Applied Soft Computing, 24 (2014) 19-30.
- [4] Villacorta, Pablo J; Masegosa, Antonio D; Castellanos, D.; Lamata, M. T.: *A linguistic approach to structural analysis in prospective studies*. Advances on Computational Intelligence, (2012) 150-159.
- [5] Villacorta, P.J., Masegosa, A. D., Castellanos, D., Novoa, P., Pelta, D.: *Sensitivity analysis in the scenario method: a multi-objective approach*. Intelligent Systems Design and Applications., (2011) 867-872.

Solving a Fuzzy Multi-Criteria Location Problem by Using Georeferenced Data

Christopher Expósito-Izquierdo, Airam Expósito-Márquez, Belén Melián-Batista, José Marcos Moreno-Vega

¹ *Department of Computer Engineering and Systems, University of La Laguna, Spain*
 {cexposit, aexposim, mbmelian, jmmoreno}@ull.edu.es

Location theory seeks to determine the best possible sites out of a finite set of eligible candidates to set up a given number of new facilities (*e.g.*, hospitals, factories, fire stations, schools, delivery centres, etc.). In this context, determining the suitability of a given site to set up a new facility is usually seen as a problematic aspect due to the fact that, in most of cases, it is subject to multiple conflicting goals [4]. These are associated with the individual interests of the stakeholders found in the application context under analysis (*e.g.*, business leaders, politicians, potential customers, financial investors, etc.). On the basis of the existence of many conflicting objectives, the practical models in this field are intrinsically tackled as optimization problems, where some known objective function must be minimized or maximized while satisfying some constraints derived from the context and which represent conditions imposed on the underlying variables. Some highlighted examples of location problems can be found in [3].

The criteria associated with the optimization problems included into the field of location theory are usually subject to some degree of vagueness and uncertainty in practice. The way through which the decision makers express their preferences can be considered as the main reason behind this fact. Particularly, in most of the cases, a given decision maker provides his expert knowledge concerning some spatial problem by means of imprecise terms, such as *near*, *not too high*, or *maybe* [1]. Additionally, the potential sites to set up a given facility do not usually verify a given constraint strictly in real-life contexts of location decisions. Traditionally, the optimization problems have been addressed by using Boolean logic so far. However, *fuzzy logic* [5] arises as a more useful alternative because it is closer to human reasoning and it is more convenient when addressing non-linearities and uncertainties.

In the present work, an illustrative fuzzy multi-objective optimization problem belonging to the field of location theory is proposed. Specifically, its optimization objective is to set up a pre-specified number of new logistic facilities on the island of Tenerife (Spain). However, all the locations are not feasible to set up a logistic facility. Instead, these locations are subject to several imprecise criteria. These criteria concern transportation, economic, and sustainability aspects of the application context at hand.

The main goal of this work is to assess the applicability of different fuzzy membership functions to satisfy the imprecise criteria proposed by the multi-objective optimization problem. The problem is here studied in a practical environment by using open georeferenced

This work has been partially funded by the Spanish Ministry of Economy and Competitiveness (project TIN2012-32608). Christopher Expósito-Izquierdo and Airam Expósito-Márquez would like to thank the Canary Government for the financial support they receive through their post-graduate grants

data extracted from the OpenStreetMap project* and handled by a Geographic Information System (GIS). The relevance of GISs in optimization problems belonging to the field of location theory arises from the fact that this information technology aids to facilitate problem understanding and decision-making process. Due to the fact that the criteria imposed by the problem are imprecise in practice, the set of points on the map is here handle as fuzzy. The fuzzy sets are used in multi-objective optimization problems with the aim of standardizing the existing criteria. This is done by assigning to each point provided by the GIS a degree of membership for each criterion. A fuzzy set is denoted as $(\tilde{L}, \mu_{\tilde{L}})$, where $\tilde{L}$ is the set of points provided by the GIS from a fuzzy perspective and $\mu_{\tilde{L}}$ is a function with the form $\mu_{\tilde{L}} : L \rightarrow [0..1]$. $\mu_{\tilde{L}}(i)$ quantifies the strength of membership of the point $i \in L$ in the fuzzy set $\tilde{L}$. Here, a value $\mu_{\tilde{L}}(\cdot) = 1$ represents a full membership in the set, $\mu_{\tilde{L}}(\cdot) = 0$ indicates not membership to the set, whereas intermediate values (i.e., $0 < \mu_{\tilde{L}}(\cdot) < 1$) represent partial membership to the set.

The degrees of membership associated with the available points in the GIS are combined appropriately to report feasible solutions for the multi-objective optimization problem under analysis. For the shake of simplicity, the points under the different fuzzy membership functions are in this work combined by using the weighted linear combination [2]. In this regard, the different criteria are weighted to represent their individual importance in the problem. The values of the weights are appropriately set by the decision maker on the basis of his preferences.

References

- [1] Mahdi Bashiri and Seyed Javad Hosseini-zhad. A fuzzy group decision support system for multifacility location problems. *The International Journal of Advanced Manufacturing Technology*, 42(5-6):533–543, 2009.
- [2] Ni-Bin Chang, G. Parvathinathan, and Jeff B. Breeden. Combining GIS with fuzzy multicriteria decision-making for landfill siting in a fast-growing urban region. *Journal of Environmental Management*, 87(1):139 – 153, 2008.
- [3] H. A. Eiselt and Vladimir Marianov. *Applications of Location Analysis*, volume 232 of *International Series in Operations Research & Management Science*. Springer International Publishing, Berlin, 2015.
- [4] Jun Gang, Yan Tu, Benjamin Lev, Jiuping Xu, Wenjing Shen, and Liming Yao. A multi-objective bi-level location planning problem for stone industrial parks. *Computers & Operations Research*, 56:8 – 21, 2015.
- [5] L.A. Zadeh. Fuzzy sets. *Information and Control*, 8(3):338 – 353, 1965.

*<http://www.openstreetmap.org/>

Método multicriterio basado en Soft Computing para evaluar la calidad de software

Yamilis Fernández Pérez¹, Carlos Cruz Corona², José Luis Verdegay Galdeano²

¹ *Universidad de las Ciencias Informáticas*
yamilisf@gmail.com, yamilisf@uci.cu

² *Universidad de Granada, España*
carloscruz@decsai.ugr.es, verdegay@decsai.ugr.es

Hoy día existe una creciente preocupación por lograr que los productos de software cumplan con ciertos criterios de calidad. Para ello, se avanza en la definición e implementación de estándares que fijan los atributos deseables del software de calidad, a la vez que se desarrollan modelos y metodologías para la evaluación de la calidad. La información almacenada durante este proceso es heterogénea, ambigua, imprecisa, y de diferentes fuentes. Todo lo anterior dificulta el proceso de agregación para obtener un índice de calidad. Para ello se han propuestos varios métodos de decisión multi-criterios basados en Soft Computing, tales como: Fuzzy TOPSIS [4], Fuzzy AHP[2]. Hay propuestas basadas en Fuzzy ANP[3] para tratar la interdependencia entre criterios; pero su aplicación práctica es engorrosa, difícil y poco intuitiva.

Todo esto nos lleva a presentar en este trabajo un nuevo método que intenta solventar algunas de estas insuficiencias y que además permite la agregación parcial y total de los criterios. Esta propuesta tiene una estructura jerárquica de cuatro niveles: en el nivel 0 se determina un índice de calidad, en el nivel 1 se evalúan las características, en el nivel 2 las sub características y en el nivel 3 las medidas. En cada nivel se tiene en cuenta la interrelación entre los miembros.

El método consta de 5 pasos:

1. Obtención del modelo de calidad, las medidas y los criterios esenciales: A partir de los requisitos de evaluación se determina el modelo de calidad y las medidas a usar. Se analizan los criterios esenciales, que son los atributos que el producto debe satisfacer y cuya ausencia debe ser penalizada de manera que todo el sub-árbol se evalúa a cero. Para cada criterio se fija un valor de umbral (α). Si el valor de la medida es menor o igual del umbral entonces se penaliza el producto con valor 0 para ese sub-árbol.
2. Determinación del peso de los criterios y la interrelación: Se utiliza la comparación a pares como propone AHP sin tener en cuenta las interdependencias, y se deriva el vector peso por el método de eigenvector(W^l). Luego se determina la relación que hay entre hermanos en cada nivel con el uso de los Mapas Cognitivos Difusos[5] y se obtiene la matriz de dependencia entre los indicadores de nivel 1 (MDI^1 de orden $ln \times ln$).

Este trabajo está soportado por los proyectos TIN2014-55024-P del Ministerio de Economía y Competitividad, y P11-TIC-8001 de la Junta de Andalucía (en todos los casos incluyendo fondos FEDER de la Unión Europea).

3. Evaluación de los productos de software: Se realizan los testing para los m softwares a evaluar, se obtiene el valor de cada una de las ln medidas seleccionadas y se conforma la matriz de decisión Md^l de orden $m \times ln$.
4. Agregación de la información: Primero se calcula la matriz de influencia de antecedentes ($B^l = Md^l \times MDI^l$). A continuación, la matriz ponderada (D^l),

$$D^l = (d_{i,j}) \in \mathbf{R}_{m \times ln} \quad d_{i,j} = w_{i,j} \times \left((x_{i,j} + b_{i,j}) / \left(\sum_{i=1}^m (x_{i,j}) + \sum_{i=1}^m (b_{i,j}) \right) \right) \quad (1)$$

Luego, se obtiene la matriz de decisión del nivel $l-1$ (Md^{l-1}),

$$Md^{l-1} = (x_{i,j}^{l-1}) \quad (x_{i,j}^{l-1}) = \sum_{k=1}^{\omega(j)} (d_{i,j}) \quad (2)$$

Por último, se penalizan los productos teniendo en cuenta los criterios esenciales, y usando el método del vector de exclusión descrito en [1]. Se repite el paso 4 desde el nivel $l-1$ hasta el nivel 0, donde se obtiene el índice de calidad ($IQ \in [0, 1]$) para cada alternativa.

5. Obtención de la Recomendación: El índice de calidad es ordenado de mayor a menor por lo que se obtiene una lista ordenada de productos de software .

Para validar la propuesta se tomaron 12 productos reales evaluados bajo un modelo de calidad basado en la ISO 25010. Se calculó el índice global de calidad usando otros métodos multi-criterios y nuestra propuesta. Se compararon los resultados y se constató la viabilidad del método propuesto.

Referencias

- [1] E. H. Cables Pérez and M. T. Lamata Jiménez. *Tesis Doctoral: Selección de personal con técnicas de Soft Computing. Propuesta de Desarrollo y de Software*. PhD thesis, Universidad de Granada, 2011.
- [2] S. K. Dubey, A. Mittal, and A. Rana. Measurement of Object Oriented Software Usability using Fuzzy AHP. *International Journal of Computer Science and Telecommunications*, 3(5):98–104, 2012.
- [3] M. L. Etaati, S. Sadi-Nezhad, A. Using Fuzzy Analytical Network Process and ISO 9126 Quality Model in Software Selection: A case study in E-learning Systems. *Journal of Applied Sciences*, 11(1):96–103, 2011.
- [4] G. A. Montazer and H. Q. Saremi. An Application of Type-2 Fuzzy Notions in Website Structures Selection : Utilizing Extended TOPSIS Method. *WSEAS TRANSACTIONS on COMPUTERS*, 7(1):8–15, 2008.
- [5] E. I. Papageorgiou and J. L. Salmeron. A review of fuzzy cognitive maps research during the last decade. *IEEE Transactions on Fuzzy Systems*, 21(1):66–79, 2013.

Análisis de patrones en predicción de series temporales fuzzy

A. Rubio¹, E. Vercher², J.D. Bermúdez³

^{1,2,3} *Departamento de Estadística e IO (Universidad de Valencia)*

Arufor@alumni.uv.es, Enriqueta.Vercher@uv.es,

Jose.D.Bermudez@uv.es

El interés de nuestro trabajo se centra en el análisis de series temporales de índices bursátiles (Figura 1), de las que se pretende realizar predicciones de su comportamiento futuro. Este problema ha sido habitualmente analizado mediante la aplicación de la metodología de series temporales fuzzy, introducida por Song y Chissom (1993), y el procedimiento de resolución establecido por Chen (1996).

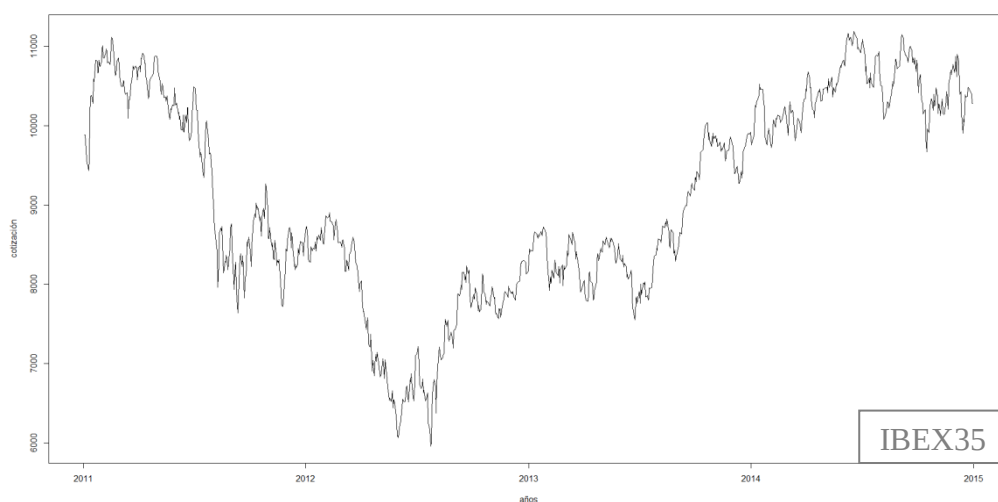


Figura 1: Valores de cotización del índice IBEX35

Siguiendo la estrategia que caracteriza el análisis de series temporales desde un enfoque fuzzy, nuestra propuesta puede considerarse un método ponderado de series temporales fuzzy para la predicción del valor futuro de un índice bursátil de referencia (Yu 2005, Cheng et al 2008). Para mejorar la precisión y robustez de la predicción hemos introducido algunas modificaciones, que afectan tanto al esquema de relaciones de las variables lingüísticas de los modelos clásicos de series temporales fuzzy como a las estrategias de ponderación, que estarían basadas en el análisis de patrones.

Hemos desarrollado un procedimiento de predicción que en líneas generales sigue el esquema básico descrito por Chen (1996), aunque presenta variaciones en diferentes etapas del procedimiento: en la partición del universo (Etapa 1), en la definición de los conjuntos fuzzy (Etapa 2), y en la construcción del operador de ponderación (Etapa 4).

Las modificaciones introducidas para la composición y asignación de los pesos obvian el uso de las relaciones fuzzy de grupo, y afectan al proceso de predicción fuzzy, proporcionando como predicción del comportamiento futuro del índice de referencia, un

número fuzzy trapezoidal (Liu 2007). La obtención de dicho número trapezoidal fuzzy permite calcular el valor y la ambigüedad del comportamiento futuro del índice, así como algunos momentos posibilísticos a partir del concepto de intervalo de medias (Dubois y Prade 1987).

Para componer nuestro operador de ponderación definimos los 'saltos', una característica intrínseca de la serie temporal fuzzy que intenta captar la tendencia de la serie, al tiempo que persigue ligar la ponderación de las relaciones fuzzy a toda la serie, en lugar de a las relaciones fuzzy de grupo. Los saltos son la manifestación de los movimientos que infiere la lógica fuzzy implícita en las relaciones fuzzy entre los conjuntos fuzzy definidos. Nuestro operador está basado en la secuencia de 'saltos' asociados a las relaciones lógicas fuzzy. A partir del análisis de dichas secuencias, se buscan patrones mediante la identificación de recurrencias en las secuencias de saltos, y finalmente se compone el operador.

Comparamos las predicciones obtenidas con nuestra propuesta con las proporcionadas por otros métodos de series temporales fuzzy que usan esquemas ponderados para la predicción puntual (Yu 2005, Cheng et al 2008), y con métodos que utilizan estrategias alternativas (Singh y Borah 2013, Wang et al 2014, Lu et al 2015). Utilizamos como serie de datos histórica el valor del índice IBEX35, entre los años 2011-2014 (Figura 1).

Referencias

- [1] Chen, S.: *Forecasting enrollments based on fuzzy time series*. Fuzzy Sets and Systems. **81(3)** (1996) 311–319.
- [2] Cheng, C. H., Chen, T. L., Teoh, H. J., Chiang, C. H.: *Fuzzy time-series based on adaptive expectation model for TAIEX forecasting*. Expert Systems with Applications. **34(2)** (2008) 1126–1132.
- [3] Dubois, D., Prade, H.: *The mean value of a fuzzy number*. Fuzzy Sets and Systems **24** (1987) 279–300.
- [4] Liu, H.-T.: *An improved fuzzy time series forecasting method using trapezoidal fuzzy numbers*. Fuzzy Optimization and Decision Making. **6(1)** (2007) 63–80
- [5] Lu, W., Chen, X., Pedrycz, W., Liu, X., Yang, J.: *Using interval information granules to improve forecasting in fuzzy time series*. International Journal of Approximate Reasoning. **57** (2015) 1–18.
- [6] Singh, P., Borah, B.: *An efficient time series forecasting model based on fuzzy time series*. Engineering Applications of Artificial Intelligence, **26(10)** (2013) 2443–2457.
- [7] Song, Q., Chissom, B. S.: *Forecasting enrollments with fuzzy time series-Part I*. Fuzzy Sets and Systems. **54(1)** (1993) 1–9.
- [8] Wang, L., Liu, X., Pedrycz, W., Shao, Y.: *Determination of temporal information granules to improve forecasting in fuzzy time series*. Expert Systems with Applications **41(6)** (2014) 3134–3142.
- [9] Yu, H.-K.: *Weighted fuzzy time series models for TAIEX forecasting*. Physica A: Statistical Mechanics and Its Applications. **349(3-4)** (2005) 609–624

Investigación parcialmente subvencionada por el Ministerio de Economía y Competitividad de España (Proyecto de investigación MTM2014-56233-P).

Banzhaf values for cooperative games with a proximity relation among the agents

Inés Gallego¹, Julio R. Fernández¹, Andrés Jiménez-Losada¹ and Manuel Ordóñez¹

¹ *Departamento de Matemática Aplicada II. Escuela Técnica Superior de Ingenieros.
Universidad de Sevilla. España*
ajlosada@us.es

Cooperative games [1] study situations of cooperation among a finite set of agents. A *cooperative game* is a pair (N, v) where N is the finite set of agents and $v : 2^N \rightarrow \mathbb{R}$ is the characteristic function with $v(\emptyset) = 0$, which establishes a worth for each coalition of players (subset of N). The worth for each coalition can be seen as profits, costs, debts... A payoff vector for a cooperative game (N, v) is a vector $x \in \mathbb{R}^N$ such that each component x_i is interpreted as the payoff obtained by player i at the end of the game for his cooperation in N . A *value* for cooperative games assigns to each game a payoff vector. A well known value is the *Banzhaf value* [2] b defined for any game (N, v) and player $i \in N$ as (if $n = |N|$)

$$b_i(N, v) = \frac{1}{2^n} \sum_{S \subseteq N \setminus \{i\}} [v(S \cup \{i\}) - v(S)].$$

An axiomatization of a value is a family of conditions which can be reasonable and determine our value. Axiomatizations allow us to choose the appropriate value for our situation. There are several axiomatizations of the Banzhaf value ([3],[4],...).

Owen [5] introduced situations for cooperative games with an additional information: a priori unions among the players defined by their relationships. A *game with a priori unions* is a triple (N, v, P) where (N, v) is a cooperative game and P is a partition of N determining the groupings by interest of the players which must be present in the bargaining of the payoffs of the players in a cooperation situation among all of them. Two different versions of the Banzhaf value modified with the information of the a priori unions were proposed: the Banzhaf-Owen value [5] and the symmetric coalitional Banzhaf value [6], both of them extending the classical Banzhaf value. These new values are obtained in two levels, bargaining among the unions and later bargaining into each union. Later Casajus [7] considered a more general model where in each union is also known a graph representing how the union is formed, namely the connections among the agents which determined the group. So, a *game with a cooperation structure* is a triple (N, v, L) where (N, v) is a game and (N, L) is a graph. The connected components of the graph are the unions of agents. Fernández et al. [8] extend the models of Owen and Casajus using proximity relations [9]. A *game with a proximity relation* is a triple (N, v, ρ) where (N, v) is a game and ρ is a proximity relation over N . If $i, j \in N$ then $\rho(i, j)$ is interpreted as the level of the relation between these players, the closeness of their ideas for instance.

Acknowledgements must be included in a footnote.

In this paper we extend the Banzhaf value to games with a proximity relation among the agents. First, following the Casajus model [7], we introduce two different versions for games with a cooperation structure in the same way as a priori unions: the graph Banzhaf-Owen value η and the Banzhaf-Myerson value ξ . Each of them defines for a given cooperative game and a player $i \in N$ a set function, $\alpha_i(N, v), \beta_i(N, v)$. Second we use the Choquet integral [10] as an aggregation function to introduce the following values, the prox-Banzhaf value A and the prox-symmetric Banzhaf value B as

$$\begin{aligned} A_i(N, v, \rho) &= \int \rho d\alpha_i(N, v), \\ B_i(N, v, \rho) &= \int \rho d\beta_i(N, v), \end{aligned}$$

for all (N, v, ρ) game with a proximity relation among the players and $i \in N$. Finally we provide both of the values with an axiomatization using logical properties in the fuzzy context of the situations.

References

- [1] T.S.H. Driessen: *Cooperative Games, Solutions and Applications*. Theory and Decision Library C. Springer (1988).
- [2] G. Owen: *Multilinear extensions and the Banzhaf value*. Naval Research Logistics Quarterly **22** (4) (1975) 741–750.
- [3] E. Lehrer: *An axiomatization of the Banzhaf value*. International Journal of Game Theory **17** (2) (1988) 89–99.
- [4] A. S. Nowak: *On an axiomatization of the Banzhaf value without the additivity axiom*. International Journal of Game Theory **26** (1) (1997) 137–141.
- [5] G. Owen: *Values of games with a priori unions*. Lecture Notes in Economics and Mathematical Systems **141** (1977) 76–88.
- [6] J.M. Alonso-Mejide and M.G. Fiestras-Janeiro: *Modification of the Banzhaf value for games with a coalition structure*. Annals of Operations Research **109** (1) (2002) 137–141.
- [7] A. Casajus: *Networks and outside options*. Social Choice Welfare **32** (2009) 1–13.
- [8] J.R. Fernández, I. Gallego, A. Jiménez Losada and M. Ordóñez: *Cooperation among agents with a proximity relation*. European Journal of Operational Research **250** (2) (2016) 555–565.
- [9] S. Ovchinnikov: *An introduction to fuzzy relations*. Fundamentals of Fuzzy Sets. Ed. D. Dubois and H. Prade. The Handbooks of Fuzzy Sets. Springer (2000).
- [10] G. Choquet: *Theory of capacities*. Annals of the Institute Fourier of Grenoble **5** (1953) 131–295.

Un enfoque multicriterio difuso para la selección de carteras socialmente responsables para inversores con aversión a las pérdidas.

Amelia Bilbao Terol¹, Mariano Jiménez López², Mar Arenas Parra¹, Verónica Cañal Fernández¹, Pablo Nguema Obama¹

¹ Universidad de Oviedo

ameliab@uniovi.es, mariammar@uniovi.es, vcanal@uniovi.es

² Universidad del País Vasco

mariano.jimenez@ehu.es

El concepto de “Responsabilidad Social Corporativa” (RSC) puede definirse como el compromiso de las empresas para contribuir a un desarrollo económico sostenible, con una mejora de la calidad de vida de los trabajadores, de la comunidad local y de la sociedad en general. La inversión Socialmente Responsable (SR) busca empresas con buen desempeño en RSC. El objetivo de este trabajo es diseñar un modelo de selección de carteras que considere tanto los objetivos financieros como de sostenibilidad [1,2]. La “Global Reporting Initiative” (GRI) sirve de marco para informar sobre el desempeño económico, ambiental y social de las organizaciones. En este trabajo elaboramos un sistema jerárquico de las categorías diseñadas en la GRI, a partir del cual se obtienen las puntuaciones de las empresas respecto a su cumplimiento en RSC. Hemos aplicado nuestra propuesta a ocho grandes empresas del IBEX-35. La información financiera ha sido obtenida de la agencia financiera Morningstar Ltd. y la relativa a la RSC de las empresas a partir del análisis de sus memorias de sostenibilidad [4].

Metodología.

Nuestro modelo trata de adaptarse a las preferencias particulares de un inversor SR con dos etapas en el proceso de construir la cartera adaptada tanto al perfil financiero como al de RS [2]. **En la primera etapa**, se obtiene una cartera convencional, en la que sólo se consideran los objetivos financieros. Para ello utilizamos la teoría de decisión bajo incertidumbre (Prospect Theory) de Kahneman-Tversky [3] con el beneficio neto como objetivo financiero. Un elemento clave de esta teoría es la función de valor en forma de S: cóncava (aversión al riesgo) en el dominio de las ganancias y convexa (gusto por riesgo) en el dominio de las pérdidas, ambos medidos con relación a un punto de referencia.

Hemos usado la siguiente función de valor:

$$u(f(x)) = \begin{cases} f(x)^\alpha, & f(x) \geq 0 \\ -\lambda(-f(x))^\alpha, & f(x) < 0 \end{cases} \quad \text{con } 0 < \alpha \leq 1 \text{ y } \lambda \geq 1.$$

Determinar los valores de los parámetros α (aversión al riesgo) y λ (aversión a las pérdidas) que se ajusten al perfil del inversor es muy difícil. Por ello optamos por utilizar etiquetas lingüísticas del tipo, “poco/moderado/muy adverso al riesgo” y “poco/moderado/muy adverso a las pérdidas”, modeladas por subconjuntos difusos.

Para obtener la cartera convencional, que nos servirá de referencia, se fija el perfil del inversor mediante una pareja de etiquetas lingüísticas para la aversión al riesgo (α) y a las pérdidas (λ). A partir de esta información asignamos valores discretos a α y a λ obteniendo las correspondientes carteras de máximo equivalente cierto sobre la frontera eficiente

asociada al valor final esperado (VFE) y al valor en riesgo condicionado (CVaR) [5]. De esta manera obtenemos un conjunto de carteras, cada una con un grado de pertenencia a las etiquetas lingüísticas. A continuación, utilizando los valores objetivo (VFE, CVaR) de cada una de ellas y su correspondiente grado de pertenencia, determinamos el centro de gravedad del conjunto de carteras convencionales. Elegimos como referencia aquella cartera que esté más próxima al centro de gravedad.

En la segunda etapa planteamos un modelo multiobjetivo que resolvemos usando la metodología de Programación por Metas Extendida [6]. Los valores financieros (VFE*, CVaR*) de la cartera de referencia elegida en la anterior etapa son los niveles de aspiración de las dos primeras metas. La tercera meta está relacionada con la responsabilidad social (SR) de la cartera y su nivel de aspiración lo determina el inversor:

$$\begin{aligned} \min \quad & \lambda[\omega_1 p_1 + \omega_2 n_2 + \omega_3 n_3] + (1-\lambda)D \\ \text{s.a} \quad & \\ & \xi + (1-a)^{-1} \sum_{j=1}^N \pi_j z_j + n_1 - p_1 = CVaR^*, \quad z_j \geq \sum_{i=1}^N (-P_{ij} x_i) + C_0 - \xi, \quad z_j \geq 0 \\ & \sum_{i=1}^N E[P_i] x_i + n_2 - p_2 = VFE^*, \quad \sum_{i=1}^N SR[P_i] x_i + n_3 - p_3 = SR^*, \\ & \sum_{i=1}^N P_{IT} x_i = C_0, \quad \omega_1 p_1 \leq D, \quad \omega_2 n_2 \leq D, \quad \omega_3 n_3 \leq D, \quad x \in X, \quad n_r, p_r \geq 0, r = 1, 2, 3 \end{aligned}$$

Conclusiones

La metodología propuesta permite obtener una cartera que satisface las preocupaciones financieras y de sostenibilidad de los inversores. El modelado mediante la teoría de los subconjuntos difusos permite un tratamiento flexible de los parámetros que definen la utilidad (o valor) financiera del inversor con un comportamiento que sigue la Teoría de Kahneman-Tversky. Con nuestra propuesta el inversor conoce el posible sacrificio financiero que conlleva su deseo de invertir en empresas comprometidas con la responsabilidad social.

Referencias

- [1] A. Bilbao-Terol, M. Arenas-Parra, V. Cañal-Fernández: *Selection of Socially Responsible Portfolios using Goal Programming and fuzzy technology*. Information Sciences 189 (2012), 110-125.
- [2] A. Bilbao-Terol, M. Arenas-Parra, V. Cañal-Fernández: *A model based on Copula Theory for sustainable and social responsible investments*. Revista de Contabilidad (2014) (en prensa).
- [3] D. Kahneman, A. Tversky: *Prospect Theory: An Analysis of Decision under Risk*. Econometrica XLVII (1979), 263-291.
- [4] P. Nguema Obama Eyang: *Modelos Multicriterio para Inversiones Socialmente Responsables*. Trabajo Fin de Máster en Modelización e Investigación Matemática Estadística y Computación. (2014).
- [5] R.T. Rockafellar, S. Uryasev: *Optimization of conditional value-at-risk*. Journal of Risk 2 (2000) 21-41.
- [6] C. Romero: *Extended lexicographic goal programming approach: an unifying approach*. OMEGA, International Journal of Management Science 29 (2001) 63-71.

Ordenación de fondos de inversión mediante valoración difusa de su responsabilidad social

María Teresa Lamata¹, Vicente Liern², Blanca Pérez-Gladish³

¹ Dpto. Ciencias de la Computación e Inteligencia Artificial,
Universidad de Granada, España
mtl@decsai.ugr.es

² Dpto. Matemáticas para la Economía y la Empresa,
Universidad de Valencia, España
vicente.liern@uv.es

³ Dpto. Economía Cuantitativa, Universidad de Oviedo, España
bperez@uniovi.es

1. Introducción

En este trabajo hemos abordado el problema de la evaluación del desempeño social de los fondos de inversión de renta variable. Como son activos financieros que cambian muy a menudo su composición (en términos de empresas en las que invierten), la evaluación del grado de responsabilidad social de este tipo de producto financiero resulta complicada. La naturaleza cualitativa, imprecisa e incierta de los criterios de inversión social, dificulta la valoración del comportamiento de los fondos, puesto que depende de las percepciones subjetivas de un experto o un grupo de expertos. Ante esta perspectiva, nos enfrentamos a un problema multicriterio en el que buscamos la solución teniendo en cuenta la importancia y credibilidad de la información proporcionada por agencias y expertos.

En un intento de superar este problema, hemos evaluado la responsabilidad social de un conjunto de fondos de inversión de renta variable mediante números difusos triangulares que incluyen el conocimiento y la experiencia de una experta en Inversión Socialmente Responsable (ISR) y de un analista. Estas evaluaciones difusas han sido luego integradas en un método de ordenación basado en TOPSIS [1, 2] que da lugar a diferentes ordenaciones de los fondos de inversión en función de la confianza del experto en las evaluaciones de los fondos. Aunque existen trabajos en los que se realiza una propuesta de evaluación del grado de responsabilidad social de activos financieros (por ejemplo en [3] se utiliza el método AHP o en [4] el método TOPSIS), presentamos un enfoque que resulta novedoso tanto en la construcción de los números difusos como en su incorporación del método TOPSIS.

2. Medición del grado de responsabilidad social impreciso de un fondo de inversión

Consideramos m fondos de inversión en renta variable F_i , $i=1, \dots, m$ y n criterios de responsabilidad social C_j , $j=1, \dots, n$. El comportamiento de cada fondo ha sido analizado por una experta en ISR que valoró el comportamiento de cada fondo en cada criterio social mediante un número real comprendido entre 0 y 1.

Dado que en la práctica un mismo activo financiero puede obtener distintas puntuaciones en función de la agencia de rating social consultada, le pedimos a la experta que el valor que nos proporcionase fuera el que en su opinión tuviera mayor posibilidad de ocurrencia. A partir de la puntuación de mayor posibilidad de ocurrencia, un analista

financiero contruyó números difusos triangulares para cada fondo y cada criterio de la siguiente forma:

$$T_{ij} = (a_{ij}^L, a_{ij}, a_{ij}^R), \quad \text{donde } a_{ij}^L = a_{ij} - k_{ij}^L b_{ij}^L, \quad a_{ij}^R = a_{ij} + k_{ij}^R b_{ij}^R,$$

donde a_{ij} es el valor con mayor posibilidad de ocurrencia, proporcionado por la experta y, k_{ij}^L, k_{ij}^R son respectivamente, las tolerancias a la izquierda y a la derecha asignados por el analista junto con la experta, a partir de datos de la serie histórica de F_i .

Por propia construcción, el valor de la función de pertenencia de T_{ij} , $\mu_{T_{ij}}(x)$, representa el grado de verdad que se concede a la afirmación: “El fondo i obtiene una valoración x con respecto al criterio social j ”.

3. Método de ordenación basado en TOPSIS

El método propuesto para la ordenación de los fondos valorados mediante números difusos se basa en TOPSIS [3] y en el cálculo de los alfa-cortes de los números T_{ij} .

- STEP 1. Fijar un grado de confianza alfa.
- STEP 2. Determinar la matriz de decisión (formada por intervalos) para alfa
- STEP 3. Construir la matriz ponderada.
- STEP 4. Establecer el fondo deseado y el no deseado (ideal positivo, negativo).
- STEP 5. Calcular las medidas de separación.
- STEP 6. Calcular la proximidad relativa al fondo deseado (ideal).
- STEP 7. Ordenar las alternativas.

Hemos aplicado nuestro método a 25 fondos del US SIF, US Social Investment Forum, (<http://www.ussif.org>) y al inversor, se le proporciona una tabla que contiene las ordenaciones para diferentes grados de confianza. Con los ejemplos analizados, hemos comprobado que las primeras y últimas posiciones se mantienen estables para los diferentes niveles, sin embargo las posiciones intermedias varían.

4. Conclusiones

El modelo propuesto proporciona un método para ordenar los fondos de inversión de acuerdo con los criterios de responsabilidad social. Además de ser fácil de usar y de contar con gran flexibilidad, es capaz de establecer un ranking, con el que hasta la fecha no cuentan los inversores. Sin duda, esto puede facilitar una selección adecuada de fondos socialmente responsables sin que la responsabilidad recaiga exclusivamente en la información proporcionada por las agencias de *rating* social.

Referencias

- [1] Hwang, C.L., Yoon, K.: Multiple Attributes Decision Making Methods. Springer. Berlín (1981).
- [2] García-Cascales, M. S., Lamata, M. T.: *On rank reversal and TOPSIS method*. Mathematical and Computer Modelling **56** (2012), 123-132.
- [3] Pérez-Gladish, B., M'Zali, B.: *An AHP-based approach to mutual funds social performance measurement*. International Journal of Multicriteria Decision Making **1(1)** (2010) 103-127.
- [4] Bilbao-Terol, A., Arenas-Parra, M., Cañal, V., Antomil, J.: *Using TOPSIS for assessing the sustainability of government bond funds*. Omega **49** (2014) 1-17.

Programación Lineal Difusa y Modelos de Localización con Cobertura para Diseminación de Información en VANETs

Antonio D. Masegosa^{1,2}, Idoia de la Iglesia¹, Unai Hernandez-Jayo¹

¹ *Instituto Tecnológico de Deusto (DeustoTech), Universidad de Deusto, Bilbao, España*
{ad.masegosa,idoia.delaiglesia,unai.hernandez}@deusto.es

² *IKERBASQUE, Fundación Vasca para la Ciencia, Bilbao, España*

La VANETs son redes ad-hoc en las que los nodos que intercambian información son vehículos. Este tipo de redes han despertado el interés tanto de la comunidad científica como de la industria automovilística o de instituciones gubernamentales por el gran número de aplicaciones innovadoras que podrían hacer posible [4]. Algunas de las áreas de aplicación más importantes son la seguridad (avisos de frenada de emergencia, de colisión en intersección, de cambio de carril, etc.); la gestión de tráfico (semáforos virtuales, zonas de acceso limitado, pago electrónico de peaje, etc.); la asistencia al conductor (diagnóstico remoto de averías, consejos de conducción eficiente, etc.); o el ocio y entretenimiento (descarga de contenidos multimedia, puntos de interés cercanos, etc). Las comunicaciones que tienen lugar dentro de las VANETs pueden clasificarse en Infraestructura-a-Vehículo (I2V), Vehículo-a-Infraestructura (V2I) y Vehículo-a-Vehículo (V2V). Por otra parte, muchos de estos servicios requieren de un servidor central que canaliza toda la información, lo que genera una necesidad de recolección y diseminación de datos desde y hacia los vehículos [1], respectivamente.

En este trabajo nos centramos en la diseminación de información desde un servidor central hacia los vehículos. Esta tarea presenta importantes retos [1] en un entorno vehicular, especialmente cuando se hace a gran escala, por diferentes razones como: la diversidad de la topología de la red, la alta velocidad de los vehículos, la desigual densidad de los mismos, la falta de infraestructura en las carreteras, o las tecnologías de comunicación disponibles y su diferente grado de penetración. Uno de los enfoques que se han propuesto para abordar este problema es el uso de infraestructuras virtuales (IVs) [1]. Las IVs son nodos estáticos (RSUs) o dinámicos (Vehículos), que recibirían la información del servidor central y la difundirían a los vehículos cercanos mediante comunicaciones I2V o V2V, respectivamente.

La correcta selección de los vehículos que actúan como IVs juega un papel importante en este tipo enfoques ya que puede evitar la necesidad de infraestructuras fijas (RSUs), reducir la sobrecarga de la red, o influir en la calidad de la comunicación. Esta selección depende del servicio concreto que se quiera implementar y de los requisitos fijados por el estándar: tipo de mensaje a enviar, tiempo de respuesta, prioridad, área a cubrir, tecnologías de comunicación disponibles o grado de alcance del mensaje. En muchos de los servicios mencionados anteriormente, la selección de la IV tiene como objetivo determinar el número mínimo de vehículos que permiten diseminar la información a un porcentaje prefijado de los vehículos

Este proyecto ha sido financiado por los proyectos de investigación TEC2013-45585-C2-2-R y TIN2014-56042-JIN del Ministerio Español de Economía y Competitividad, y PC2013-71A del Gobierno Vasco.

que se encuentran en una determinada área geográfica, cumpliendo con los requerimientos del estándar.

El proceso de selección que acabamos de describir se puede modelar como un Problema de Cobertura de Conjuntos (PCC)[2]. Dado un conjunto de nodos de demanda, y un conjunto de posibles ubicaciones para instalaciones, el objetivo del PCC consiste en determinar el número mínimo de instalaciones necesarias para cubrir toda la demanda, y su ubicación. En el caso que aquí nos concierne, los nodos de demanda se corresponderían con los vehículos a los que se les quiere enviar la información, y las instalaciones con aquellos vehículos que cuentan con comunicaciones I2V, y por tanto pueden actuar como IV.

Sin embargo, la aplicación de los modelos clásicos de PCC puede no ser adecuado, ya que suponen todos los datos del problema como perfectamente conocidos. Como ocurre en otras muchas situaciones prácticas, la información disponible en estos escenarios suele presentar incertidumbre, imprecisión o ambigüedad. Para abordar dicho problema proponemos el uso de modelos de localización de cobertura de conjuntos con restricciones difusas [3]. De esta manera, queremos diseñar modelos que permitan establecer una IV con cierta tolerancia en el cumplimiento de requerimientos, como el porcentaje de vehículos a los que debe llegar el mensaje, el número máximo de vehículos asociados a un nodo IV y la potencia máxima de la señal de comunicación. Para resolver el modelo planteado hemos usado el enfoque paramétrico propuesto por Verdegay en [5]. Este enfoque lleva a la resolución de un problema auxiliar de programación lineal paramétrico para el que se aplicó un Algoritmo Genético.

Los resultados obtenidos muestran que el modelo propuesto ofrece un marco teórico para el diseño de IVs en VANETS en las que existe imprecisión o vaguedad en los requisitos que ha de cumplir la diseminación de información, y que puede ser aplicado para la implementación de diferentes servicios que requieren distintos niveles de tolerancia en el cumplimiento de las restricciones.

Referencias

- [1] P. M. D'Orey, N. Maslekar, I. de la Iglesia, and N. K. Zahariev. NAVI: Neighbor-Aware Virtual Infrastructure for Information Collection and Dissemination in Vehicular Networks. In *2015 IEEE Vehicular Technology Conference (VTC Spring)*, pages 1–6, 2015.
- [2] M. L. Fisher and P. Kedia. Optimal Solution of Set Covering/Partitioning Problems Using Dual Heuristics. *Management Science*, 36(6):674–688, jun 1990.
- [3] V. C. Guzmán, A. D. Masegosa, D. Pelta, and J. L. Verdegay. Fuzzy models and resolution methods for covering location problems: an annotated bibliography. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 2016. En prensa.
- [4] P. Papadimitratos, A. La Fortelle, K. Evenssen, R. Brignolo, and S. Cosenza. Vehicular communication systems: Enabling technologies, applications, and future outlook on intelligent transportation. *IEEE Communications Magazine*, 47(11):84–95, 2009.
- [5] J. L. Verdegay. Fuzzy mathematical programming. In *Fuzzy information and decision processes*, pages 231–237. 1982.

Optimización de la fase de aprendizaje del Sistema de Clasificación Basado en Reglas Difusas para Big Data Chi-FRBCS-BigDataCS

Mikel Elkano¹, Mikel Galar¹, José Sanz¹, Humberto Bustince¹

¹ Dpto. de Automática y Computación, Universidad Pública de Navarra, Campus Arrosadía
s/n, 31006 Pamplona, España

{mikel.elkano, mikel.galar, joseantonio.sanz, bustince} @unavarra.es

En este trabajo presentamos una optimización del único Sistema de Clasificación Basado en Reglas Difusas que es capaz de afrontar problemas Big Data hasta la fecha, esto es, Chi-FRBCS-BigDataCS [1]. Este método se basa en el algoritmo original de Chi et al. [2] y consiste en el aprendizaje concurrente de varios clasificadores Chi cuyas bases de reglas se agregan al final. La paralelización del método se lleva a cabo con el popular framework *Apache Hadoop*.

En esta propuesta presentamos varias optimizaciones del algoritmo Chi-FRBCS-BigDataCS que agilizan el proceso de aprendizaje del mismo sin alterar el modelo obtenido (y por tanto la precisión de la clasificación). Las modificaciones llevadas a cabo son las siguientes:

- *Aprovechando los pares clave-valor.* Proponemos emplear las fases de *sorting* y *merging* de MapReduce para resolver los duplicados y los conflictos de las bases de reglas y evitar la búsqueda exhaustiva que se realiza en el método original. Esto es posible gracias a que en nuestra aproximación cada *mapper* envía al *reducer* las reglas por separado (un par clave-valor para cada regla).
- *Permitiendo la ejecución de múltiples reducers.* Chi-FRBCS-BigDataCS no permite la ejecución de más de un reducer, lo que puede suponer un cuello de botella. El flujo de datos presentado en la anterior optimización nos permite la ejecución de múltiples reducers, aumentando el grado de paralelización tanto como se desee.
- *Pre-cálculo de los grados de pertenencia mediante Look-Up-Tables.* Proponemos añadir una fase de pre-cálculo mediante Look-Up-Tables para asegurar que el grado de pertenencia de un ejemplo a una determinada etiqueta lingüística se calcula una sola vez durante el cálculo de los pesos de las reglas llevado a cabo en los *mapper*.

Los 6 conjuntos de datos empleados para evaluar la mejora obtenida por nuestra propuesta se han obtenido del repositorio UCI [3]. Sin embargo, con el objetivo de obtener más conjuntos de datos y de trabajar con el mismo marco experimental que en [1], hemos seleccionado 15 problemas de clasificación binarios diferentes convirtiendo los conjuntos de datos originales multi-clase en múltiples problemas binarios *One-vs-Rest*. Para ello, en cada

Este trabajo ha sido financiado parcialmente por el Ministerio de Ciencia y Tecnología de España con el proyecto TIN2013-40765-P y por la red TIN2014-56381-REDT.

conjunto de datos original seleccionamos una clase positiva y consideramos el resto como la clase negativa.

Los resultados de la comparativa entre nuestra optimización y el método original se muestran en la Tabla 1. Hemos omitido Susy de la media debido al tiempo requerido para ejecutar todas las combinaciones de mappers y particiones. No obstante, en la Tabla 1 incluimos una comparativa para Susy empleando 32 mappers (el máximo que podemos ejecutar de forma concurrente en el cluster utilizado). De acuerdo a la Tabla 1, podemos apreciar que el algoritmo de aprendizaje de nuestra aproximación se ejecuta entre 10 y 100 veces más rápido que el original. Cabe destacar que la mejora obtenida será cada vez mayor conforme el número de ejemplos y reglas aumenta.

Tabla 1: Promedio del tiempo de ejecución de los mappers (mm:ss).

Dataset	32 mappers		64 mappers		128 mappers		256 mappers	
	Original	Propuesta	Original	Propuesta	Original	Propuesta	Original	Propuesta
AVG.	03:05	00:16	00:46	00:04	00:11	00:00	00:02	00:00
Susy	2562:07	26:49	-	-	-	-	-	-

Referencias

- [1] V. López, S. del Río, M. Benitez, y F. Herrera: *Cost-sensitive linguistic fuzzy rule based classification systems under the MapReduce framework for imbalanced big data*. Fuzzy Sets and Systems. **258** (Año) 5–38.
- [2] Z. Chi, H. Yan, y T. Pham: *Fuzzy algorithms with applications to image processing and pattern recognition*. World Scientific.
- [3] M. Lichman: *UCI Machine Learning Repository*. <http://archive.ics.uci.edu/ml>.

Sistemas de Clasificación Basados en Reglas Difusas para Big Data con MapReduce: Análisis de la Granularidad

Alberto Fernández¹, Sara del Río², Francisco Herrera²

¹ *Departamento de Informática, Universidad de Jaén, Jaén, España*

alberto.fernandez@ujaen.es

² *Depto. de Ciencias de la Computación e I.A., Universidad de Granada, Granada, España*

srio@decsai.ugr.es, herrera@decsai.ugr.es

En este trabajo vamos a estudiar el uso de los Sistemas de Clasificación Basados en Reglas Difusas (SCBRDs) [3] en el contexto de los problemas Big Data [2]. En concreto, haremos uso del primer modelo adaptado al esquema MapReduce [1], denominado como Chi-FRBCS-BigData [4]. Nuestro objetivo final será el de estudiar la dependencia entre el nivel de granularidad asociado a las etiquetas difusas, y el número de “Maps” seleccionados para la ejecución paralela del modelo.

Nuestra hipótesis se basa en que a priori algunos problemas pueden alcanzar resultados de calidad sólidos utilizando una baja granularidad. Sin embargo, conforme el número de Mapas se incrementa, también lo hace la dispersión de los datos. Así, en este caso el modelo puede beneficiarse de un mayor número de etiquetas, de acuerdo a una representación más fina del espacio del problema.

Los conjuntos de datos seleccionados para este estudio se presentan en la Tabla 1. Los resultados experimentales han sido obtenidos con una configuración para 3, 5 y 7 etiquetas, y utilizando 32, 64 y 128 Maps. Éstos se muestran en la Tabla 2, utilizando la media de 10 particiones para la medida de precisión estándar.

Tabla 1: Resumen de los conjuntos de datos

Conjunto de datos	#Ej.	#Atrs.	Clases seleccionadas	#Ejemplos por clase
Kddcup_DOS_vs_normal	4856151	41	(DOS; normal)	(3883370; 972781)
Poker_0_vs_1	946799	10	(0; 1)	(513702; 433097)
Covtype_2_vs_1	495141	54	(2; 1)	(283301; 211840)
Census	141544	41	(-50000.; 50000+.)	(133430; 8114)
Fars_Fatal_Inj_vs_No_Inj	62123	29	(Fatal_Inj; No_Inj)	(42116; 20007)

Observamos que los mejores resultados se alcanzan con el más alto nivel de granularidad. Además, podemos comprobar que se obtiene una mejor estabilidad para la dispersión de datos (al seleccionar más Maps) también en aquéllos casos con mayor número de etiquetas.

Ello se debe a que una representación más fina del espacio del problema permite un mejor reconocimiento en las áreas frontera, incluso con menos ejemplos sobre los que aprender. Además, la fusión de reglas en este caso resulta en una BR más diversa. En concreto, consideramos que cuánto menor sea el nivel de granularidad, hay una mayor cantidad de reglas

Este trabajo ha sido parcialmente financiado por el Ministerio Español de Ciencia y Tecnología bajo el proyecto TIN2014-57251-P; y el Plan de Investigación Andaluz P12-TIC-2958.

Tabla 2: Resultados en precisión para Chi-FRBCS-BigData. El mejor resultado por problema y número de Maps está marcado en negrita, el mejor total se indica con subrayado.

Conjunto de Datos	32 Maps		
	3 etiquetas	5 etiquetas	7 etiquetas
Kddcup_DOS_vs_normal	99.92498	99.95345	99.95500
Poker_0_vs_1	58.92973	56.87383	59.98018
Covtype_2_vs_1	74.61723	83.24885	87.43990
Census	93.48355	94.68176	95.54433
Fars_Fatal_Inj_vs_No_Inj	94.25642	99.86419	99.94522
Promedio	84.24238	86.92441	88.57293
	64 Maps		
	3 etiquetas	5 etiquetas	7 etiquetas
Kddcup_DOS_vs_normal	99.92490	99.95167	99.95543
Poker_0_vs_1	57.94608	56.78075	59.94370
Covtype_2_vs_1	74.52195	82.84057	87.20926
Census	93.30058	94.54349	95.50047
Fars_Fatal_Inj_vs_No_Inj	93.98393	99.82303	99.94522
Promedio	83.93549	86.78790	88.51082
	128 Maps		
	3 etiquetas	5 etiquetas	7 etiquetas
Kddcup_DOS_vs_normal	99.92525	99.95155	99.95425
Poker_0_vs_1	56.96169	56.79547	59.93157
Covtype_2_vs_1	74.01238	82.46134	86.77500
Census	92.97315	94.38729	95.45372
Fars_Fatal_Inj_vs_No_Inj	93.81633	99.81723	99.94522
Promedio	83.53776	86.68257	88.41195

repetidas para cada BR_i independiente. Esto implica que al incrementar el número de etiquetas, permite que las reglas se adscriban localmente a su propio espacio de datos, resultando en una mejor cobertura global tras la etapa Reduce.

Como trabajo futuro, planeamos mejorar el esquema de Chi-FRBCS-BigData mediante: (1) diferentes niveles de granularidad durante la búsqueda de las reglas para una mejor contextualización del problema; y (2) aprovechar las características del modelo MapReduce para extender la función Reduce y así decrementar el tiempo de cómputo del algoritmo.

Referencias

- [1] Jeffrey Dean and Sanjay Ghemawat. MapReduce: A flexible data processing tool. 53(1):72–77, 2010.
- [2] A. Fernández, S. Río, V. López, A. Bawakid, M.J. del Jesus, J.M. Benítez, and F. Herrera. Big data with cloud computing: An insight on the computing environment, mapreduce and programming framework. 4(5):380–409, 2014.
- [3] H. Ishibuchi, T. Nakashima, and M. Nii. *Classification and modeling with linguistic information granules: Advanced approaches to linguistic data mining*. Springer-Verlag, Berlin, Germany, 2004.
- [4] S. Río, V. López, J.M. Benítez, and F. Herrera. A mapreduce approach to address big data classification problems based on the fusion of linguistic fuzzy rules. *International Journal of Computational Intelligence Systems*, 8(3):422–437, 2015.

Un estudio preliminar sobre el uso de técnicas de aprendizaje incremental para problemas de big data

Juan Carlos Gámez¹, David García², Antonio González², Raúl Pérez²

¹ *Departamento de Arquitectura de Computadores, Electrónica y Tecnología Electrónica
Universidad de Córdoba, Córdoba, Spain*
jcgamez@uco.es

² *Departamento de Ciencias de la Computación e Inteligencia Artificial
Universidad de Granada, Granada, Spain*
{dgarcia, A.Gonzalez, fgr}@decsai.ugr.es

Habitualmente los algoritmos de aprendizaje de reglas tienden a deteriorar su comportamiento cuando tienen que procesar una cantidad masiva de datos. Los problemas que presentan esta característica se engloban dentro de los llamados "big data". El modelo de programación MapReduce [2] se propuso como una forma de resolver este tipo de problemas, y se basa en el uso de algoritmos paralelos y distribuidos. Básicamente consiste en la aplicación de dos operaciones básicas, la función *map* que es responsable de dividir la base de datos original y procesar cada problema por separado, y la función *reduce* que junta y agrega los resultados de la función anterior.

Recientemente [4] se ha propuesto un algoritmo de clasificación basado en reglas difusas lingüísticas para problemas de big data. Este algoritmo usa el modelo MapReduce y toma como algoritmo base el método de aprendizaje propuesto por Chi et al [1].

En este trabajo exploramos el uso de algoritmos de aprendizaje incremental como herramienta alternativa al modelo MapReduce. Esto supone la utilización de un modelo secuencial en vez de un modelo paralelo.

Un algoritmo de aprendizaje de reglas difusas incremental, procesa en distintas etapas o episodios, nuevos ejemplos que se van recibiendo a lo largo del tiempo.

Con la idea de utilizar el potencial de estos algoritmos para problemas de big data, definimos un procedimiento que inicialmente divide el conjunto de ejemplos en varios subconjuntos, que definen los distintos episodios. Esta etapa coincide con la división que se propone en el modelo MapReduce, sin embargo el procesamiento es completamente diferente, como se observa en la Figura 1. A continuación, en cada paso se lanza un algoritmo de aprendizaje incremental que tiene como entrada un nuevo episodio y el conjunto de reglas aprendidas en la etapa anterior, y ofrece como salida un nuevo conjunto de reglas que añade al conocimiento de partida, la información aportada por el nuevo episodio. Obviamente en el primer paso el conjunto de partida de reglas difusas es el conjunto vacío.

En el trabajo proponemos un algoritmo de aprendizaje incremental de reglas difusas que en una primera etapa comprueba si el comportamiento del conocimiento previamente

Este trabajo está parcialmente financiado por Ministerio de Economía y Competitividad a través del proyecto de investigación TIN2015-71618-R. Agradecemos a los autores del trabajo [4] por su colaboración en el desarrollo de esta propuesta, permitiéndonos acceder a las particiones que utilizaron en su estudio experimental.

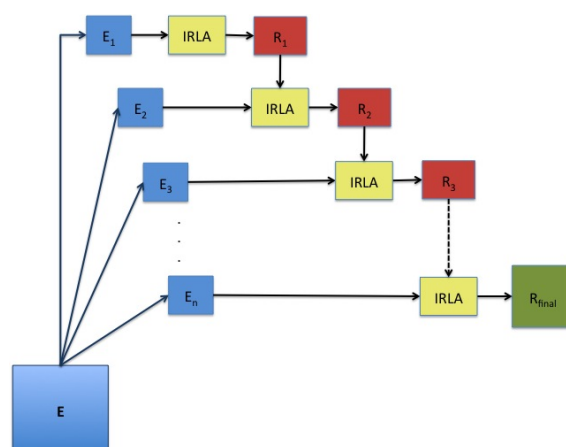


Figura 1: Uso de un modelo incremental de aprendizaje de reglas para problemas big data. IRLA representa un algoritmo general de aprendizaje incremental de reglas difusas.

adquirido es aceptable para el nuevo episodio (la aceptabilidad viene definida por un umbral que trata de estimar el porcentaje final de acierto) en cuyo caso no es necesario realizar ninguna tarea y se procesa el siguiente subconjunto de ejemplos. En otro caso, será necesario realizar un ajuste del conocimiento. Para ello se inicia una segunda etapa en la que se hace una llamada a un algoritmo de aprendizaje de reglas, en nuestro caso NSLV-AR [3], para que revise y actualice el conocimiento. En una etapa final se combinan las reglas del conocimiento previo con las nuevas reglas obtenidas por el algoritmo base. Esta etapa es fundamental y está relacionada con la función *reduce* del modelo MapReduce. Sin embargo, la diferencia fundamental es que el modelo es secuencial en vez de paralelo.

El estudio experimental realizado muestra que nuestra propuesta es capaz de obtener un conjunto reducido de reglas difusas, con una predicción similar al modelo paralelo, pero con muy buenos resultados en tiempo.

Referencias

- [1] Z. Chi, H. Yan, T. Pham: *Fuzzy algorithms: with applications to image processing and pattern recognition*. World Scientific (1996) **vol. 10**.
- [2] J. Dean, S. Ghemawat: *MapReduce: a flexible data processing tool*. Communications of the ACM, **vol. 53**, No 1 (2010) 72–77.
- [3] D. García, J.C. Gámez, A. González, R. Pérez: *An interpretability improvement for fuzzy rule bases obtained by the iterative rule learning approach*. International Journal of Approximate Reasoning. **vol. 67**, (2015) 37–58.
- [4] S. del Río, V. López, J. M. Benítez, F. Herrera: *A MapReduce approach to address big data classification problems based on the fusion of linguistic fuzzy rules*. International Journal of Computational Intelligence Systems. **vol. 8**, No 3, (2015) 422–437.

A MapReduce Implementation of a Genetic Fuzzy System for Regression

I. Rodríguez-Fdez¹, M. Mucientes¹, and A. Bugarín¹

¹ *Centro Singular de Investigación en Tecnoloxías da Información (CiTIUS)
Universidade de Santiago de Compostela. Santiago de Compostela, Spain*

ismael.rodriguez@usc.es, manuel.mucientes@usc.es,
alberto.bugarin.diz@usc.es

In this paper we present S-FRULER (Scalable Fuzzy Rule Learning through Evolution for Regression), a scalable MapReduce version of FRULER [1] which is a Genetic Fuzzy System (GFS) that learns simple and linguistic TSK-1 knowledge bases for regression problems. S-FRULER obtains models with high accuracy and low complexity, whilst reducing the algorithm runtime. S-FRULER focuses on splitting the problem into smaller partitions and incorporates a feature selection process for reducing the number of variables used in each partition. Each partition is then solved independently using the FRULER algorithm. The Reduce function obtains linguistic TSK fuzzy rule bases from the information generated in each partition.

The main contributions of this work are: i) a distributed GFS approach which uses a MapReduce technique and the Spark software, ii) a random feature selection process to reduce the number of variables used in each dataset partition, and iii) an aggregation function to obtain linguistic TSK fuzzy rule bases from the rule bases obtained in each dataset partition.

In GFSs the size of the problem has a huge influence in the performance of the obtained models, since i) the fuzzy rule bases learned suffer from exponential rule explosion when the number of variables increases, and ii) the convergence time increments with the number of examples. S-FRULER addresses the problem following a MapReduce approach that uses FRULER [1]. FRULER is composed of a two-stage preprocessing — formed by an instance selection and a multi-granularity fuzzy discretization—, and a genetic algorithm, which contains an ad-hoc TSK-1 rule generation module. Although the runtime of this approach is acceptable for medium size datasets, it does not scale properly when solving large scale problems as it may not converge. To solve that, S-FRULER divides the problem into small problems that are tractable using the MapReduce approach. Each of the divisions is solved independently in the Map phase using FRULER. Then, the solutions obtained in each Map are combined in the Reduce phase in order to obtain a final solution for the original data.

The algorithm structure is shown in figure 1. First, the multi-granularity fuzzy discretization process is performed using the whole training dataset. Then, the training dataset is splitted into n_{map} partitions, which correspond to the tasks to be distributed into the Working Nodes.

This research was supported by the Spanish Ministry of Economy and Competitiveness (grant TIN2014-56633-C3-1-R) and the Galician Ministry of Education (grants EM2014/012, CN2012/151 and GRC2014/030). All grants were co-funded by the European Regional Development Fund (FEDER program).

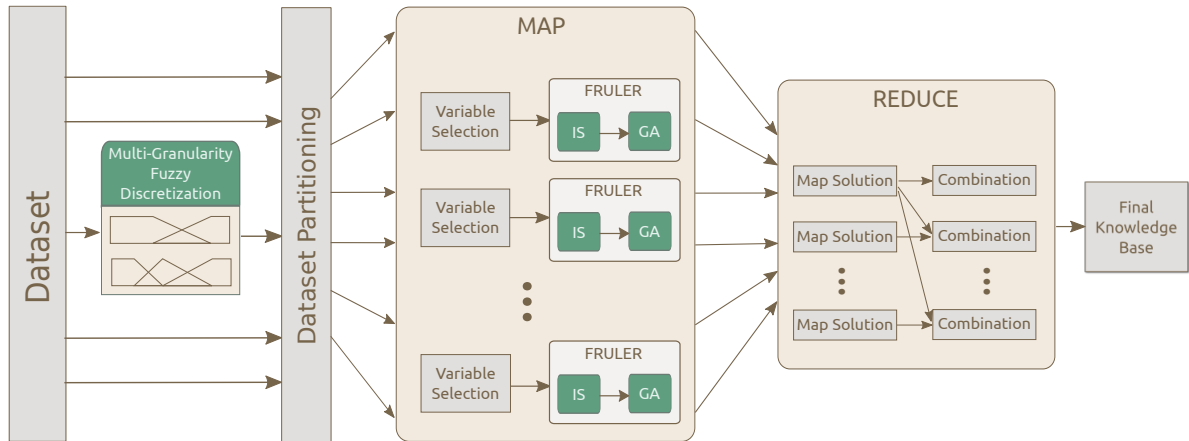


Figure 1: S-FRULER architecture showing the preprocessing, Map and Reduce phases.

The Map function is applied independently and distributively to each training dataset partition. For each of these partitions, a subset of randomly selected variables is used. Thus, all the procedure done in each task uses the corresponding subset of examples and subset of variables, which is going to be referred as the dataset partition. For each task, FRULER is applied without the discretization step, since it was already performed before the partition of the data. It is composed by an instance selection of the most representative examples and a genetic algorithm, which contains an ad-hoc TSK-1 rule generation module. The evolutionary learning process obtains a definition of the data base. Then an ad-hoc TSK-1 rule generation module obtains the antecedents and consequents of each possible rule using only the representative examples.

After the execution of FRULER for each dataset partition the reduction phase is applied in order to get the final solution. Each of the FRULER executions obtains a TSK-1 Knowledge Base, composed by a linguistic partition of the input variables and the TSK fuzzy rule set. To combine the solutions generated in the Map phase without increasing significantly the complexity, S-FRULER takes into account the properties of each of the obtained Knowledge Bases, and combines their granularities in a similar way to a crossover operation for integer valued individuals. Finally, the rule base can be generated using the ad-hoc TSK rule base Generation process of FRULER.

S-FRULER has speedups usually larger than the number of dataset partitions used, showing a scalability higher than linear in both standalone and cluster modes. It was compared with three GFSs in terms of complexity of the models (number of rules) and precision (mean squared error) using 10 large-scale datasets, showing good results with less complex models. Results demonstrate the capability of S-FRULER to obtain precise and simple models in large scale problems.

References

- [1] I Rodríguez-Fdez, M Mucientes, and A Bugarín. FRULER: Fuzzy rule learning through evolution for regression. *arXiv preprint arXiv:1507.04997*, 2015.

Detección de estructuras curvilíneas usando la transformación morfológica borrosa todo-nada

Pedro Bibiloni, Manuel González-Hidalgo, Sebastia Massanet

Universitat de les Illes Balears, Dept. Matemàtiques i Informàtica

Crta. Valldemossa km 7.5, 07122 Palma, Illes Balears

{p.bibiloni,manuel.gonzalez,s.massanet}@uib.es

La extracción de estructuras curvilíneas en imágenes es un paso esencial en muchas aplicaciones. En particular, tienen un papel importante en la detección de grietas en edificios, tuberías o carreteras, en la detección de los vasos sanguíneos en imágenes médicas o en el reconocimiento de huellas dactilares. La detección de estas estructuras no es una tarea sencilla como demuestra la publicación de numerosos algoritmos basados en las más diversas teorías. En [2] se puede consultar una exhaustiva recopilación de las diferentes técnicas. Una de las principales razones por las que no existe aún un algoritmo que supere al resto en todas las aplicaciones es la dificultad de definir formalmente qué se entiende por *estructura curvilínea*. Aunque en términos generales pueden ser definidas como regiones estrechas, alargadas y localmente lineales cuyos píxeles tienen intensidades diferentes a los píxeles del vecindario, la definición debe concretarse para ser adaptada a cada aplicación concreta.

Entre las distintas técnicas existentes para la detección de estructuras curvilíneas destaca la morfología matemática en niveles de gris. Esta teoría dispone de un conjunto de operadores locales simples que, combinados, son capaces de extraer información de la imagen basada en las características geométricas y fotométricas del entorno de un píxel mediante el uso de elementos estructurantes adecuados. En particular, en [1] se propone un detector de estructuras curvilíneas basado en la transformación morfológica todo-nada con elementos estructurantes lineales multidireccionales y multiescala. Este operador morfológico es capaz de detectar grupos de píxeles que satisfacen unas restricciones geométricas determinadas o que forman una configuración específica con notable éxito.

Sin embargo, la morfología matemática en niveles de gris no es capaz de manejar adecuadamente la incertidumbre y ambigüedad presente en las imágenes naturales. Por este motivo, se introdujo la teoría de conjuntos borrosos y de la lógica borrosa en los operadores morfológicos, generando la morfología matemática borrosa que mejora en muchas aplicaciones (detección de contornos, filtrado de ruido, etc.) a los otros paradigmas morfológicos. Dentro de esta teoría, en [3] se introdujo la transformación morfológica borrosa todo-nada (FMHMT) mediante el uso de t-normas y negaciones borrosas. Este operador ha resultado útil en el diseño de algoritmos capaces de detectar determinados patrones en imágenes (ver [3]).

Definición 1 ([3]). Sea T una t-norma, I una implicación borrosa y N una negación borrosa fuerte. La *transformación morfológica borrosa todo-nada* (FMHMT) de la imagen A en

Este trabajo ha sido financiado por el proyecto TIN2013-42795-P concedido por el Ministerio de Economía y Competitividad. P. Bibiloni es beneficiario de la beca FPI-CAIB con código FPI/1645/2014 de la Conselleria d'Educació, Cultura i Universitats del Govern de les Illes Balears, cofinanciada por fondos europeos.

niveles de gris mediante el elemento estructurante $B = (B_1, B_2)$ en niveles de gris se define, para cada y en el dominio d_A de A , como

$$FMHMT_{T,I,N}(A,B)(y) = T(E_I(A,B_1)(y), E_I(N(A), B_2)(y)),$$

donde $N(A)(x) = N(A(x))$ para todo $x \in d_A$ y E_I es la erosión borrosa dada por

$$E_I(A,B)(y) = \inf_{x \in d_A \cap T_y(d_B)} I(B(x-y), A(x)),$$

donde $x \in T_y(d_B) \Leftrightarrow x - y \in d_B$.

Así, en esta comunicación, se propone un algoritmo basado en el uso del operador FMHMT para la detección de estructuras curvilíneas en imágenes de diversa índole. Para una detección exhaustiva, se usan, de forma similar a [1], elementos estructurantes con forma lineal de distintos tamaños y distintas direcciones ($0^\circ, 45^\circ, 90^\circ, 135^\circ$). Cada uno de estos elementos estructurantes en niveles de gris genera una imagen todo-nada borrosa, que es interpretada como el grado de similitud obtenido como la agregación de los valores de verdad de las inclusiones del elemento estructurante B_1 en la imagen y de B_2 en la imagen complementaria. En este punto, y con el objetivo de conseguir una imagen todo-nada final a partir de esta familia, se propone el uso de dos funciones de agregación. La primera función agrega los valores obtenidos con las transformaciones todo-nada en una misma escala y las diferentes orientaciones. La segunda agrega los valores obtenidos con las diferentes escalas. A continuación, se aplica el algoritmo de adelgazamiento *Non-maxima suppression*, un algoritmo de histéresis y un algoritmo de eliminación de regiones pequeñas para obtener la imagen final.

Una vez definido el algoritmo, se estudia el comportamiento del mismo en función de la t-norma e implicación borrosa usadas en el operador FMHMT, en función del número de escalas consideradas en los elementos estructurantes y en función de las funciones de agregación utilizadas. Finalmente, la comparación visual de los resultados obtenidos en imágenes de diversa naturaleza indica el potencial de este método en la detección de las estructuras curvilíneas, obteniendo en varios casos mejores resultados que los obtenidos por el método presentado en [1].

Referencias

- [1] X. Bai, T. Wang, and F. Zhou. Linear feature detection based on the multi-scale, multi-structuring element, grey-level hit-or-miss transform. *Computers & Electrical Engineering*, 46:487 – 499, 2015.
- [2] P. Bibiloni, M. González-Hidalgo, and S. Massanet. A survey on curvilinear object segmentation in multiple applications. 2016. Sometido a Pattern Recognition.
- [3] M. González-Hidalgo, S. Massanet, A. Mir, and D. Ruiz-Aguilera. A fuzzy morphological hit-or-miss transform for grey-level images: A new approach. *Fuzzy Sets and Systems*, 286:30 – 65, 2016. Theme: Images and Clustering.

FPGA Implementation of the Two-Dimensional Fuzzy-ELA Algorithm for Image Enlargement

M. Brox¹, S. Sánchez-Solano², P. Brox², A. Gersnoviez¹, I. Baturone²

¹ Department of Computer Architecture, University of Córdoba
[mbrox, andresgm]@uco.es

² Instituto de Microelectrónica de Sevilla, IMSE-CNM (CSIC/Universidad de Sevilla)
[santiago, brox, lumi]@imse-cnm.csic.es

Resolution improvement of images is today required for many applications, as medical or satellite imaging, where it is very important to distinguish details [1]. The interpolation capability provided by Fuzzy Logic, in addition to its effectivity to incorporate heuristic knowledge into numeric procedures, has motivated its usage in recent algorithms for image enlargement [2]. This paper presents the hardware implementation of the two-dimensional Fuzzy-ELA algorithm proposed in [3]. The 2D Fuzzy-ELA method is a generalization of the basic Fuzzy-ELA algorithm [4], which uses a fuzzy system to adapt the interpolation to the presence of edges in images, achieving better results than many other approaches. The hardware implementation described in this paper includes parameters to allow selecting different scale factors for the image enlargement. In order to simplify the description, Fig.1a illustrates the process when a factor of two is chosen.

Cells labeled with capital letters correspond to original pixels. New interpolated pixels can belong to a new row (*r* pixel), a new column (*c* pixel), or a new row and a new column simultaneously (*rc* pixel). Let us consider the *r*, *c*, and *rc* pixels shown in the green area of Fig.1a. The *r* and *c* pixels are interpolated using two Fuzzy-ELA blocks. The first one uses as input the three nearest original pixels from the upper and lower rows (pixels D, E, F, G, H, and X in the blue area of the Fig.1a), whilst the second uses the three nearest original pixels from the left and right columns (pixels B, E, H, C, F, and X in the red area of the Fig.1a). Fuzzy-ELA interpolators use the rulebases shown in Table 1 to detect possible edges in the three main directions analyzed in each case, where input values are absolute luminance differences in the considered directions and the terms “VS”, “S”, and “L” are

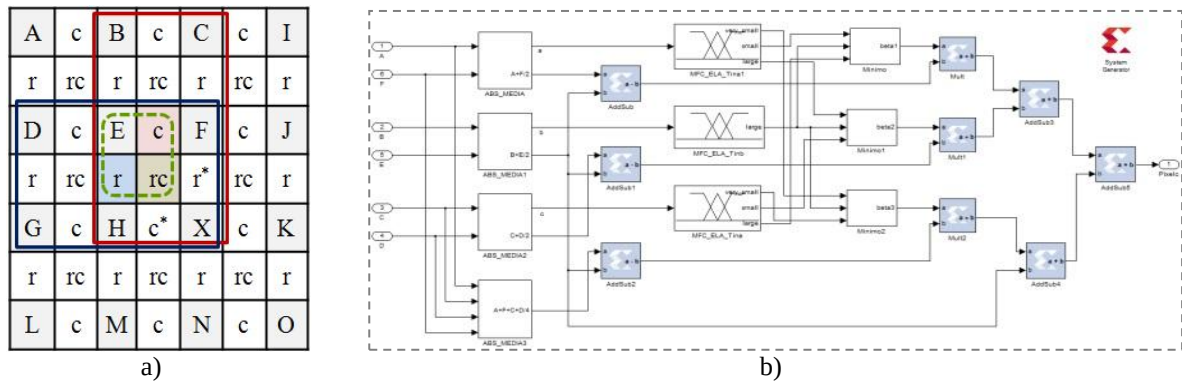


Fig. 1: (a) 2D Fuzzy-ELA algorithm; (b) Hardware implementation of the Fuzzy-ELA block

¹ This work was supported by the projects TEC2014-57971-R and RTC-2014-2932-8 from the Spanish Government (with support from FEDER).

Table 1. Fuzzy rules for r and c pixels.

r pixel		c pixel	
Antecedent	Consequent	Antecedent	Consequent
$ D-X =S \ \& \ E-H =L \ \& \ F-G =L$	$(D+X)/2$	$ B-X =S \ \& \ E-F =L \ \& \ C-H =L$	$(B+X)/2$
$ D-X =L \ \& \ E-H =L \ \& \ F-G =S$	$(F+G)/2$	$ B-X =L \ \& \ E-F =L \ \& \ C-H =S$	$(C+H)/2$
$ D-X =VS \ \& \ E-H =L \ \& \ F-G =VS$	$(D+X+F+G)/4$	$ B-X =VS \ \& \ E-F =L \ \& \ C-H =VS$	$(B+X+C+H)/4$
Otherwise	$(E+H)/2$	Otherwise	$(E+F)/2$

linguistic labels of the fuzzy sets used to represent the concepts “very small”, “small”, and “large”, respectively.

Once r and c pixels have been calculated, the rc pixel is inferred by applying vertical Fuzzy-ELA to pixels (E, r, H, F, r^*, X) and horizontal Fuzzy-ELA to pixels (E, C, F, H, c^*, X) . Values of r^* and c^* cannot be calculated until K and N have been processed. To avoid this inconvenience, values associated to these pixels are obtained as the average of their vertical $[r^*=(F+X)/2]$, and horizontal $[c^*=(H+X)/2]$ neighbors, respectively. The final value assigned to rc is the average of the outputs of the two Fuzzy-ELA interpolators.

As a step previous to its implementation as a configurable IP-module, the 2D Fuzzy-ELA algorithm has been verified on FPGA using a model-based design methodology for rapid development of video and image processing systems [5]. Fig.1b shows the block diagram of one of the four Fuzzy-ELA blocks included in the design (one for r and c pixels, and two for calculating rc pixel). The shapes of the fuzzy sets appearing in Table 1 were adequately selected in order to simplify the evaluation of the system output by means of a Fuzzy-Mean defuzzification method.

The hardware implementation of the algorithm for image size of 210x210 pixels consumes less than 600 slices in Spartan-6 or Zynq-7000 devices. A BRAM block and 32 hardware multipliers are also employed in both FPGAs. The circuit is able to operate over 150 MHz in a Spartan-6, thus allowing to process images with SXGA (1280x1024) and HD 1080P60 (1920x1080) resolutions. In comparison with other traditional interpolation methods or edge-dependent algorithms, the results are better, both quantitatively (RMSE or PSNR) and from a qualitative point of view (visual perception of image quality).

References

- [1] R.C. Gonzalez, R.E. Woods: *Digital image processing*. Prentice Hall, 3rd. edition, 2007.
- [2] H. Tamukoh, N. Suetake, H. Kawano, R. Kubota, B. Cha, T. Aso: “Fuzzy-rule-embedded reduction image construction method for image enlargement with high magnification”. Proceedings of the 9th International Conference on Computer Vision Theory and Applications, (VISAPP 2014), vol. 1, pp. 228-233, Jan. 2014
- [3] P. Brox, I. Baturone, S. Sánchez-Solano: “Image enlargement using the Fuzzy-ELA algorithm”, Proceedings of the Information Processing and Management of Uncertainty in Knowledge-based Systems (IPMU 2006), pp. 866-873, Jul. 2006.
- [4] P. Brox, I. Baturone, S. Sánchez-Solano: *A Fuzzy Edge-Dependent Interpolation Algorithm*, in *Soft Computing in Image Processing: Recent Advances*. Springer Verlag, 2007.
- [5] S. Sánchez-Solano, M. Brox, E. del Toro, P. Brox, I. Baturone: “Model-Based Design Methodology for Rapid Development of Fuzzy Controllers on FPGAs”, *IEEE Transactions on Industrial Informatics*, vol. 9, no. 3, pp. 1361-1370, Aug. 2013.

On Fuzzy Modelling of Shape Vertices

J. Chamorro-Martínez¹, L. Suárez-Llorens¹

¹ *Department of Computer Science and Artificial Intelligence, University of Granada, Spain*
jesus@decsai.ugr.es, luissuarez@correo.ugr.es

Shape analysis is often required for detection and recognition of objects in an image. In this analysis, the selection of interest points, and in particular the vertices, is a very important task; they are useful, for example, to trace the curve or to describe the contour. Nevertheless, the location of this interest points are imprecise, so fuzzy approaches are needed [1, 2].

Let Ω be a connected and bounded image region. Let $C = \partial\Omega$ be the discrete contour of Ω , defined as a cyclically ordered set of points:

$$C = \{\mathbf{p}_i = (x_i, y_i)\}_{1 \leq i \leq n} \quad (1)$$

where successive boundary pixels in Ω are arranged in a cyclic (clockwise) order. We shall extend the notation to any natural number i by assuming $p_i = p_{i \bmod n}$. Let $\mathcal{C}$ be the set of all discrete contours.

Let $S_{ij}^C = \{\mathbf{p}_k \in C\}_{i \leq k \leq j}$ be the subset of C representing the segment between $\mathbf{p}_i$ and $\mathbf{p}_j$, and let $\Theta^C = \{S_{ij}^C\}_{i \neq j}$ the set of all the segments of C . Let l_S be a line estimated by means of a line fitting process from the set of points $S \in \Theta^C$. We model the linearity of a segment $S \in \Theta^C$ by means of a fuzzy set $\tilde{L}$:

$$\tilde{L} : \Theta^C \rightarrow [0, 1] \quad (2)$$

where the membership degree is defined on the basis of the coefficient of determination as:

$$\tilde{L}(S_{ij}^C) = 1 - \left(\frac{\sum_{k=i}^j \|\mathbf{p}_k - \hat{\mathbf{p}}_k\|}{\sum_{k=i}^j \|\mathbf{p}_k - \bar{\mathbf{p}}\|} \right)^q \quad (3)$$

where $\|\cdot\|$ is the Euclidean norm, $\bar{\mathbf{p}}$ is the mean of the points in S_{ij}^C , $\hat{\mathbf{p}}_k$ is the projection of $\mathbf{p}_k$ on $l_{S_{ij}^C}$, and q is a range adjustment parameter (fixed to 3 in this paper). The numerator of the ratio represents the sum of squared errors of the fitted model, while the denominator can be thought of as the sum of errors from the null model, that is, a model where each $\mathbf{p}_k$ value is predicted to be the mean of the $\mathbf{p}_k$ values (without having any additional information, the mean would be the best guess if our aim is to minimize the squared differences). Therefore, the ratio is indicative of the degree to which the model parameters improve upon the prediction of the null model. The smaller this ratio, the higher the degree $\tilde{L}(S_{ij}^C)$.

A vertex is a point where two or more lines meet, so the linearity is expected to be high both in their left and right sides. In addition, the lines meeting in a vertex form an angle, so

Part of this research was funded by the Spanish Ministry of Economy and Competitiveness and the European Regional Development Fund (FEDER) under project grant TIN2014-58227-P

the linearity is expected to be low in any segment centered around it. Taking into account this idea, we model the verticity of a point $p_i \in C$ by means of a fuzzy set $\tilde{V}$:

$$\tilde{V} : \mathcal{C} \times \mathbb{N} \rightarrow [0, 1] \quad (4)$$

where the membership degree is defined as:

$$\tilde{V}(C, i) = \otimes \{ \tilde{L}(S_{i-w, i}^C), \tilde{L}(S_{i, i+w}^C), 1 - h_{vvery}(\tilde{L}(S_{i-w, i+w}^C)) \} \quad (5)$$

with $\otimes$ being a t-norm (in this paper, the minimum has been used), h_{vvery} being a linguistic modifier modeled in this paper by $h_{vvery}(x) = x^4$, and w being the size of a neighbourhood around p_i (in this paper, $w = 0.1 * n$). The semantics of this expression is the following: a point in a contour C is a vertex to the extent that the points in both its left and right neighbourhoods match a line, and the union of both neighborhoods do not strictly match a line, the condition “strict” being modelled by the linguistic modifier h_{vvery} .

Let us define $\tilde{V}_C : C \rightarrow [0, 1]$ with $\tilde{V}_C(p_i) = \tilde{V}(C, i)$. Let $\tilde{V}^*(C)$ be the fuzzy subset that includes the local maxima of $\tilde{V}(C, i)$ defined as:

$$\tilde{V}^*(C) = \left(\arg \max_{p_i \in S_{i-w, i+w}} \tilde{V}(p_i) \right) \cap \tilde{V}_C \quad (6)$$

Figure 1 shows an example with several shapes representing a transition from a square to a triangle. The membership degree is drawn using a gradient color from red (associated to degree 1) to black (degree 0). In addition, the local maxima in verticity are marked with a cross. In Figure 1(a) the four vertices has a degree of 1; in Figures 1(b)-(d) the membership degree of the top-left vertex decreases as the square becomes a triangle (with values 1.0, 0.8, 0.55, 0.3 and 0 respectively), while the others keep the value 1.0. If, as first approximation, we use the sigma-count to count the number of vertices in $\tilde{V}^*$, and we use this data to obtain a membership degree to the fuzzy sets “Square” and “Triangle” (with triangular membership function centered on 4 and 3, respectively) we obtain the following possibility distributions: (a) 1.0/Sq, (b) 0.8/Sq + 0.2/Tr, (c) 0.55/Sq + 0.45/Tr, (d) 0.3/Sq + 0.7/Tr and (e) 1.0/Tr.

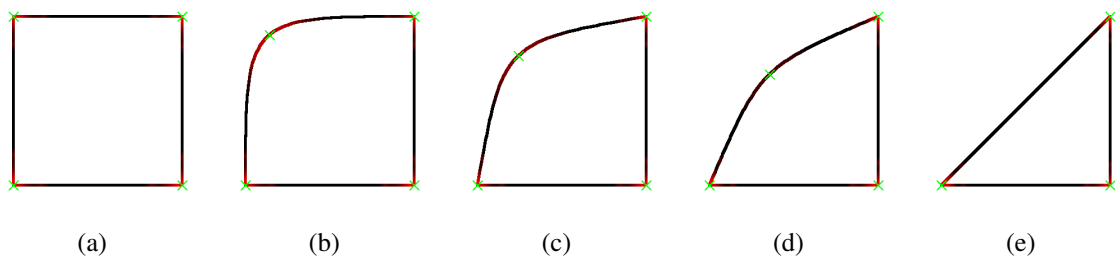


Figure 1: Shapes representing a transition from a square to a triangle.

References

- [1] B. Lazzerini, F. Marcelloni: *A Fuzzy Approach to 2-D Shape Recognition*. IEEE Transaction on Fuzzy Systems, Vol.9, No.1 (2001) 5–16.
- [2] G. Castellano, A.M. Fanelli, M.A. Torsello: *Shape annotation by semi-supervised fuzzy clustering*. Information Sciences No.289 (2014) 148–161.

Face Recognition Using Color Local Binary Patterns: A Comparative Study

F. Dornaika^{1,2}, I. Gonzalez³, and Y. Ruichek³

¹ UPV/EHU, Manuel Lardizabal, 1, 20018 San Sebastian, Spain

² IKERBASQUE, Basque Foundation for Science, Bilbao, , Spain
fadi.dornaika@ehu.es

³ University of Technology of Belfort-Montbéliard, Belfort, France
ignacio.gonzales@utbm.fr, yassine.ruichek@utbm.fr

In order to tackle image classification tasks, several variants of Local Binary Patterns (LBP) have been proposed in the literature [1, 2, 3]. In this work, we mainly study the performance of color Local Binary Patterns variants when applied to the problem of face recognition. In particular, we compare the Opposite Color LBP with the recent Quaternionic Local Binary Pattern.

Image descriptors: The basic LBP descriptor was introduced for texture classification in digital images. It is calculated on an area of 3×3 pixels. Every pixel in the original image is replaced by a sequence of 8 bits. Each bit value (0 or 1) is the result of comparing the central pixel value to that of its neighbors. LBP descriptors can be given by the image of binary codes or by their histograms. LBP descriptors are invariant under a monotonic change in the value of the pixels. Since its introduction in the 90s, many LBP descriptors have been proposed. We give three main LBP descriptors.

Complete LBP (CLBP) [1]: This descriptor is an improvement of the classic LBP. In a Complete LBP three maps of codes are generated for coding an image: (i) CLBP_C for central pixels, (ii) CLBP_S (Sign) which is the classic LBP, and (iii) CLBP_M for coding the magnitude differences.

Opposite Color LBP (OC-LBP) [2]: In OC-LBP, the operator is used in each color channel independently. This yields three maps. For color channel pairs, the central pixel is taken from a neighboring channel and the neighboring pixels from other channel. This yields six maps. Thus, a total of nine maps or histograms are obtained (R, G, B, RG, RB, GB, GR, BR, RG), making the descriptor slower than the grayscale histogram.

Quaternionic Local Binary Pattern (QLBP) [3]: This is a new model proposed for the re-identification of persons (person recognition by non-overlapping cameras). First, color pixels are represented by imaginary quaternions. Pseudo quaternion rotation is then applied on the color image using several reference quaternions. The phases of the resulting operation are then fed to a classic LBP operator.

Dataset: The image database "Essex University Facial Recognition" is used to perform face recognition with different LBP descriptors. This dataset is divided into four directories

Acknowledgements must be included in a footnote.

<http://cswwww.essex.ac.uk/mv/allfaces/index.html>

(faces94, faces95, faces96, and face) in order of increasing difficulty. Faces96 and face are the most difficult to different reasons: background variation, scale variation and extreme expressions. Overall, the characteristics are as follows: Total number of persons: 393, Images per person: 20, Total number of images: 7860 color 180 x 200 images

Performance evaluation: The evaluation protocol is as follows. First, the descriptor in hand is extracted from all images. Second, the whole set of images is randomly split into a training part, and a testing part (this process is repeated ten times). Third, we evaluate the recognition rate over the test images using the 1-NN classifier adopting Euclidean distance or χ^2 distances, and Support Vector Machines. We consider several percentages for the training part of data. Table 1 summarizes the average recognition rates for several LBP descriptors and for several classifiers. The number in parenthesis denotes the number of maps used by the descriptor. The upper part of the table corresponds to the use of histograms and the lower part to the use of maps.

We can conclude that, in general, the results obtained by the OC-LBP descriptor were better than those obtained by the QLBP. Both descriptors were superior to grey-level CLBP (obtained with NN classifier). For χ^2 distance, the performance obtained with LBP histograms was better than the one provided by LBP images/maps. However, when Euclidean distance or SVM is used the use of maps provided superior results in general. Future work may investigate the estimation of optimal reference quaternions for QLBP descriptors.

Table 1: Recognition performance (%) using several descriptors and several classifiers.

Descriptor \ Training size	20%			30%			40%		
CLBP	83.2			88.8			91.6		
Histogram \ Classifier	NN	χ^2	SVM	NN	χ^2	SVM	NN	χ^2	SVM
QLBP (3)	86.5	87.8	67.7	89.3	92.1	78.9	92.4	94.7	87.0
OC-LBP (3)	87.8	95.7	66.9	93.6	97.7	77.9	95.4	99.2	85.2
QLBP (6)	87.0	90.6	70.2	91.1	93.6	82.7	93.4	95.1	89.8
OC-LBP (6)	92.9	97.5	68.7	96.4	99.2	81.2	97.7	99.7	88.8
Map \ Classifier	NN	χ^2	SVM	NN	χ^2	SVM	NN	χ^2	SVM
QLBP (3)	88.5	88.0	75.3	92.4	91.3	85.2	94.5	94.4	88.8
OC-LBP (3)	89.6	89.8	87.3	94.4	94.1	92.1	95.4	96.1	94.8
QLBP (6)	89.0	88.6	75.6	93.6	91.3	85.5	96.4	95.9	91.1
OC-LBP (6)	89.8	91.9	88.0	94.9	94.4	93.1	96.2	97.2	95.2

References

- [1] Z. Guo, L. Zhang, and D. Zhang: *A completed modeling of local binary pattern operator for texture classification* IEEE Trans. Image Process. **19-6** (2010) 1657–1663.
- [2] T. Maenpaa: *The local binary pattern approach to texture analysis: extensions and applications*. Ph.D. thesis, University of Oulu. (2003).
- [3] R. Lan and Y. Zhou and Y. Tang and C. Chen: *Person re-identification using quaternionic local binary pattern*. IEEE Int. Conf. on Multimedia and Expo (2014) 1–6.

Agregación de imágenes de contornos usando OWAs

Manuel González-Hidalgo, Sebastia Massanet, Arnau Mir, Daniel Ruiz-Aguilera

*Departament de Matemàtiques i Informàtica, Universitat de les Illes Balears
Carretera de Valldemossa, km 7.5, E-07122, Palma (Illes Balears), España
{manuel.gonzalez,s.massanet,arnau.mir,daniel.ruiz}@uib.es*

La detección de contornos es un paso fundamental en muchas tareas de procesamiento de imágenes de alto nivel. Esto conlleva la propuesta constante en últimos años de nuevos detectores de contornos. Es sabido que no hay un detector de contornos óptimo para todo tipo de imágenes y, además, no es una misión sencilla establecer el conjunto óptimo de parámetros de un detector de contornos concreto. Así, en este trabajo, para hacer frente a estos inconvenientes, se presenta un algoritmo basado en la obtención de una imagen de contornos de consenso a partir de las imágenes de contornos obtenidas por varios detectores de contornos o configuraciones del mismo detector. El proceso es similar al descrito en [1] con métodos de reducción de ruido. En nuestro caso, las imágenes de contornos de entrada se combinan utilizando conjuntos borrosos multidimensionales y funciones de agregación compensatorias, aplicando posteriormente una función de penalización, para obtener la imagen de contornos de consenso. Los resultados experimentales muestran que este enfoque conduce a un detector de contornos robusto que es capaz de obtener resultados notables para diferentes tipos de imágenes. A continuación se detallan los pasos del algoritmo:

Descripción del algoritmo:

Para una imagen dada I de la cual queremos determinar los contornos:

- **Paso 1** [*obtención de contornos*]. Aplicamos n detectores de contornos diferentes, siendo $I_1, I_2, \dots, I_n$ sus respectivas imágenes de contornos obtenidas. Las representamos por medio de un conjunto borroso multidimensional.
- **Paso 2** [*Agregación usando OWAs*]. En cada píxel (i, j) , se aplican q funciones de agregación compensatorias (concretamente OWAs) diferentes $F_1, F_2, \dots, F_q$ aplicadas a los valores $I_1(i, j), I_2(i, j), \dots, I_n(i, j)$. Obtenemos, para cada píxel, q valores diferentes $Y = \{y_1, \dots, y_q\}$ que han de ser considerados.
- **Paso 3** [*Obtención de la imagen consensuada*]. Por cada y_k , se calcula su nivel de desacuerdo con respecto a los valores obtenidos por los otros métodos, mediante una medida de desacuerdo D . Elegimos el valor consensuado para el píxel (i, j) como aquél que da el mínimo valor de desacuerdo:

$$E(i, j) = \arg \min_{y \in Y} (D(y, (I_1(i, j), \dots, I_n(i, j)))),$$

Trabajo financiado por el proyecto TIN2013-42795-P.

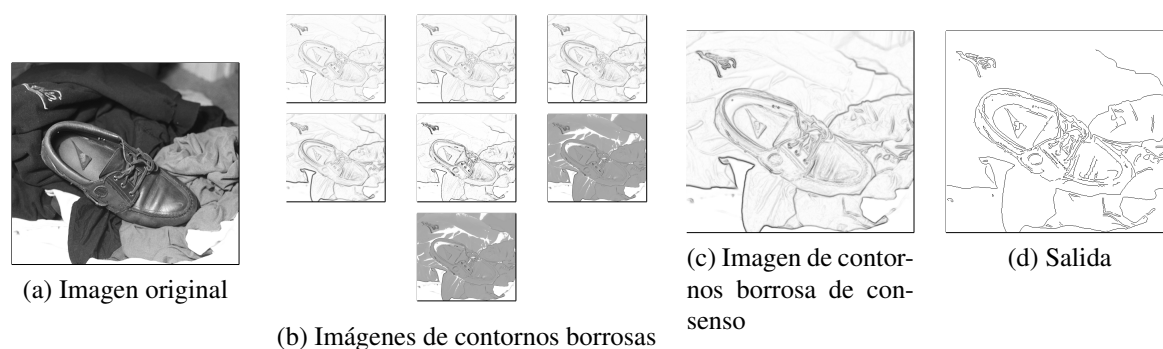


Figura 1: Ilustración de los pasos del algoritmo propuesto.

En este trabajo hemos considerado 7 detectores de contornos. Concretamente, el método de Canny para $\sigma \in \{0.5, 1.0, 1.5, 2.0\}$, y los métodos descritos en [2, 3, 4]. Hemos aplicado 3 operadores OWAs, calculando sus pesos usando cuantificadores lingüísticos borrosos de Yager, y como función de desacuerdo hemos utilizado el error cuadrático medio.

Se ha realizado también una comparación objetiva de los resultados obtenidos utilizando tres medidas objetivas, como son la medida de Pratt, el ρ -coeficiente y la F -medida, para valorar la similitud entre los contornos detectados y la *Ground-Truth* de la base de datos. Se ha utilizado la base de datos de la *University of South Florida*. El test estadístico no paramétrico de Wilcoxon demuestra que existen diferencias significativas en el 40% de las comparaciones entre el método de consenso propuesto y los 7 métodos de detección de contornos considerados, usando las tres medidas objetivas anteriores. Además, en el resto de casos, el método de consenso obtiene una media superior en las medidas a las de los otros métodos. Finalmente, la desviación típica de las medidas objetivas usando el método de consenso es la más reducida, mostrando que el método de consenso es el más robusto entre los considerados.

Referencias

- [1] L. González Jaime, E. E. Kerre, M. Nachtegaele, and H. Bustince, *Consensus image method for unknown noise removal*. Knowledge-Based Systems, V. 70 (2014), 64 – 77.
- [2] M. González-Hidalgo, S. Massanet, A. Mir, and D. Ruiz-Aguilera, *On the generalization of the fuzzy morphological operators for edge detection*. Proceedings of IFSA-EUSFLAT 2015, Atlantis Press. 1082–1089.
- [3] M. González-Hidalgo, S. Massanet, A. Mir, and D. Ruiz-Aguilera, *On the generalization of the uninorm morphological gradient*. Proceedings of IWANN 2015, Part II. Springer International Publishing, 2015. 436–449.
- [4] M. González-Hidalgo, S. Massanet, A. Mir, and D. Ruiz-Aguilera. *A new edge detector based on uninorms*. Proceedings of IPMU 2014, Part II. Springer International Publishing, 2014, 184–193.

Una metodología para evaluar algoritmos de segmentación jerárquica

Carely Guada¹, J. Tinguaro Rodríguez¹, Daniel Gómez²,
Javier Yáñez¹, Javier Montero^{1,3}

¹ Facultad de Ciencias Matemáticas, Universidad Complutense de Madrid, 28040 Madrid
{cguada, jtrodrig, jayage, jamonter}@ucm.es

² Facultad de Estudios Estadísticos, Universidad Complutense de Madrid, 28040 Madrid
dagomez@estad.ucm.es

³ Instituto de Geociencias (CSIC, UCM), Plaza de Ciencias 3, 28040 Madrid

Este artículo propone un método para evaluar procedimientos de segmentación jerárquica, con la finalidad de permitir comparar distintos algoritmos jerárquicos entre sí, y a su vez, con otras técnicas de segmentación no jerárquicas. La segmentación de imágenes es actualmente una herramienta muy utilizada en las aplicaciones de procesamiento de imágenes, consiste en un proceso que divide una imagen en múltiples regiones u objetos. Existen una gran variedad de técnicas de segmentación de imágenes [6], así como técnicas para evaluar su desempeño [7]. Dentro de las técnicas de segmentación de imágenes está la segmentación jerárquica, la cual produce una serie de particiones o una secuencia consistente de segmentaciones, identificando los objetos jerárquicamente en una imagen con distintos niveles de detalle, es decir, genera desde una sola partición que contiene un solo objeto y a medida que aumenta el número de particiones va aumentando el número de objetos identificados. Así, la segmentación jerárquica es una extensión relevante de la segmentación, puesto que varias aplicaciones requieren diferentes niveles de detalle [2], de tal manera, que la identificación de objetos es consistente a través de los diferentes niveles de detalle.

No obstante, a diferencia de la segmentación no jerárquica, existen relativamente pocos procedimientos de segmentación jerárquica de imágenes [1,3,5], y éstos usualmente encuentran el problema de que no existen técnicas claramente aceptadas en cuanto a cómo evaluar estos algoritmos de segmentación jerárquica. Así, en este trabajo se propone un método que permite evaluar dichos procedimientos de segmentación jerárquica mediante una técnica de evaluación supervisada de segmentación basada en bordes, debido a que ésta es una de las metodologías de evaluación más utilizada y aceptada actualmente [1].

En más detalle, el método propuesto se construye considerando como referencia un conjunto de segmentaciones realizadas por humanos. Así, dada una segmentación humana y una secuencia jerárquica de segmentaciones por medio de un algoritmo, existe un elemento de la secuencia que mejor aproxima el nivel de detalle obtenido de un humano de referencia. Por lo tanto, debido a que una secuencia jerárquica de segmentaciones se aproxima a la población de diferentes niveles de detalle de segmentación humana, esas segmentaciones de la secuencia servirán como base para la evaluación de la secuencia jerárquica como un todo.

El proceso de evaluación se realiza utilizando las curvas *precision-recall* (PR) la cual es una técnica de evaluación estándar y apropiada para técnicas de segmentación que consideran un conjunto de parámetros. Para ello, sea un conjunto L de segmentaciones humanas independientes GT^l ($1 \leq l \leq L$) como referencia, y sea un conjunto de K segmentaciones S^k ($1 \leq k \leq K$) organizadas jerárquicamente. Dada una segmentación

humana GT^l , éstas se emparejan sucesivamente contra cada segmentación generada por el algoritmo jerárquico, obteniendo así el conteo de los *verdaderos positivos* (TP_{kl}), *falsos positivos* (FP_{kl}) y *falsos negativos* (FN_{kl}), con los cuales se calculan los valores de *precision* y *recall*:

$$P_{kl} = \frac{TP_{kl}}{TP_{kl} + FP_{kl}}, \quad R_{kl} = \frac{TP_{kl}}{TP_{kl} + FN_{kl}}$$

A su vez, estos valores se utilizan para calcular el valor F_{kl} para cada comparación entre la segmentación humana y del algoritmo, el cual representa la proporción compuesta que relacionan esos dos valores, calculado como la media armónica entre ellos:

$$F_{kl} = 2 \cdot \frac{P_{kl} \cdot R_{kl}}{P_{kl} + R_{kl}}$$

Esta medida indica lo bien que una segmentación de la secuencia jerárquica se aproxima a los bordes y al nivel de detalle dados por la segmentación humana. Por lo que, para cada imagen de referencia humana GT^l , se selecciona la segmentación $S^{k_l^*}$ de la secuencia $S^1, \dots, S^K$ con el valor de F_{kl} más alto, esto es $k_l^* = \arg \max_{1 \leq k \leq K} F_{kl}$. Entonces, a partir de las segmentaciones $S^{k_l^*}$ que mejor capturan cada segmentación humana GT^l , se obtienen los valores de *precision* (P) y *recall* (R) globales para la imagen considerada, con las cuales se calcula el valor F -global para medir el rendimiento global del procedimiento de segmentación jerárquica evaluado en una imagen de referencia.

$$P = \frac{1}{L} \sum_{l=1}^L P_{k_l^* l} \quad R = \frac{1}{L} \sum_{l=1}^L R_{k_l^* l}$$

Además, en este trabajo se muestran resultados computacionales para ilustrar la viabilidad del método propuesto, utilizando el algoritmo jerárquico *Divide-and-Link* [4] y otros algoritmos de segmentación no jerárquica (por ejemplo *Canny*, *Sobel* y *Prewitt* [6]).

References

- [1] Arbelaez, P., Maire, M., Fowlkes, C., Malik, J.: *Contour detection and hierarchical image segmentation*. IEEE Trans. On Pattern Analysis and Machine Intelligence. **33** (2011).
- [2] Castillo-Ortega, R., Chamorro-Martínez, J., Marín, N., Sánchez, D., Soto-Hidalgo, J.M.: *Describing Images Via Linguistic Features and Hierarchical Segmentation*. Procs. Of WCCI 2010 IEEE World Congress on Computational Intelligence. (2010) 1104–1111.
- [3] Cheng, H., Sun, Y.: *A hierarchical approach to color image segmentation using homogeneity*. IEEE Trans. on Image Processing. **9** (2000) 12:2071-2082.
- [4] Gómez, D., Zarrazola, E., Yáñez, J., Rodríguez, J., Montero, J.: *A Divide-and-Link Algorithm for Hierarchical Clustering in Networks*. Information Sciences. **316** (2015) 308-328.
- [5] Gómez, D., Zarrazola, E., Yáñez, J., Rodríguez, J., Montero, J.: *Fuzzy Image Segmentation based upon Hierarchical Clustering*. Knowledge-Based Systems. **87** (2015) 26-37.
- [6] Guada, C., Gómez, D., Rodríguez, J., Yáñez, J., Montero, J.: *Classifying image analysis techniques from their output*. IJCIS (in press).
- [7] Martin, D., Malik, J., Fowlkes, C.: *Learning to detect natural image boundaries using local brightness, color and textures cues*. IEEE Trans. On Pattern Analysis and Machine Intelligence. **26** (2004) 5:530-549.

Esta investigación ha sido parcialmente financiada por el Gobierno de España, proyecto TIN2015-66471-P, y por el Gobierno de la Comunidad de Madrid, proyecto S2013/ICE-2845 (CASI-CAM).

Reducción del efecto vignetting usando medidas de borrosidad

Laura Lopez-Fuentes^{1,2,3}, Sebastia Massanet¹, Manuel González-Hidalgo¹

¹ *Universitat de les Illes Balears, Dept. Matemàtiques i Informàtica
Ctra. Valldemossa km 7.5, 07122 Palma, Illes Balears, Espanya*

² *AnsuR Technologies, Martin Linges Vei 25, 1364 Fornebu, Noruega*

³ *Computer Vision Centre Barcelona, Campus UAB, 08193 Bellaterra (Barcelona), Espanya
{l.lopez,s.massanet,manuel.gonzalez}@uib.es*

El vignetting es un efecto no deseado en imágenes digitales que debe ser corregido como un primer paso en muchas aplicaciones de visión por computador. Se basa en el decaimiento radial de la luminosidad desde el centro óptico de la imagen hacia los vértices de la misma. Aunque existen diversas causas que provocan este fenómeno, la mayoría de los algoritmos diseñados para reducir este efecto se dedican a la corrección de los conocidos como vignetting natural y vignetting de píxel. Mientras que el primero es debido a como la luz alcanza localizaciones distintas en el sensor de la cámara, el segundo sólo afecta a las cámaras digitales y es provocado por la dependencia angular de los sensores digitales.

Se han propuesto diversos algoritmos para reducir el vignetting en imágenes digitales. Estos algoritmos, a partir de una única imagen de una cámara no predeterminada, deben ser capaces de diferenciar las variaciones globales de luminosidad causadas por el vignetting de aquellas causadas por la presencia de textura o luz. En particular, en un artículo no publicado [3], se propuso un método de corrección de vignetting basado en la minimización de la entropía de la intensidad logarítmica a partir de la observación de que la adición de vignetting en una imagen aumenta la entropía de la misma. Sin embargo, en algunas imágenes, este método obtenía resultados no deseados debido a unas deficiencias matemáticas que fueron subsanadas en [1], donde se propuso una nueva versión del método resolviendo los problemas matemáticos e introduciendo una mejora basada en la búsqueda del centro óptico de la imagen. Esta nueva versión del método se demostró superior a la original obteniendo resultados competitivos con otros métodos existentes en la literatura (ver [1]).

El objetivo de esta comunicación es la propuesta de un nuevo algoritmo de reducción de vignetting en imágenes digitales a partir de la maximización de la medidas de borrosidad de la imagen. Las medidas de borrosidad fueron introducidas en [2] e intentan determinar el grado de borrosidad de un conjunto borroso. Las medidas de borrosidad han sido utilizadas en el campo del procesamiento de imágenes y especialmente, en la segmentación de imágenes. Estos algoritmos suelen basarse en una minimización de la medida de borrosidad como criterio para la obtención del mejor umbral para la segmentación. En cambio, la adición de vignetting en una imagen tiende a disminuir la medida de borrosidad de una imagen

Este trabajo ha sido financiado por el proyecto TIN2013-42795-P concedido por el Ministerio de Economía y Competitividad. L. Lopez-Fuentes es beneficiaria de la beca NAERINGSPHD del Norwegian Research Council para la realización de un doctorado industrial en AnsuR Technologies bajo el acuerdo de colaboración Ref.3114 con la Universitat de les Illes Balears.

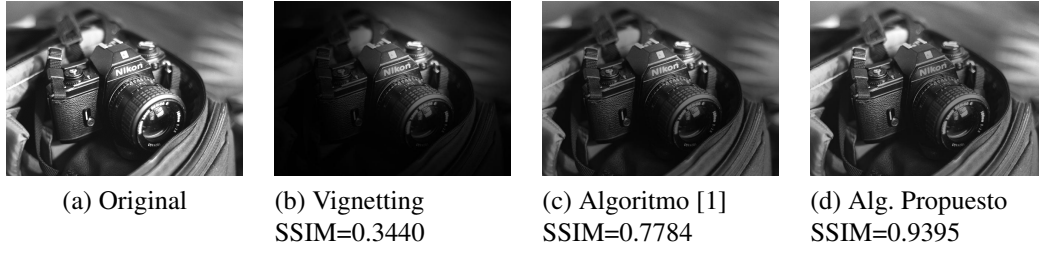


Figura 1: Comparación visual de los resultados obtenidos por el algoritmo propuesto y por el introducido en [1]. Se incluyen los valores obtenidos de la medida SSIM.

y por lo tanto, la maximización bajo ciertas restricciones de esta medida se demuestra útil para la reducción de vignetting. Así, tras encontrar el centro óptico de la imagen, se propone corregir la imagen con vignetting mediante una función de ganancia polinomial de grado seis dada por $g_{a,b,c}(r) = 1 + ar^2 + br^4 + cr^6$ donde r es la distancia del centro óptico de la imagen al píxel tratado, y en la que los valores a , b y c se determinan mediante la maximización, bajo la restricción $g'_{a,b,c}(r) > 0$ para todo $r \in (0, 1)$, de una medida de borrosidad de la imagen corregida $\text{Imagen corregida} = \text{Imagen vignetting} \cdot g_{a,b,c}(r)$. Concretamente, se estudia el rendimiento del algoritmo para dos medidas de borrosidad como son los índices lineal y cuadrático de borrosidad dados por

$$\gamma_p = \frac{2}{MN} \left(\sum_{m=1}^M \sum_{n=1}^N (\min\{\mu(I(m,n)), 1 - \mu(I(m,n))\})^p \right)^{\frac{1}{p}}$$

con $p = 1$ y $p = 2$, respectivamente, donde I es una imagen $M \times N$ y μ es la función de pertenencia sigmoidea de Zadeh. Finalmente, se lleva a cabo una comparación objetiva del rendimiento de ambas configuraciones con respecto al algoritmo propuesto en [1] basada en el uso de la medida SSIM (ver Figura 1) para 15 imágenes de la base de imágenes de la *University of South Florida* a las que se les ha añadido vignetting artificial de forma aleatoria. Los resultados muestran el potencial del método en la reducción del vignetting, especialmente cuando la intensidad del mismo es muy elevada.

Referencias

- [1] L. Lopez-Fuentes, G. Oliver, and S. Massanet. Revisiting image vignetting correction by constrained minimization of log-intensity entropy. In I. Rojas, G. Joya, and A. Catala, editors, *Advances in Computational Intelligence. Proc. of IWANN 2015. Part II*, pages 450–463. Springer International Publishing, 2015.
- [2] A. D. Luca and S. Termini. A definition of a nonprobabilistic entropy in the setting of fuzzy sets theory. *Information and Control*, 20(4):301 – 312, 1972.
- [3] T. Rohlfing. Single-image vignetting correction by constrained minimization of log-intensity entropy. 2012. No publicado. Disponible en: <http://citeseerx.ist.psu.edu/viewdoc/summary?doi=10.1.1.258.4780>.

Membrane proteins Detection Using Unsupervised Fuzzy Competitive Learning

Mohammed Madiafi^{1,2}, Abdelaziz Bouroumi²

¹ *National School of Applied Sciences, Safi
Cadi Ayyad University, Morocco*

m.madiafi@uca.ma, madiafi.med@gmail.com

² *Information Processing Laboratory, Ben M'sik Faculty of Sciences
Hassan II University of Casablanca, Morocco*
a.bouroumi@gmail.com

Membrane proteins play important roles in the functioning and communication of biological cells. They represent about 30% of the proteome and are the target of many medical drugs, but only a small number of their structures have been determined. Hence the need for efficient methods for detecting these proteins that are hard to recognize due to their localisation inside cell membranes. In this paper we present a method for detecting membrane proteins by analyzing transmission electron microscopy, or TEM, images. The method is suited to non-stacked membrane proteins that are generally harder to detect than stacked ones, because of the heterogeneity of gray levels and the very low signal to noise ratio that characterize TEM images [1]. It is based on a two-layer competitive neural network [2] whose synaptic connections weights are learned from training data using a dynamic optimal training algorithm (DOTCNN) [3]. The learning data are iteratively presented to the input layer in the form of p -dimensional object vectors $x_k, k = 1, 2, \dots, n$. The c output units that each represents a cluster compete for each input datum. For this, the membership degree of the k^{th} object to the i^{th} cluster at iteration t is first calculated using the expression:

$$u_{ik,t} = \frac{\|x_k - w_{i,t}\|^{-\alpha}}{\sum_{j=1}^c \|x_k - w_{j,t}\|^{-\alpha}} \quad (1)$$

where α is a user defined parameter and $w_{i,t}$ is the vector of synaptic weights characterizing the connections between the i^{th} output unit and the p input units. The c vectors $w_1, w_2, \dots, w_c$ model the prototypes or code vectors of the c clusters, and the ultimate goal of the learning process is to find optimal values for the components of these vectors, i.e., components that minimize the global error defined by:

$$E_{k,t} = \sum_{i=1}^c u_{ik,t-1}^m \|x_k - w_{i,t-1}\|^\alpha \quad (2)$$

where $m > 1$ is a parameter for controlling the fuzziness degree of the c -partition.

We solve this search and optimization problem by starting from randomly chosen weights $w_{rq}, r = 1, 2, \dots, c$ and $q = 1, 2, \dots, p$ and applying the gradient descent rule in order to iteratively adapt them so that the process converges to a minimum of the error criterion. This leads to the learning rule:

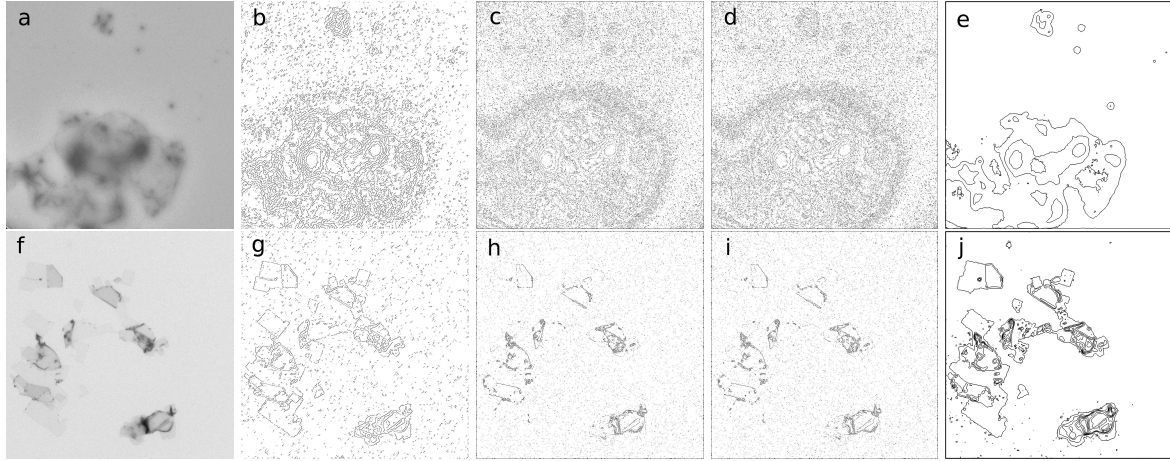


Figure 1: Edge detection results obtained for two original TEM images (a and f) by the proposed method (e and j) and three other methods: Canny (b and g), Prewitt (c and h), and Sobel (d and i).

$$w_{rq,t} = w_{rq,t-1} + \eta_t \alpha u_{rk,t-1}^2 (x_{kq} - w_{rq}) \|x_k - w_{r,t-1}\|^{\alpha-2} [(1-m)u_{rk,t-1}^{m-2} + \sum_{i=1}^c m u_{ik,t-1}^{m-1}] \quad (3)$$

where η_t denotes the learning rate at iteration t .

DOTCNN is a fuzzy method because more than one can win the competition. Its dynamic aspect is modeled using a threshold ξ_t that fixes, for each iteration t , the output neurons that can compete for winning the current input object and benefit from the adjustment of their synaptic weights. Hence, only neurons that satisfy the condition:

$$u_{ik,t} > \xi_t \quad (4)$$

are qualified to compete. As the competition difficulty increases with iterations, this condition prevents unnecessary adjustment of cluster prototypes that are very far from the current input object. As an illustration, we present in Figure 1 the edge detection results obtained by DOTCNN and three other methods for two test examples of TEM images. A detailed description of these results, as well as a more complete description of the proposed method, will be provided in the full paper.

Main references

- [1] N. Coudray, F. Beck, J.-L. Buessler, A. Korinek, A. Karathanou, H. Rmigy, H. Kihl, A. Engel, J.M. Plitzko, J.-P. Urban: *Automatic acquisition and image analysis of 2d crystals*. Microscopy Today. Vol 16 **No 1** (2008) 48-49.
- [2] M. Madiafi, A. Bouroumi: *A New Fuzzy Learning Scheme for Competitive Neural Networks*. Applied Mathematical Sciences. Vol 6 **No 63** (2012) 3133-3144.
- [3] M. Madiafi, A. Bouroumi: *Dynamic Optimal Training for Competitive Neural Networks*. Computing and Informatics. Vol 33 **No 2** (2014) 237-258.

Modelado lingüístico de relaciones espacio-temporales de los vectores de movimiento H264/AVC para la detección de eventos en vídeos de tráfico.

Luis Rodriguez-Benitez¹, Juan Giralt¹, Juan Moreno-Garcia¹, Luis Jimenez¹

¹ *Departamento de Tecnologías y Sistemas de Información, Univ. de Castilla-La Mancha,*
 luis.rodriguez@uclm.es, juan.giralt@uclm.es, juan.moreno@uclm.es,
 luis.jimenez@uclm.es

Algunos sistemas de control de tráfico on-board deben de trabajar en tiempo real y analizar la información procedente de la señal de vídeo de manera continua adaptándose a los cambios constantes que ocurren [1, 3]. Estos cambios pueden ser debidos tanto a la propia naturaleza de estas grabaciones en las que no hay un escenario fijo sino que este cambia continuamente o a otro tipo de circunstancias como pueden ser cambios en la velocidad del vehículo desde el que se captura el vídeo, entre otros. Por todo esto, cualquier aproximación que pretenda extraer conocimiento de estos vídeos debe tener entre sus propiedades principales la de la adaptabilidad. El objetivo de este trabajo es la detección de eventos en tiempo real procesando exclusivamente los vectores de movimiento H264/AVC [2] obtenidos directamente del propio codificador. Una vez detectado el evento se va a caracterizar el mismo utilizándose para ello una representación lingüística [5]. Esta caracterización se realiza mediante la utilización de conjuntos difusos a los que hay que dotar de la propiedad de la adaptabilidad, que les permitiría conseguir una caracterización del evento lo más precisa posible. Por ejemplo, si se pretende detectar un adelantamiento, una vez identificado ese tipo concreto de maniobra, ésta se puede caracterizar para conocer si se ha realizado a baja velocidad, velocidad normal, alta o muy alta, por ejemplo.

Como base de la detección de eventos se van a estudiar los cambios en los módulos de los vectores de movimiento. El análisis de los módulos no se realiza de manera global para cada imagen del vídeo (frame) sino que éste se divide en una serie de áreas o zonas con un conjunto de características diferenciadas. Es decir, los vectores de movimiento presentes en vídeos capturados desde un vehículo en movimiento tienen valores diferentes en función de la zona o posición en la que se encuentran en el frame. En función de este comportamiento podríamos dividir el frame en diferentes áreas y esta decisión también vendrá determinada por el tipo de evento concreto que se quiere identificar. Se estudia por tanto la variación de los módulos de los vectores entre zonas adyacentes en un conjunto de frames consecutivos. Para establecer esta comparación se calcula, para cada una de las zonas seleccionadas, un vector único que caracteriza el movimiento en ese área de la imagen. Ese vector se denomina vector de movimiento global y para un frame cualquiera representa la magnitud del campo de movimiento calculado a partir de todos los vectores en una zona de la imagen que llamaremos Z . Posteriormente, se estudia cómo evolucionan desde el frame F_j hasta el frame F_{j+m-1} (m

Este trabajo está financiado por el proyecto TIN2015-64776-C3-3-R del Ministerio de Economía y Competitividad.

frames) los valores de los vectores globales vinculados a dos zonas Z y Z' . Para ello se va a utilizar la covarianza ya que permite determinar si los dos vectores globales crecen de manera similar u opuesta y a partir de esta información se podrá interpretar que es lo que sucede en los m frames en función de los valores de la covarianza y de las zonas comparadas. Para determinar el tipo de evento que sucede, se calcula un intervalo de confianza sobre la covarianza, en los m frames siguientes, y el hecho de que este valor se encuentre por debajo del límite inferior del intervalo, dentro del intervalo, o por encima del límite superior del intervalo es el que permite discriminar entre diferentes tipos de eventos. Posteriormente, se transformarán los valores de los estadísticos anteriormente calculados para así obtener una representación difusa de los mismos que permita caracterizar de manera lingüística los eventos, además de obtener una información complementaria de los mismos. Esta transformación de dominio no utilizará elementos difusos que no vean modificado su valor a lo largo del procesamiento del vídeo, es decir estáticos, sino que éstos modificarán sus características para adaptarse así a las condiciones cambiantes del vídeo de tráfico. La idea es la de definir dinámicamente la respuesta ante cambios en la velocidad, dirección, luminosidad, etc. A continuación, a partir del diseño de un sistema de reglas [4] se detectará cuál, de entre el conjunto de eventos predeterminado, es el que mejor se ajusta o mejor corresponde con aquello que ocurre en la subsecuencia de frames procesada en cada momento. Finalmente, se validará la propuesta anteriormente descrita mediante el procesamiento de vídeos de tráfico con diferentes características de resolución, tasa de frames por segundo, tipo de vía por la que se circula, etc.

Referencias

- [1] E. Ohn-Bar, A. Tawari, S. Martin, M. Trivedi, *On surveillance for safety critical events: In-vehicle video networks for predictive driver assistance systems*, Computer Vision and Image Understanding, Volume 134, May 2015, Pages 130-140
- [2] Iain E. Richardson. *The H.264 Advanced Video Compression Standard*, Wiley Publishing, 2nd ed., 2010.
- [3] P. Sirichai, S. Kaviya, Y. Fujii, P. Yupapin, *Smart Car with Security Camera for Road Accidence Monitoring*, Procedia Engineering, Volume 8, 2011, Pages 308-312.
- [4] R.R. Yager, L.A. Zadeh. *An introduction to fuzzy logic applications in intelligent systems*. Springer Science and Business Media, Vol. 165, 2012.
- [5] L.A. Zadeh. *The concept of a linguistic variable and its application to approximate reasoning II*. Information Sciences, Volume 8, Issue 4, 1975, Pages 301-357.

Hardware implementation of fuzzy inference systems for real-time video processing applications

S. Sánchez-Solano¹, M. Brox², E. Calvo-Gallego¹, A. Gersnoviez², P. Brox¹

¹ *Instituto de Microelectrónica de Sevilla, IMSE-CNM (CSIC/Universidad de Sevilla)*
[santiago,calvo,brox]@imse-cnm.csic.es

² *Department of Computer Architecture, University of Córdoba*
[mbrox, andresgm]@uco.es

As a consequence of its ability to handle inaccurate or incomplete information and its capacity to mimic the human reasoning schema, Fuzzy Logic-based techniques have been widely used in many image processing algorithms [1] [2]. Incorporating these algorithms into current embedded video processing systems requires the use of specific hardware designs capable of providing the data-transfer rates demanded by existing applications [3]. This paper presents a hardware architecture for fuzzy inference systems which is able to provide an inference in each clock cycle, thus allowing the usage of fuzzy techniques in low- and middle-level video processing tasks. Even though the architecture imposes some limitations on the membership functions and inference mechanisms that can be used, it may be suitable for hardware implementation of many fuzzy solutions proposed in the literature.

The block diagram of the proposed architecture for a 2-input system is depicted in Fig.1a. It consists of three layers that implement, respectively, the fuzzification, inference, and defuzzification stages of a typical fuzzy inference system. Membership function circuits (MFCs) at the first stage generate families of normalized triangular functions, where the term “normalized” means that the overlapping degree of two consecutive functions is 2 along all the universe of discourse. For this kind of membership functions, given an input value, the pertinence degree of one of the two active (non-zero) fuzzy labels can be calculated by means of a straight line equation, whilst the other one is obtained by subtracting that value from 1. In order to provide their outputs one clock cycle after the change of the inputs, MFCs include a set of registers to store the parameters (breakpoints and slopes) defining the function family and use a parallel structure, composed of comparators and multiplexors, which evaluates the range of active functions. Finally, each MFC also incorporates the logic required to determine the labels of the active rules that will be evaluated in the inference stage.

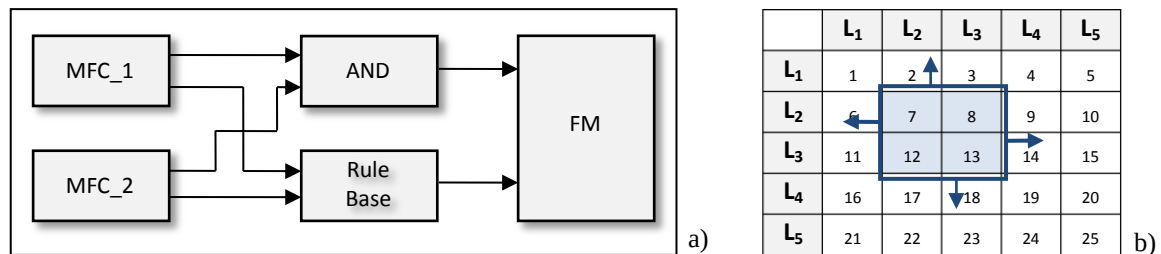


Fig. 1: Block diagram (a) and active-rule based mechanism (b) for fuzzy inference systems.

¹ This work was supported by the projects TEC2014-57971-R and RTC-2014-2932-8 from the Spanish Government (with support from FEDER).

Using membership functions with a limited overlapping degree, has the advantage of setting an upper bound for the number of possible active rules for a given value of the inputs, thus allowing the use of efficient active-rule based strategies to carry out the inference process. This concept is illustrated in Fig.1b, where the four active rules of a 2-input fuzzy system using 5 labels per input with overlapping degree 2 can be selected by shifting, according to the input values, the highlighted window through the rulebase. The combinational block in charge of performing this functionality is composed of a set of registers storing the rules' consequent values and a chain of multiplexors that provide the outputs for a given input label. The inference stage also includes a block to implement the antecedent connective (AND). This element uses as many multipliers as active rules in order to provide simultaneously all the terms required to carry out the Fuzzy-Mean (FM) defuzzification method given by (1).

$$\hat{y} = \sum \alpha_i c_i / \sum \alpha_i \quad (1)$$

being α_i and c_i the activation degree and the consequent of the i -th active rule, respectively.

Using families of triangular functions and the product as antecedent connective ensures that the denominator of (1) is always the unity. The inference can be then evaluated by calculating only the summation that appears in the numerator of the expression. For the example illustrated in Fig.1 this calculus is performed in the defuzzification stage by using four 2-input multipliers blocks and a 4-input adder.

The use of pipeline registers to uncouple the three processing stages allows that all the system components can operate simultaneously over data corresponding to different input values, thus providing a new inference result in each clock cycle. With minimal impact on system latency, this throughput can also be maintained for hierarchical combinations of basic rulebases capable of providing more complex I/O behaviors. In practical realizations, the pipeline strategy is usually also applied inside the processing stages (mainly in the multipliers) to increase the overall system operating frequency.

Results obtained in the development of a fuzzy background modeling application based on this architecture [4] show that its hardware implementation on programmable devices and integrated circuits can be considerably accelerated by using design methodologies and CAD tools similar to those proposed by the authors in [5] and [6].

References

- [1] E. E. Kerre, M. Nachtegaele, eds.: *Fuzzy techniques in image processing*, vol. 52. Physica, 2013.
- [2] Z. Chi, H. Yan, T. Pham: *Fuzzy algorithms: with applications to image processing and pattern recognition*. World Scientific, 1996.
- [3] D. G. Bailey: *Design for embedded image processing on FPGAs*. John Wiley & Sons, 2011.
- [4] E. Calvo-Gallego, P. Brox, S. Sánchez-Solano: "A Fuzzy System for Background Modeling in Video Sequences". IEEE Int. Workshop on Fuzzy Logic and Applications (WILF'2013), Genova (Italia), Nov. 2013
- [5] S. Sánchez-Solano, M. Brox, E. del Toro, P. Brox, I. Baturone: "Model-Based Design Methodology for Rapid Development of Fuzzy Controllers on FPGAs". *IEEE Transactions on Industrial Informatics*, vol. 9, no. 3, pp. 1361-1370, Aug. 2013.
- [6] M. Brox, S. Sánchez-Solano, E. del Toro, P. Brox, F. J. Moreno-Velo, "CAD Tools for Hardware Implementation of Embedded Fuzzy Systems on FPGAs". *IEEE Transactions on Industrial Informatics*, vol. 9, no. 3, pp. 1635-1644, Aug. 2013.

Acciones de Grupos Borrosas

D. Boixader, J. Recasens

Secció Matemàtiques i Informàtica
ETS Arquitectura del Vallès
Universitat Politècnica de Catalunya
Pere Serra 1-15
08190 Sant Cugat del Vallès
 {dionis.boixader,j.recasens}@upc.edu

Resumen

Este trabajo generaliza (fuzzifica) las acciones de un grupo sobre un conjunto para tratar situaciones donde hay imprecisión e incertidumbre. Las acciones borrosas pueden tratar la granularidad de un conjunto o incluso generarla definiendo una relación de equivalencia borrosa en él.

Palabras clave: acción borrosa, aplicación borrosa, relación de T -indistinguibilidad.

Las acciones de un grupo G sobre un conjunto I son una herramienta útil en varias ramas de las matemáticas y de la informática. Ejemplos típicos son las acciones de subgrupos del grupo lineal general $GL(n, \mathbb{R})$ en el espacio vectorial $\mathbb{R}^n$, la acción del grupo lineal proyectivo $PGL(n+1, \mathbb{R})$ en el espacio proyectivo $\mathbb{P}^n(\mathbb{R})$ y la acción de grupos de simetría de polígonos regulares, frisos y mosaicos [4].

También son útiles en Pattern Recognition, en especial en el reconocimiento de caracteres [3, 5], donde una imagen x puede ser comparada con un prototipo p buscando un elemento g de un determinado grupo de modo que la acción de g sobre x lo transforme en p ($gx = p$). Sin embargo es muy poco probable que podamos encontrar un elemento g de un grupo “razonable” con el que gx sea exactamente p . De hecho, consideraremos que x corresponde al carácter p cuando podamos encontrar una g con gx próximo o similar a p . En este caso y en otros donde exista imprecisión y ruido la acción debe relajarse. Las acciones borrosas introducidas en este trabajo permiten esta relajación.

Como veremos, las acciones borrosas están íntimamente ligadas a las relaciones de indistinguibilidad y a las aplicaciones borrosas. El lector puede hallar información sobre estos dos conceptos en [6] y en [1, 2] respectivamente.

Definición 1. Sea G un grupo con elemento neutro e , I un conjunto no vacío y T una t -norma. $\alpha : G \times I \times I \rightarrow [0, 1]$ es una T -acción borrosa de G sobre I si, y sólo si, $\forall g, h \in G, \forall x \in I$

$$1a) \quad T(\alpha(hg, x, y), \alpha(g, x, z)) \leq \alpha(h, z, y)$$

$$1b) \quad T(\alpha(g, x, z), \alpha(h, z, y)) \leq \alpha(hg, x, y)$$

$$2. \quad \alpha(e, x, x) = 1.$$

α es una acción borrosa perfecta si, y sólo si, para todo $g \in G, x \in I$ existe un único $y_x \in I$ tal que $\alpha(g, x, y) = 1$.

$\alpha(g, x, y)$ se interpreta como el grado en que y es la imagen de x bajo la acción de g .

A continuación presentamos algunas de las propiedades más relevantes de las acciones borrosas.

$\alpha(e, x, y)$ da el grado con el cual x se transforma en y por el elemento neutro e de G . Mide por tanto el grado con el cual x e y se pueden considerar equivalentes o indistinguibles y refleja la granularidad de I . De hecho, esta relación es una relación de T -indistinguibilidad en I que denotaremos por E_I . De forma más general, se tiene que fijando $g \in G$, la relación borrosa $R_g(x, y) = \alpha(g, x, y)$ es una aplicación borrosa inyectiva de I en I respecto a la T -indistinguibilidad E_I . Además R_g y $R_{g^{-1}}$ son aplicaciones borrosas inversas.

En una acción borrosa tendemos a considerar equivalentes objetos que son indistinguibles salvo por la acción de un $g \in G$. Así tenemos que la relación borrosa E_α de I definida por $E_\alpha(x, y) = \sup_{g \in G} \alpha(g, x, y)$ es una relación de T -indistinguibilidad, (si T es continua por la izquierda).

Si consideramos una acción borrosa α sobre I , las relaciones borrosas R en I útiles deben ser compatibles (invariantes) por la acción.

Definición 2. R es invariante por α si, y sólo si, $T(R(x, y), \alpha(g, x, x'), \alpha(g, y, y')) \leq R(x', y')$.

Proposición 3. E_I y E_α son invariantes por α . Además, E_I es la menor relación de T -indistinguibilidad invariante por α .

La relación entre acciones borrosas y crisp sobre un conjunto I nos la proporcionan los dos siguientes resultados.

Proposición 4. Si α es una acción borrosa perfecta de G sobre I y, para $g \in G$ y $x \in I$, $y_x \in I$ es el único elemento de I tal que $\alpha(g, x, y_x) = 1$, entonces $gx = y_x$ es una acción crisp¹.

Proposición 5. Sea E una relación de T -indistinguibilidad en I invariante por la acción (crisp) $x \rightarrow gx$. La aplicación $\alpha : G \times I \times I \rightarrow [0, 1]$ definida por $\alpha(g, x, y) = E(gx, y)$ es una acción borrosa perfecta de G sobre I . Además, $E(x, y) = \alpha(e, x, y) = E_I(x, y)$.

Referencias

- [1] Demirci, M. *Fuzzy functions and their fundamental properties*. Fuzzy Sets and Systems 106, 1999, 239–246.
- [2] Demirci, M., Recasens, J. *Fuzzy groups, fuzzy functions and fuzzy equivalence relations*. Fuzzy Sets and Systems 144, 2004, 441–458.
- [3] Grenader, U. *General Pattern Theory*. Oxford Mathematical Monographs. 1993.
- [4] Martin, G.E. *Transformation Geometry: An introduction to symmetry*. Springer-Verlag, 1982.
- [5] Miller, M.I., Younes, L. *Group Actions, Homeomorphisms, and Matching: A general Framework*. International Journal of Computer Vision 41, 2001, 61–84.
- [6] Recasens, J. *Indistinguishability Operators. Modelling Fuzzy Equalities and Fuzzy Equivalence Relations*. Studies in Fuzziness and Soft Computing. Springer, 2011.

¹Recordemos que una acción crisp de G sobre I es una aplicación $\alpha : G \times I \rightarrow I$ que satisface 1: $(hg)x = h(gx)$ 2: $ex = x$, donde $\alpha(g, x)$ se denota por gx .

Medidas de proximidad ordinal: aplicaciones

José Luis García Lapresta¹, Raquel González del Pozo¹, David Pérez Román²

¹ Dep. de Economía Aplicada, Universidad de Valladolid
lapresta@eco.uva.es, ragopo@live.com

² Dep. de Organización de Empresas y C.I.M., Universidad de Valladolid
david@emp.uva.es

Los individuos manifiestan sus opiniones de diversas formas ante los diferentes problemas de decisión que se les plantean. Pueden hacerlo seleccionando una o varias alternativas, elaborando un ranking de las mismas o calificándolas mediante el uso de escalas numéricas o cualitativas, intervalos, números difusos, etc. En este trabajo se supondrá que los agentes evalúan las alternativas a través de una escala cualitativa no necesariamente uniforme (la proximidad psicológica entre términos lingüísticos consecutivos puede ser distinta).

Siguiendo a García-Lapresta y Pérez-Román [1], se considerará que un conjunto de agentes $A = \{1, \dots, m\}$, con $m \geq 2$, evalúa un conjunto de alternativas $X = \{x_1, \dots, x_n\}$, con $n \geq 2$, mediante términos lingüísticos de una escala finita $\mathcal{L} = \{l_1, \dots, l_g\}$ ordenada de menor a mayor, donde la granularidad de $\mathcal{L}$ es al menos 3 ($g \geq 3$).

El modelo considera que la proximidad psicológica entre $l_r \in \mathcal{L}$ y $l_s \in \mathcal{L}$ viene representada por un grado abstracto $\pi_{rs} \in \Delta$. Así, $\Delta = \{\pi_{rs} \mid r, s \in \{1, \dots, g\}\} = \{\delta_1, \dots, \delta_h\}$ será el conjunto de todas las proximidades psicológicas entre términos lingüísticos, con $\delta_1 \succ \dots \succ \delta_h$.

Definición 1 ([1, Def. 1]) Una *medida de proximidad ordinal* sobre $\mathcal{L}$ con valores en Δ es una función $\pi : \mathcal{L}^2 \rightarrow \Delta$, donde $\pi(l_r, l_s) = \pi_{rs}$ denota el grado de proximidad entre l_r y l_s , que satisface las condiciones siguientes:

1. *Exhaustividad*: Para cada $\delta \in \Delta$, existen $l_r, l_s \in \mathcal{L}$ tales que $\delta = \pi_{rs}$.
2. *Simetría*: $\pi_{sr} = \pi_{rs}$, para cualesquiera $r, s \in \{1, \dots, g\}$.
3. *Máxima proximidad*: $\pi_{rs} = \delta_1 \Leftrightarrow r = s$, para cualesquiera $r, s \in \{1, \dots, g\}$.
4. *Monotonía*: $\pi_{rs} \succ \pi_{rt}$ y $\pi_{st} \succ \pi_{rt}$, para cualesquiera $r, s, t \in \{1, \dots, g\}$ tales que $r < s < t$.

Toda medida de proximidad ordinal puede representarse mediante la *matriz de proximidades*,

$$\begin{pmatrix} \pi_{11} & \cdots & \pi_{1s} & \cdots & \pi_{1g} \\ \cdots & \cdots & \cdots & \cdots & \cdots \\ \pi_{r1} & \cdots & \pi_{rs} & \cdots & \pi_{rg} \\ \cdots & \cdots & \cdots & \cdots & \cdots \\ \pi_{g1} & \cdots & \pi_{gs} & \cdots & \pi_{gg} \end{pmatrix} = (\pi_{rs}),$$

Los autores agradecen la financiación del Ministerio de Economía y Competitividad (proyecto ECO2012-32178) y de la Consejería de Educación de la Junta de Castilla y León (proyecto VA066U13).

que es simétrica y de dimensión $g \times g$, con coeficientes en Δ , $\pi_{rr} = \delta_1$, con $r = 1, \dots, g$, y $\pi_{1g} = \delta_h$ (véase García-Lapresta y Pérez-Román [1, Prop. 2]).

La determinación y comparación del grado de consenso en diversos grupos de agentes, así como el uso de procedimientos de clustering jerárquico aglomerativo para agrupar a agentes cuyas opiniones sean lo más homogéneas posibles, son dos de las aplicaciones de las medidas de proximidad ordinal recogidas en García-Lapresta y Pérez-Román [1]. En ambos casos el punto de partida será el *perfil* que recoge las opiniones de los agentes sobre las alternativas: una matriz $V = (v_i^a)$, donde $v_i^a \in \mathcal{L}$ es la valoración otorgada por el agente $a \in A$ a la alternativa $x_i \in X$.

En el presente trabajo se presenta un refinamiento de [1] donde, para evitar la pérdida de información, en lugar de seleccionar una mediana de los grados de proximidad entre las valoraciones asignadas por los agentes, se conservan las dos. Sean $\Delta_2 = \{(\delta_r, \delta_s) \in \Delta^2 \mid r \leq s\}$ y M el operador que asigna el correspondiente par de medianas a cada vector de grados de proximidad, tal como se expone en los siguientes ejemplos: $M(\delta_1, \delta_2, \delta_3, \delta_4, \delta_4) = (\delta_3, \delta_3)$ (caso impar, donde se repite la mediana); $M(\delta_1, \delta_2, \delta_3, \delta_4, \delta_5, \delta_5) = (\delta_3, \delta_4)$ (caso par, donde se consideran las dos medianas).

Definición 2 Dado por un perfil $V = (v_i^a)$, el *grado de consenso* en un subconjunto de al menos dos agentes $I \subseteq A$ sobre un subconjunto de alternativas $\emptyset \neq Y \subseteq X$ se define como

$$C(I, Y) = M \left(\pi \left(v_i^a, v_i^b \right)_{\substack{a, b \in I, a < b \\ x_i \in Y}} \right) \in \Delta_2.$$

Para comparar el grado de consenso entre diferentes subconjuntos de agentes en relación a un subconjunto de alternativas se considerará un orden lineal apropiado sobre Δ_2 . Cuando se produzcan empates, se recurrirá a un proceso secuencial de desempate basado en eliminar las dos medianas de grados de proximidad y la repetición del proceso.

El procedimiento de clustering introducido tiene como elemento determinante las funciones de similitud basadas en el consenso, de forma que dos agentes o clústers se unen para formar un nuevo clúster cuando se maximice el consenso.

Definición 3 Dado un perfil $V = (v_i^a)$, la *función de similitud* relativa a un conjunto de alternativas $\emptyset \neq Y \subseteq X$, $S_Y : (\mathcal{P}(A) \setminus \{\emptyset\})^2 \rightarrow \Delta_2$, se define como

$$S_Y(I, J) = \begin{cases} C(I \cup J, Y), & \text{si } \#(I \cup J) \geq 2, \\ (\delta_1, \delta_1), & \text{si } \#(I \cup J) = 1. \end{cases}$$

En el presente trabajo también se incluye un estudio de campo basado en la aplicación del análisis clúster jerárquico aglomerativo a una muestra de individuos, los cuales manifestaron sus opiniones sobre un conjunto de alternativas a través de una escala cualitativa no uniforme.

Referencias

- [1] García-Lapresta, J.L., Pérez-Román, D.: *Ordinal proximity measures in the context of unbalanced qualitative scales and some applications to consensus and clustering*. Applied Soft Computing **35** (2015) 864–872.

From S-implications to similarities

Erick González-Caballero¹, Susana Díaz², Susana Montes², Rafael Espín³

¹ Higher Polytechnic Institute of Havana, Cuba
erickgc@yandex.com

² UNIMODE research unit, University of Oviedo, Spain
diazsusana@uniovi.es, montes@uniovi.es

³ Autonomous University of Coahuila, Mexico
rafaelalejandrospeinandrade@gmail.com

Fuzzy implications are one of the main operations in fuzzy logic. They appeared as a generalization of the classical implication to fuzzy logic and have been widely studied ([1]). At the beginning, we could find several different definitions in the literature, but lately the usual definition is given by:

Definition 1. ([2]) A map $I : [0, 1] \times [0, 1] \rightarrow [0, 1]$ is called a fuzzy implication if it satisfies the following conditions: $I(\cdot, y)$ is decreasing, $I(x, \cdot)$ is increasing, $I(0, 0) = I(1, 1) = 1$ and $I(1, 0) = 0$.

There are a lot of ways of defining fuzzy implications, some of them using t-norms, t-conorms and negations. One of the most important family of such implications is related to the formalism of Boolean logic (see [1, 2] among many others).

Definition 2. An S-implication associated with a t-conorm S and a strong negation n is defined by $I_{S,n}(x, y) = S(n(x), y)$, $\forall x, y \in [0, 1]$.

Here, we will consider the most usual involutive negation $n(x) = 1 - x, \forall x \in [0, 1]$ and therefore the notation will be simplified to I_S . Moreover, as not any other kind of fuzzy implication will be used, we will call S-implications just implications. These implications can be symmetrized to obtain a logical equivalence (see [5]). Thus, for any pair of elements $x, y \in [0, 1]$ we have that a logical equivalence as: $E_T(x, y) = T(I_S(x, y), I_S(y, x))$, where t-norm T is the dual t-norm of the t-conorm S .

From here we can measure the degree of equivalence between two fuzzy sets as follows: $E_S(A, B) = T_{x \in X} E_T(A(x), B(x))$ for any A, B in $F(X)$, where $F(X)$ denotes the set of all the fuzzy subsets of the finite referential X .

We are interested in the properties of E_S . In particular, we would like to check in which cases it is a similarity, since they are one of the most usual measures of comparison between fuzzy sets.

Definition 3. ([4]) A map: $E : F(X) \times F(X) \rightarrow [0, 1]$ is a similarity measure if it satisfies, for any A, B, C in $F(X)$, the following axiomatic: (1) Reflexivity: $E(A, A) = 1$; (2) Symmetry: $E(A, B) = E(B, A)$ and (3) $E(A, B) \geq E(A, C)$ and $E(B, C) \geq E(A, C)$ if $A \subseteq B \subseteq C$.

The research reported in this paper has been partially supported by the Spanish Ministerio de Economía y Competitividad under Project TIN2014-59543-P and EUREKA SD project (agreement number 2013-2591), that is supported by the Erasmus Mundus programme of the European Union.

Reflexivity is totally characterized as it is followed from the following results.

Proposition 1. *Let T be a t-norm. The associated logical equivalence E_T is reflexive if and only if $T(a, 1 - a) = 0$, for all $a \in [0, 1]$.*

Thus, if the logical equivalence is reflexive, we have that the t-norm is nilpotent. Unfortunately, the converse is not true. Thus, for instance, the Schweizer-Sklar t-norm $T_{SS}(x, y) = (\max(\sqrt{x} + \sqrt{y} - 1, 0))^2$ is nilpotent, but $E_{T_{SS}}(0.3, 0.3) = 0.7163 \neq 1$, that is, it does not satisfy the axiom of reflexivity.

Based on the properties of the nilpotent t-norms, we can express the condition in Proposition 1 as follows:

Corollary 1. *Let T be a t-norm. The associated logical equivalence E_T is reflexive if and only if T is nilpotent and its associated automorphism fulfils that $\varphi(x) + \varphi(1 - x) \leq 1$ for any $x \in [0, 1]$.*

For the symmetry, the behaviour of any t-norm is satisfactory, since any of them are symmetric. Thus,

Proposition 2. *Let T be a t-norm. The associated logical equivalence E_T is symmetric.*

The last step is to study when a reflexive and symmetric logical equivalence is a similarity. The answer is always, as we can see from the following proposition.

Proposition 3. *Let T be a t-norm, whose associated logical equivalence E_T is reflexive. Then E_S also fulfils the third axiom in Definition 3.*

Thus, any reflexive logical equivalence fulfils Axiom 3, but the converse is not true in general. Thus, for instance, the minimum t-norm fulfils Axiom 3, but $E_{T_M}(x, x) \neq 1$ for any $x \in (0, 1)$.

From the previous results, we can analyzed the behaviour of the main families of t-norms: Schweizer-Sklar, Hamacher, Frank, etc. (see [3]) with respect to the coherence of the logical equivalence associated to them.

References

- [1] M. Baczynski and B. Jayaram. *Fuzzy Implications*. Springer, 2008.
- [2] J. Fodor and M. Roubens. *Fuzzy preference modelling and multicriteria decision support*. Kluwer Academic Publishers, 1994.
- [3] E.P. Klement, R. Mesiar and E. Pap. *Triangular Norms*. Kluwer Academic Publishers, 2000.
- [4] X. Lui. Entropy, distance measure and similarity measure of fuzzy sets and their relations. *Fuzzy Sets and Systems* 52, 305–318, 1992.
- [5] J. Recasens. *Indistinguishability Operator*. Springer, 2010.

Functions preserving a class of fuzzy binary relations

Isabel Aguiló¹, Tomasa Calvo², Pilar Fuster-Parra¹, Gaspar Mayor¹, Jaume Suñer¹

¹ *University of Balearic Islands, Palma de Mallorca, Spain*
 {isabel.aguilo,pilar.fuster,gmayor,jaume.sunyer}@uib.es

² *University of Alcalá, Alcalá de Henares, Spain*
 tomasa.calvo@uah.es

In this contribution we consider the class of fuzzy binary relations defined on a set, satisfying the asymmetry, transitivity and negative transitivity conditions. We study those functions that preserve the class of such fuzzy relations, obtaining a very explanatory characterization of them.

Binary relations appear in many knowledge areas, they play a central role in decision making because the technique used is based on pairwise comparison of alternatives. In the classical theory of preference modeling (see [1]) a binary relation R on a given set A is considered as a weak preference relation: aRb if and only if “ a is not worse than b ”. Three binary relations associated to the given preference relation R : strict preference P , indifference I and incomparability J are also defined on A : aPb if and only if aRb and not bRa ; aIb if and only if aRb and bRa ; aJb if and only if not aRb and not bRa . If R is a complete preorder then P is asymmetric, transitive and negatively transitive. Fuzzy logic provide a natural framework for extending the concept of crisp binary relation, by assigning to each ordered pair of elements in the universe of discourse a number from the real unit interval (the degree to which the elements in question are in relation, or in other words, the strength of the link between any two elements). This approach was already used by Zadeh in his first paper on fuzzy sets (1965). In this more general setting, one can describe preferences between alternatives by considering a fuzzy weak preference relation R such that the value $R(a,b)$ in $[0,1]$ is understood as the degree to which the proposition “ a is not worse than b ” is true. The problem of defining fuzzy versions of strict preference, indifference and incomparability relations associated to a fuzzy weak preference relation has been extensively studied by many authors. A detailed monograph on this topic is (see [2]).

In this paper we deal with the class of fuzzy binary relations on a set A satisfying three conditions: asymmetry, transitivity and negative transitivity. For convenience, we denote this class by $\Omega(A)$. Then we study those functions $F : [0,1]^m \rightarrow [0,1]$ that transform any m -dimensional vector $(P_1, \dots, P_m)$ in $\Omega(A)^m$ into an element $F(P_1, \dots, P_m)$ in the same class $\Omega(A)$. The main goal of the contribution is the characterization of such functions. Note that we are mainly interested in preserving the asymmetry and both transitivity conditions, so $F(P_1, \dots, P_m)$ should not be understood as the global fuzzy relation in $\Omega(A)$ summarizing somehow $P_1, \dots, P_m$.

This work was supported by the projects TIN2013-42795-P, and TIN2014-56381-REDT (LODISCO) (Spanish Government).

Definition 1 (The class $\Omega(A)$)

A fuzzy binary relation on a non-empty set A , $P : A \times A \rightarrow [0, 1]$, is in the class $\Omega(A)$ if the following conditions hold for all $x, y, z \in A$:

- P1) $P(x, y) \wedge P(y, x) = 0$ (asymmetric),
P2) $P(x, z) \wedge P(z, y) \leq P(x, y) \leq P(x, z) \vee P(z, y)$ (transitive and negatively transitive respectively).

Note that if $\text{Im}P \subset \{0, 1\}$, i.e. P is a crisp binary relation on A , then Definition 1 is nothing more than the definition of a (crisp) strict preference relation associated to a (crisp) complete preorder. So Definition 1 could be considered an extension of that concept in the fuzzy setting.

Definition 2 (Preserving the class $\Omega(A)$)

A function $F : [0, 1]^m \rightarrow [0, 1]$ preserves the class $\Omega(A)$ if $F(P_1, \dots, P_m)$ is in $\Omega(A)$ for all set A and all $(P_1, \dots, P_m) \in \Omega(A)^m$ where $F(P_1, \dots, P_m)(x, y) = F(P_1(x, y), \dots, P_m(x, y))$ for all $x, y \in A$.

Proposition 1 (First characterization)

A function $F : [0, 1]^m \rightarrow [0, 1]$ preserves the class $\Omega(A)$ if and only if for all $a, b, c \in [0, 1]^m$,

- i) If $F(a) > 0$ then $F(b) = 0$ where b is such that $b_i = 0$ whenever $a_i > 0$,
ii) $b \wedge c \leq a \leq b \vee c \Rightarrow F(b) \wedge F(c) \leq F(a) \leq F(b) \vee F(c)$.

Proposition 2

If $F : [0, 1]^m \rightarrow [0, 1]$ preserves the class $\Omega(A)$, then it is increasing with $F(0) = 0$.

Observe that $F = \text{Min}$ does not preserve the class $\Omega(A)$. Thus, the conditions in Proposition 2 are not sufficient to ensure that F preserves that class.

Proposition 3 (Second characterization)

A function $F : [0, 1]^m \rightarrow [0, 1]$ preserves the class $\Omega(A)$ if and only if there exist $k \in \{1, \dots, m\}$ and an increasing function $f : [0, 1] \rightarrow [0, 1]$ with $f(0) = 0$ such that for all $(a_1, \dots, a_k, \dots, a_m) \in [0, 1]^m$, $F(a_1, \dots, a_k, \dots, a_m) = f(a_k)$.

Thus, a one-variable function $F : [0, 1] \rightarrow [0, 1]$ preserves the class $\Omega(A)$ if and only if it is increasing with $F(0) = 0$.

Remark 1

More general treatment can be done on the basis of a De Morgan triplet (T, S, N) other than that we have used in Definition 1: $(\text{Min}, \text{Max}, 1 - \text{id})$.

References

- [1] M. Roubens, Ph. Vincke : *Preference Modelling*. Springer-Verlag. (1985).
[2] J. Fodor, M. Roubens: *Fuzzy Preference Modelling and Multicriteria Decision Support*. Kluwer Academic Publishers. (1994).

A study on dissimilarities and negative transitivity

Agnieszka Siwiec¹, Susana Díaz², Bernard De Baets³, Susana Montes²

¹ *Instytut Matematyki, University of Gdansk, Poland*
agnieszka.siwiec.zaba@gmail.com

² *UNIMODE research unit, University of Oviedo, Spain*
diazsusana@uniovi.es, montes@uniovi.es

³ *KERMIT, Ghent University, Belgium*
Bernard.DeBaets@UGent.be

In this contribution we study dissimilarities. We first compare them to distances, a similar but different family of operators. We also present two new families of dissimilarities defined using distances and cardinalities of sets respectively. Finally, we carry out a first study about the properties they satisfy. We focus on negative transitivity.

Let us first recall that a dissimilarity is an operator $D : \mathcal{F}(X) \times \mathcal{F}(X) \rightarrow \mathbb{R}$ that takes the value 0 when a set is compared to itself ($D(A, A) = 0$), that is symmetric and satisfies some kind of monotonicity: $D(A, B) \leq D(A, C)$ and $D(B, C) \leq D(A, C)$ for any A, B, C such that $A \subseteq B \subseteq C$, where X is a set, $\mathcal{F}(X)$ denotes the set of fuzzy subsets of X and $A, B, C \in \mathcal{F}(X)$.

In [2] we can find an example of a dissimilarity that is not a distance. It is also known that not every distance is a dissimilarity. Another property that allows to establish a new difference between distances and dissimilarities is the good behavior of dissimilarities when transformed by an increasing function.

Proposition 1 *If $D : \mathcal{F}(X) \times \mathcal{F}(X) \rightarrow \mathbb{R}$ is a dissimilarity and $\varphi : \mathbb{R} \rightarrow \mathbb{R}$ is an increasing map such that $\varphi(0) = 0$, then $\varphi \circ D$ is a dissimilarity.*

Distances are preserved if the function φ is a positive constant, but Proposition 1 does not hold for distances in general.

Despite they are not the same type of functions, distances and dissimilarities are closely related and in fact, we can build new dissimilarities using distances as we next show. We will consider that X is a finite universe and we will denote its cardinality by n .

Definition 1 *Let d be a distance defined on $[0, 1] \times [0, 1]$ into $[0, 1]$ and S a t -conorm, we define the operator $D : \mathcal{F}(X) \times \mathcal{F}(X) \rightarrow [0, 1]$ as follows*

$$D_S(A, B) = S_{x \in X} \left(\frac{d(A(x), B(x))}{n \cdot \sup_{a, b \in [0, 1]} d(a, b)} \right).$$

The research reported in this paper has been partially supported by the Spanish Ministerio de Economía y Competitividad under Project TIN2014-59543-P.

The previous operators are dissimilarities as long as the distance used in the definition is *monotone* as we next show:

Theorem 2 *Let X be a finite referential and let d be a distance in $[0, 1]$. If d is such that $d(a, c) \geq \max\{d(a, b), d(b, c)\}$ for any $a, b, c \in [0, 1]$ such that $a \leq b \leq c$, then the operator introduced in Definition 1 is a dissimilarity.*

This general family of dissimilarities includes some of the most popular in the literature, as for instance the Hung and Yang in [1] ones for the particular case of fuzzy sets.

Another way to obtain new dissimilarities is using cardinalities [3]. A function $\text{card} : \mathcal{F}(X) \rightarrow [0, \infty)$ is a *scalar cardinality* if the following postulates are satisfied for each $a, b \in [0, 1]$, $x, y \in X$, $A \in \mathcal{F}(X)$ and each finite family $\{A_i\}_{i \in I} \subset \mathcal{F}(X)$ s. t. $A_i \cap A_j = \emptyset$ for $i \neq j$:

(P1) $a \leq b \Rightarrow \text{card}(a\mathbf{I}_{\{x\}}) \leq \text{card}(b\mathbf{I}_{\{y\}})$, where $\mathbf{I}_{\{x\}}$ is the indicator in x ,

(P2) $\text{card}(\bigcup_{i \in I} A_i) = \sum_{i \in I} \text{card}(A_i)$,

(P3) $A \in \mathcal{P}(X) \Rightarrow \text{card}(A) = |\text{supp}(A)|$.

Definition 2 *Given a cardinality function card and $\phi : [0, 1] \rightarrow [0, 1]$ a decreasing map such that $\phi(0) = 1$ and $\phi(1) = 0$, we define the operator $D : \mathcal{F}(X) \times \mathcal{F}(X) \rightarrow [0, 1]$ as follows:*

$$D(A, B) = \phi \left(\frac{\text{card}(A \cap B)}{\text{card}(A \cup B)} \right). \quad (1)$$

Theorem 3 *The operator introduced in Def. 2, where the union and intersection are defined by the maximum and minimum respectively, is a dissimilarity for any cardinality, $\text{card}(\cdot)$.*

Next we study the negative transitivity of the families described above. Let us recall that a binary relation R is S -negatively transitive, where S is a t-conorm, if $S(R(a, b), R(b, c)) \geq R(a, c)$ for every a, b, c . We have found that most of the operators included in our families satisfy the negative transitivity defined by the Łukasiewicz t-conorm S_L , but not stronger types of negative transitivity.

Theorem 4 *The operator D_S from Def. 1 is S_L -negatively transitive if S_L dominates S .*

The operator D_{pk1} defined as $D_{pk1}(A, B) = 1 - \frac{\sum_{x \in X} \min(A(x), B(x))}{\sum_{x \in X} \max(A(x), B(x))}$ is a particular case of the family introduced in Def. 2 and we have proven that it is S_L -negatively transitive. As a consequence, some known types of dissimilarities are S_L -negatively transitive as:

$$D(A, B) = 1 - \frac{|\text{core}(A \cap B)|}{|\text{core}(A \cup B)|}, \quad D(A, B) = 1 - \frac{|\text{supp}(A \cap B)|}{|\text{supp}(A \cup B)|} \quad \text{or} \quad D(A, B) = 1 - \frac{|(A \cap B)_{\geq \alpha}|}{|(A \cup B)_{\geq \alpha}|}, \alpha \leq 1.$$

References

- [1] W-L. Hung, M-S. Yang: *Similarity measures of intuitionistic fuzzy sets based on hausdorff distances*. Pattern Recognition Letters **25** (2004) 1603–1611.
- [2] S. Montes, I. Couso, P. Gil, C. Bertoluzza: *Divergence measure between fuzzy sets*. International Journal of Approximate Reasoning **30** (2002) 91–105.
- [3] M. Wygalak: *An axiomatic approach to scalar cardinalities of fuzzy sets*. Fuzzy Sets and Systems **110** (2000) 175–179.

Partial Metric Spaces and Indistinguishability Operators: A Duality Relationship

Pilar Fuster-Parra¹, Javier Martín¹, Oscar Valero¹

¹ *Department of Mathematics and Computer Science,
Balearic Islands University, Palma de Mallorca, Illes Balears, Spain.
{pilar.fuster, javier.martin, o.valero}@uib.es*

In 1982, E. Trillas introduced the notion of T -indistinguishability operator with the purpose of fuzzifying the classical (crisp) notion of equivalence relation ([7]). According to [7] (see also [6]), given a t -norm T , a T -indistinguishability operator on a (nonempty) set X is a fuzzy relation $E : X \times X \rightarrow [0, 1]$ satisfying for all $x, y, z \in X$ the following: (i) $E(x, x) = 1$ (Reflexivity), (ii) $E(x, y) = E(y, x)$ (Symmetry), and (iii) $T(E(x, y), E(y, z)) \leq E(x, z)$ (T -Transitivity).

In the study of indistinguishability, the relationship between metrics (distinguishability operators) and T -indistinguishability operators have been tackled for many authors (see [1, 2, 6] and the references therein). With the aim of exposing the aforesaid relationship, let us recall the well-known notion of metric space. A pseudo-metric on a (nonempty) set X is a function $d : X \times X \rightarrow [0, \infty[$ such that for all $x, y, z \in X$: (i) $d(x, x) = 0$, (ii) $d(x, y) = d(y, x)$, and (iii) $d(x, z) \leq d(x, y) + d(y, z)$ (see [1]). Of course, a pseudo-metric d on X is called a metric provided that it satisfies, in addition, the following axiom for all $x, y \in X$: (i') $d(x, y) = 0$ if and only if $x = y$ (see [2, 5]).

An exhaustive study about the aforesaid relationship between metrics and T -indistinguishability operators has been performed by B. De Baets and R. Mesiar in [1, 2]. Concretely, they proved the following method for constructing pseudo-metrics and metrics from T -indistinguishability operators.

Theorem 1 *Let X be a set such that $\text{card}(X) > 2$. If T^* is a t -norm with additive generator f_{T^*} and T is a t -norm, then the following assertions are equivalent:*

- (i) T^* is weaker than T (i.e., $T^*(x, y) \leq T(x, y)$ for all $x, y \in X$),
- (ii) For any T -indistinguishability operator E on X , the function $d_E : X \times X \rightarrow [0, \infty[$ is a pseudo-metric on X , where $d_E(x, y) = f_{T^*}(E(x, y))$ for all $x, y \in X$.
- (iii) For any T -indistinguishability operator E on X that separates points (i.e., for any $x, y \in X$, $E(x, y) = 1$ if and only if $x = y$), the function $d_E : X \times X \rightarrow [0, \infty[$ is a metric on X , where $d_E(x, y) = f_{T^*}(E(x, y))$ for all $x, y \in X$.

Besides, the next result also was proved in order to construct T -indistinguishability operators from pseudo-metrics and metrics.

The authors acknowledge the support of the Spanish Ministry of Economy and Competitiveness, under grants TIN2013-42795-P and TIN2014- 56381-REDT (LODISCO).

Proposition 2 *Let T^* be a continuous Archimedean t -norm on X with additive generator f_{T^*} . The following assertions hold:*

- (i) *If d is a pseudo-metric on X , then the fuzzy relation $E_d : X \times X \rightarrow [0, 1]$ is a T^* -indistinguishability operators where $E_d(x, y) = f_{T^*}^{(-1)}(d(x, y))$ for all $x, y \in X$.*
- (ii) *If d is a metric on X , then the fuzzy relation $E_d : X \times X \rightarrow [0, 1]$ is a T^* -indistinguishability operators that separates points where $E_d(x, y) = f_{T^*}^{(-1)}(d(x, y))$ for all $x, y \in X$.*

In the light of the preceding result, it seems natural to wonder whether the continuity of the t -norm can be removed from its statement. However, in [1] it was proved that such an assumption can not be taken out.

In 1994, a generalization of the metric notion was introduced by S. Matthews in order to develop a suitable mathematical tool for quantitative models in denotational semantics ([5]). The new aforementioned metric notion was called partial metric and it is defined as follows: a partial metric on a (nonempty) set X is a function $p : X \times X \rightarrow [0, \infty[$ such that for all $x, y, z \in X$: (i) $p(x, x) = p(x, y) = p(y, y)$ if and only if $x = y$, (ii) $p(x, x) \leq p(x, y)$, (iii) $p(x, y) = p(y, x)$, and (iv) $p(x, z) \leq p(x, y) + p(y, z) - p(y, y)$. Notice that a metric on a (nonempty) set X is a partial metric p on X such that $p(x, x) = 0$ for all $x \in X$.

Recently, it has been showed that the notions of T -indistinguishability operator and partial metric are closely related. In fact, on the one hand, M. Demirci and, on the other hand, R. Bukatin, R. Kopperman, and S. Matthews have elucidated that the logical counterpart for partial metrics is, in some sense, a generalized T -indistinguishability operator (see [3, 4]). Inspired by the exposed deed, the purpose of the present work is to extend the way of constructing, on the one hand, metrics from T -indistinguishability operators (given in Theorem 1) and, on the other hand, T -indistinguishability operators from metrics (given in Proposition 2) to the framework of partial metrics and generalized T -indistinguishability operators.

References

- [1] B. De Baets, R. Mesiar, *Pseudo-metrics and T -equivalences*. Fuzzy. Math. **5** (1997) 471–481.
- [2] B. De Baets, R. Mesiar, *Metrics and T -equalities*. J. Math. Anal. Appl. **267** (2002) 531–547.
- [3] M. Bukatin, R. Kopperman, S. Matthews, *Some corollaries of the correspondence between partial metric and multivalued equalities*. Fuzzy Set. Syst. **256** (2014), 57–72.
- [4] M. Demirci, *The order-theoretic duality and relations between partial metrics and local equalities*, Fuzzy Set. Syst. **192** (2011) 45–57.
- [5] S. Matthews, *Partial metric topology*. Ann. New York Acad. Sci. **728** (1994), 183–197.
- [6] J. Recasens, *Indistinguishability Operators: Modelling Fuzzy Equalities and Fuzzy Equivalence Relations*, Springer, Berlin. (2010).
- [7] E. Trillas, *Assaig sobre les relacions d'indistingibilitat*. Proc. Primer Congrés Català de Lògica Matemàtica, Barcelona, (1982) 51–59.

Clasificación difusa aplicada a la estimación de la intensidad del esfuerzo físico a partir de rasgos de expresión facial

Andrés Cencerrado¹, Jordi Moreno¹, Lluís Capdevila^{1,2}

¹ *Health&SportLab R+D department. Eureka Building, UAB Campus. 08193 Barcelona.*
andres.cencerrado@healthsportlab.com,
jordi.moreno@healthsportlab.com

² *Departament de Psicologia Bàsica, Universitat Autònoma de Barcelona.*
lluis.capdevila@uab.cat

El análisis de la expresión facial (AEF) ha demostrado ser una potente herramienta para la medición de diferentes indicadores fisiológicos y determinados estados de actividad (por ejemplo, somnolencia), así como estados emocionales (sorpresa, frustración, etc.) [1,2]. Este estudio tiene como objetivo la identificación objetiva de la intensidad de esfuerzo físico de un sujeto en movimiento a través del registro en vídeo de su rostro. La valoración del esfuerzo físico a través de medidas "clásicas" como la frecuencia cardiaca y la autopercepción de esfuerzo presentan sesgos individuales significativos que evidencian la necesidad de recurrir a otros indicadores complementarios objetivos. El análisis de la expresión facial resulta una técnica muy adecuada para lidiar con esta problemática ya que tiene un carácter completamente no invasivo.

En este trabajo de investigación se evalúa la idoneidad de la aplicación del algoritmo *fuzzy c-means* a este problema. Para ello, se han analizado 383 registros de sujetos realizando la prueba de esfuerzo físico *Test UKK* [3], en la cual se obtuvieron diversas medidas de frecuencia cardiaca (FC) y de variables psico-emocionales, como el grado de esfuerzo autopercebido (escala RPE [4]), además de la grabación en vídeo del sujeto, tanto en el momento inicial de la prueba (PRE, ausencia de esfuerzo), como en el tramo final (POST, esfuerzo máximo). Posteriormente, los datos de vídeo fueron analizados mediante el *Computer Expression Recognition Toolbox* (CERT) [5], que proporciona la intensidad de cada unidad específica de acción (UA) del sistema de codificación facial (FACS) [6].

El algoritmo se aplicó en base a las 5 AUs que presentaban una mayor discrepancia entre las situaciones POST y PRE, generando agrupaciones en un espacio de dimensión 5. Se establecieron 4 agrupaciones para examinar la capacidad de diferenciación entre diferentes grados de esfuerzo (i.e., se agrupa coherentemente con respecto a FC y autopercepción de esfuerzo).

La Figura 1 muestra los grados de pertenencia a cada agrupación obtenida en función del porcentaje de frecuencia máxima teórica individual alcanzado en cada caso. Cada punto, además, está coloreado según la agrupación que presenta mayor grado de pertenencia (lo que equivaldría a una agrupación no difusa). Como se puede observar, la agrupación permite identificar diferentes grados de esfuerzo (agrupación 2 - intensidad baja, 1 y 4 - intensidades altas y moderadas altas, 3 - intermedias). De no utilizar lógica difusa, no obstante, tendríamos un elevado número de errores. Estos resultados, además, son muy similares cuando se expresan en función del esfuerzo autopercebido, y presentan

una similitud muy importante independientemente de la ejecución del algoritmo, tanto en velocidad de convergencia como en centroides obtenidos. Esto denota que el AEF se puede utilizar con los fines previstos, y que el uso de la lógica difusa nos aporta suficientes garantías para diseñar índices graduales de esfuerzo físico a partir de registros en vídeo.

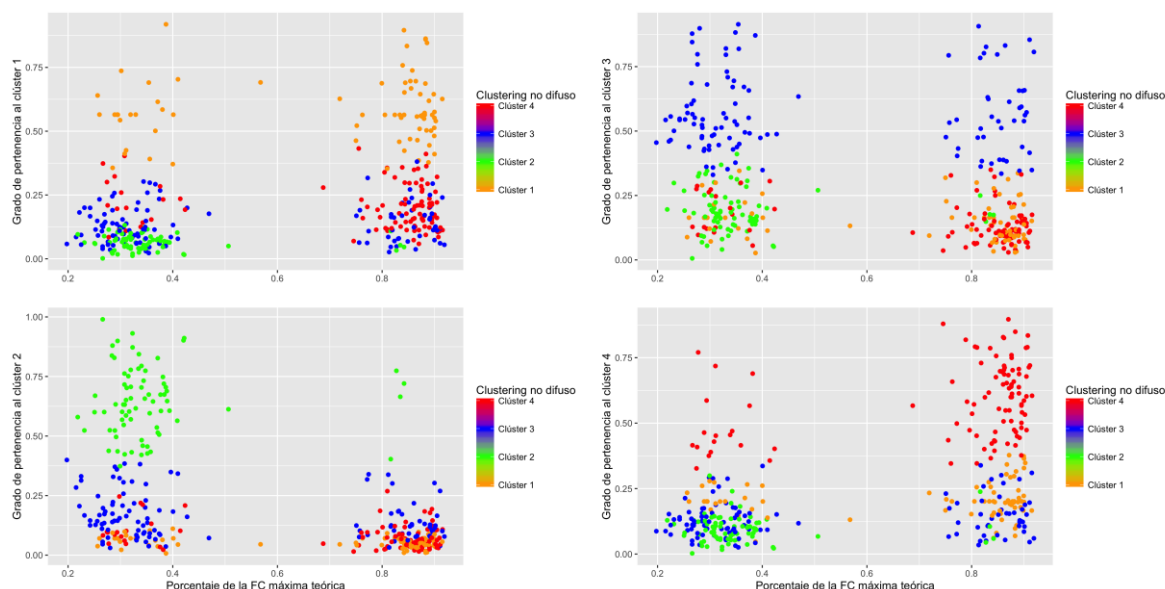


Figura 1: Grados de pertenencia a cada una de las agrupaciones. En color, las asignaciones fijas a cada agrupación en caso de realizar *clustering* no difuso.

Referencias

- [1] Moreno J, Ramos-Castro J, Movellan J, Parrado E, Rodas G, Capdevila L: *Facial Video-Based Photoplethysmography to Detect HRV at Rest*. International Journal of Sports Medicine. **36** (2015) 474-80.
- [2] Lin L, Huang C, Ni X, Wang J, Zhang H, Li X, Qian Z: *Driver fatigue detection based on eye state*. Technol Health Care. **23** (2015) S453-63.
- [3] UKK Institute: *UKK Walk Test: Tester's guide*. Tampere, Finland: UKK Institute. (2006).
- [4] Borg G: *Borg's perceived exertion and pain scales*. Champaign, IL: Human Kinetics. (1998).
- [5] Littlewort G, Whitehill J, Tingfan W, Fasel I, Frank M, Movellan J, Bartlett M: *The computer expression recognition toolbox (CERT)*. Automatic Face & Gesture Recognition and Workshops. (2011) 298-305.
- [6] Ekman P, Friesen WV: *The Facial Action Coding System*. Palo Alto: Consulting Psychologists Press. (1978).

Estudio realizado en el marco del proyecto PTQ-13-06350 del programa Torres Quevedo del Ministerio de Economía y Competitividad: *Machine Learning* con Variables Biométricas y Cognitivas: *Big Data* en soluciones *mHealth* individualizadas, y del proyecto PSI2011-29807-C02-01 del Ministerio de Ciencia e Innovación.

Análisis exploratorio gráfico borroso con Radviz para la representación de la evolución del aprendizaje

José Otero¹ , Rosario Suárez¹ , Luis Junco¹

¹ *Edificio Departamental Oeste, Campus de Viesques, Gijón, Asturias.*
jotero@uniovi.es, mrsuarez@uniovi.es, lajunco@uniovi.es

Con la llegada de las nuevas tecnologías a los centros de enseñanza, aparte de permitir un nuevo diseño de las estrategias de formación, se ha abierto la puerta a la adquisición automática de una gran cantidad de información. El estudio detallado de todo este caudal de datos permite evaluar los cambios en el comportamiento y las mejoras de los estudiantes de tal manera que es posible estimar el nivel de aprovechamiento logrado con cada acción contenida en el guion educativo [Elias 11]. A este nuevo campo de conocimiento e investigación se le ha llamado Learning Analytics o Análisis del aprendizaje, y puede definirse como “el uso de datos y modelos para predecir el progreso y el rendimiento del estudiante, y la capacidad de actuar basándose en esa información” [Siemens 10].

Uno de los problemas que más interés suscita en este campo es la adecuada selección de la información necesaria para una precisa modelización del aprendizaje, tanto a nivel personal como del grupo completo de estudiantes. Y junto con el modelo surge la necesidad de contar con herramientas adecuadas para la visualización o representación del estado de progreso en el que se encuentran los alumnos. Información valiosísima de cara a apoyar la actividad docente, que permita planificar adecuadamente los contenidos del curso y la metodología a emplear en cada momento [Wirth 05]

En consonancia con dicha necesidad hemos publicado un primer estudio [Otero 14] en el que abordábamos el problema de selección de información y modelado del aprendizaje de varios grupos de estudiantes en el campo de la enseñanza de los lenguajes de programación. Para ello y con el fin de automatizar su seguimiento se han aplicado métricas de software para caracterizar la calidad de los trabajos desarrollados durante el curso. Se recogieron y analizaron más de 2000 trabajos de más de 200 alumnos consistentes en pequeños programas de cuyo análisis, basado en una adecuada selección de las métricas más relevantes, se pretende predecir el resultado de la evaluación final del alumno. La información con la que contamos es imprecisa ya que para trabajos correspondientes a la misma temática los valores arrojados por cada métrica resultan variables y además existe una falta de conocimiento en los temas en que el alumno no ha hecho entrega de ningún trabajo.

Abordamos ahora el problema de la representación del estado del aprendizaje, mediante una herramienta de análisis exploratorio gráfico como es la transformación de Radviz aplicado a problemas con información borrosa [Martinez 13], una alternativa al escalado multidimensional MDS [Sanchez 08]. Estas herramientas proyectan la información en un mapa, haciendo que cada estudiante se posicione de acuerdo con su perfil de conocimientos, permitiendo al profesor identificar grupos con carencias similares en su formación, segmentar grupos heterogéneos y mostrar la evolución de las competencias adquiridas durante el desarrollo del curso.

Se presentan los resultados de la comparación de los mapas de evolución de varios grupos de alumnos, unos desarrollados durante el mismo curso, y otros desarrollados en años sucesivos.

Referencias

- [Elias 11] Elias, T. Learning analytics: Definitions, processes and potential. Recuperado en mayo de 2015, de <http://learninganalytics.net/LearningAnalyticsDefinitionsProcessesPotential.pdf>
- [Siemens 10] Siemens, George. What Are Learning Analytics? [Electrónico] Elearnspace, August 25, 2010. URL: <http://www.elearnspace.org/blog/2010/08/25/what-are-learning-analytics/>
- [Wirth 05] Wirth, K. and Perkins, D. Knowledge Surveys: The ultimate course design and assessment tool for faculty and students. Proceedings: Innovations in the Scholarship of Teaching and Learning Conference, St. Olaf College/Carleton College, April 1-3, 2005.
- [Sánchez 08] Sánchez, L., Palacios, A., Suárez, M. R., Couso, I. Graphical exploratory analysis of vague data in the early diagnosis of dyslexia. Proc. 12th IPMU: Information Processing and Management of Uncertainty. pp 1417-1424. 2008
- [Martinez 13] Martinez, A., Sanchez, L., & Couso, I. (2013, July). Engine Health Monitoring for engine fleets using fuzzy Radviz. In Fuzzy Systems (FUZZ), 2013 IEEE International Conference on (pp. 1-8). IEEE.
- [Otero 14] Otero, J., Palacios, A.M., Suarez, M.R., Junco, L.A., Couso, I., Sanchez, L. A procedure for extending input selection algorithms to low quality data in modelling problems with application to the automatic grading of uploaded assignments. Scientific World Journal, Hindawi Publishing Corporation, 2014, p. 1-10; doi: 10.1155/2014/468405; URI : <http://dx.doi.org/10.1155/2014/468405>, ISSN : 2356-6140

Generalizing the Choquet Integral Using Copulas: An Application to Fuzzy Rule-Based Classification Systems

Giancarlo Lucca¹, José Antonio Sanz¹, Graçaliz Pereira Dimuro^{1,2}, Benjamín Bedregal³, Humberto Bustince¹

¹ GIARA, Universidad Pública de Navarra, ²Institute of Smart Cities Spain.

lucca.112793@e.unavarra.es, {joseantonio.sanz, bustince}@unavarra.es

² C3, Universidade Federal do Rio Grande, Brazil
gracalizdimuro@furg.br.

³ DIMAP, Universidade Federal do Rio Grande do Norte, Brazil.
bedregal@dimap.ufrn.br.

A classification problem [1] is a situation that requires the classification of samples into predefined groups or classes according to their features and properties. To accomplish this task, first a classifier is learned using the available information of several samples with known classes. Then, the classifier is used to give a decision about the class new samples belong to.

Fuzzy Rule-Based Classification Systems (FRBCSs) [2] provide an interpretable model, since they are composed by a set of rules that are formed of linguistic labels. An important issue of any FRBCS is the adopted Fuzzy Reasoning Method (FRM), which consists in an inference procedure that uses the information stored in the knowledge base to determine the class in which new samples may be classified.

In the FRM the usage of the aggregation function *maximum* to obtain the global information is widely accepted. However, whenever one considers, for each class, just the information given by a single fuzzy rule, the available information provided by all the other fuzzy rules is ignored. To avoid this shortcoming, Barrenechea et al. [4] introduced a FRM (and it was extended in [3]) that takes into account the information given by all fuzzy rules using averaging operators. To do so, they applied the Choquet Integral as the aggregation operator using different fuzzy measures and evolutionary algorithms.

The aim of this work is to generalize the Choquet Integral used in the FRM defined in [4], by replacing the product operator by copulas, so that system could become more accurate. We also compare our new approach with the classical FRM of the Winning Rule (WR). To do so, the copula functions used in this work are presented in Table 1.

Table 1: Non associative copulas used in this paper

Definition	Observations
$C_F(x, y) = xy + x^2y(1 - x)(1 - y)$	Non Commutative
$C_L(x, y) = \max\{\min\{x, \frac{y}{2}\}, x + y - 1\}$	Non commutative
$C_{Div}(x, y) = \frac{xy + \min\{x, y\}}{2}$	Commutative

This work is supported by Brazilian National Counsel of Technological and Scientific Development (Proc. 233950/2014-1, 232827/2014-1, 481283/2013-7, 306970/2013-9) and by the Spanish Ministry of Science and Technology under Projects TIN2010-15055 and TIN2011-29520.

We have tested the performance of all the classifiers in 28 real world data-sets selected from the KEEL dataset repository using a five folder cross validation model. We have used FARC-HD as the fuzzy classifier [5]. The averages in all the datasets of all the approaches are presented in Table 2, which shows that all the generalizations by copulas achieve a better mean than that of the WR and the standard Choquet Integral.

Table 2: Averaged testing results

Generalization	WR	Choquet	C_F	C_L	C_{Div}
Mean	79.69	79.89	80.13	80.07	80.03

We have used the well-known Wilcoxon test, to statistically compare the copulas versus WR and the original Choquet integral. These results are reported in Table 3. According to these statistical results, we can confirm that the use of our generalization is better than that of the WR when the copula C_F is considered. Moreover, since the standard Choquet Integral (it uses the product) does not provide statistical differences versus the WR, we can also infer that our generalization by copulas allows it to obtain a leap in the performance so that the WR can be statistically enhanced.

Table 3: Wilcoxon Test to compare the different approaches

(a) Comparisons versus the classical FRM				(b) Comparisons versus the Choquet Integral			
Comparison	R^+	R^-	p-value	Comparison	R^+	R^-	p-value
C_F vs. WR	283.5	122.5	0.08	C_F vs. Choquet	227	179	0.58
C_L vs. WR	262	144	0.17	C_L vs. Choquet	217.5	188.5	0.67
C_{Div} vs. WR	256	150	0.22	C_{Div} vs. Choquet	209.5	196.5	0.86
Choquet vs. WR	241	165	0.39				

References

- [1] R. O. Duda, P. E. Hart and D. G. Stork *Pattern Classification*. Wiley, Hoboken, NJ, USA (2001).
- [2] H. Ishibuchi, T. Nakashima and M. Nii *Classification and modeling with linguistic information granules: Advanced approaches to linguistic Data Mining*. Springer-Verlag, Berlin 2004.
- [3] G. Lucca, J. A. Sanz, G. P. Dimuro, B. Bedregal, R. Mesiar, A. Kolesarova and H. Bustince *Pre-aggregation functions: construction and an application*. IEEE Transactions on Fuzzy System, 2015. DOI: 10.1109/TFUZZ.2015.2453020.
- [4] E. Barrenechea, H. Bustince, J. Fernandez, D. Paternain and J. A. Sanz *Using the Choquet Integral in the Fuzzy Reasoning Method of Fuzzy Rule-Based Classification Systems*. Axioms. **2** (2013) 208–223.
- [5] J. Alcalá-Fdez, R. Alcalá and F. Herrera *A Fuzzy Association Rule-Based Classification Model for High-Dimensional Problems With Genetic Rule Selection and Lateral Tuning* IEEE Transactions on Fuzzy Systems. (2011) 857–872.

Tackling Brain Computing Interfaz problems using Fuzzy Rule-Based Classification Systems

Maria Pérez-Zabalza¹, Jose Antonio Sanz², Aránzazu Jurío², Marisol Gomez¹,
Humberto Bustince²

¹ *Mathematics Department, Universidad Pública de Navarra, Spain*

mariapzabalza@gmail.com, marisol@unavarra.es

² *GIARA, Universidad Pública de Navarra, Spain*

joseantonio.sanz, bustince@unavarra.es

A Brain Computer Interface (BCI) is a technology based on acquisition, signal preprocessing, features extraction and classification of brain signals for communication or control of electronic devices. The classification stage has been addressed by different machine learning algorithms (decision trees, support vector machines, random forests, etc...).

The BCI problem faced in this work consist in classifying the following three different spontaneous mental activities of 3 participants:

1. Imagination of repetitive self-paced left hand movements, (*left*, class 2).
2. Imagination of repetitive self-paced right hand movements, (*right*, class 3).
3. Generation of words beginning with the same random letter, (*word*, class 7).

Each user performed 4 sessions in the same day. EEG data were not splitted in trials since the subjects were continuously performing any of the mental tasks. Each session lasted 4 minutes, with breaks of 5-10 minutes between them, where the participant performed a certain task (each corresponding to a certain class) for about 15 seconds and then he/she switched randomly to another task at the operator's request.

In this work, we concentrate in solving the BCI classification module employing soft computing techniques. Soft computing approaches have shown their ability to cope with a great variety of classification problems [1]. Among these techniques, Fuzzy Rule-Based Classification Systems (FRBCSs) [2] are widely used because they offer a good trade-off between accuracy and interpretability, which is obtained by the usage of linguistic labels in the antecedent part of the fuzzy rules.

In this work, we have selected two state-of-the-art FRBCSs: FARC-HD [3] and FURIA [4]. To solve the BCI problem presented here, classifiers have to provide an output every half a second. This is achieved by adding a sliding window of 1 second with 50% overlap after the classification output (originally obtained sample by sample).

The signal preprocessing is the same as the one applied by the winner algorithm of the BCI Competition IV in this dataset [5]. However, we did not apply any feature reduction method. The quality of the results is measured using the accuracy rate. Our aim was to find out whether FRBCSs allow to improve the performance of the method presented in [5] (coded as DBDA).

Table 1: Classification accuracy for the three subjects.

Classifier Method	Subject 1	Subject 2	Subject 3	Average
DBDA	75.14	62.18	50.83	62.72
FARC-HD	80.82	68.69	53.90	67.87
FURIA	81.05	70.75	52.75	67.95

Table 1 shows the results for each user and the corresponding average result, where for each participant the best result is stressed in **bold-face**. We can observe that both FRBCSs obtain similar results, FURIA provides a slightly better overall predictive capability due to the improvement of users 1 and 2. FARC-HD obtains better performance than FURIA with the third user. Finally, both FRBCSs outperform the results obtained by the winner method in [5].

References

- [1] Sanz, J.A., Bernardo, D., Herrera, F., Bustince, H., Hageras, H.: *A Compact Evolutionary Interval-Valued Fuzzy Rule-Based Classification System for the Modeling and Prediction of Real-World Financial Applications with Imbalanced Data*. IEEE Transactions on Fuzzy Systems, 23 (4), art. no. 6849462, pp. 973-990.
- [2] H. Ishibuchi, T. Nakashima and M. Nii.: *Classification and modelling with linguistic information granules: Advanced approaches to linguistic Data Mining*. Springer-Verlag, Berlin 2004.
- [3] J. Alcalá-Fdez, R. Alcalá, and F. Herrera: *A fuzzy association rule based classification model for high-dimensional problems with genetic rule selection and lateral tuning*. IEEE Transactions on Fuzzy Systems, vol. 19, no. 5, pp. 857–872, 2011.
- [4] J. Hühn and E. Hüllermeier: *FURIA: an algorithm for unordered fuzzy rule induction*. . Data Mining and Knowledge Discovery, vol. 19, no. 3, pp. 293–319, 2009.
- [5] Ferran Galán, Francesc Oliva, Joan Guàrdia: *Using mental tasks transitions detection to improve spontaneous mental activity classification*. Med Bio Eng Comput (2007) 45:603–609.

Graph Construction Using Adaptive Local Hybrid Coding Scheme: Application to Face Recognition

M. Tavassoli Kejani¹, F. Dornaika^{2,3}, A. Bosaghzadeh², and H. Mahvash¹

¹ *University of Isfahan, Isfahan, Iran*

mtk.tavassoli@gmail.com, h.mahvash@eng.ui.ac.ir

² *University of the Basque Country UPV/EHU, San Sebastian, Spain*

³ *IKERBASQUE, Basque Foundation for Science, Bilbao, Spain*

fadi.dornaika@ehu.es

Introduction Labeled instances are often difficult, expensive, or time consuming to obtain since they require the intervention of human experts. On the other hand, unlabeled data may be relatively easy to collect. Semi-supervised learning exploits both types of data by using large amount of unlabeled instances together with labeled instances. Graph based semi-supervised learning is a type of semi-supervised task that operates on graphs where data instances are represented by vertices and edges show the similarity between instances. In this paper, we propose a novel graph construction method and apply it to graph-based semi-supervised learning. The core idea is to use adaptive and hybrid coding to get the graph.

Related work: graph construction methods KNN method is the most common graph construction method. It consists of two separate and independent processes. First, the adjacency matrix is constructed by K-Nearest Neighbor method. Second, the weights of the obtained edges are estimated. There are several choices for edge weight calculation, for instance, [2] uses the heat kernel and [3] proposes the use of the inverse of distance as weight. Locally Linear Embedding (LLE) [4] focuses on preserving the local structure of data. The LLE graph can be obtained by applying two processes. First, the adjacency matrix is computed using the KNN method. Second, the non-zero entries of the weight matrix are estimated by reconstructing the data samples from its neighboring samples.

Local Hybrid Coding (LHC) [1] The coding of a sample with respect to a given set of samples can be obtained by applying two phases. In the first phase, the local bases are selected according to Euclidean distance. Thus, the dataset is partitioned in two disjoint subsets: Local bases (K samples) and non-local bases (K' samples). In the second phase, the sample in hand is coded by minimizing a criterion containing three terms: (i) the residual error of the sample reconstruction, (ii) the ℓ_2 regularization term of the local bases coefficients, and (iii) the ℓ_1 regularization of the non-local bases coefficients. The two sets of unknown coefficients are then iteratively obtained by alternating Locality-constrained Linear Coding (LLC) over the local bases with a sparse coding over the non-local bases.

Proposed graph construction method Inspired by LHC, we propose an adaptive coding scheme in order to build a graph. The proposed coding actually attempts to get hybrid code using LLC and sparse code. Unlike LHC [1] which uses Euclidean distance to select the local bases, our proposed method exploits the coefficients provided by LLC [5] in order to adaptively select the local bases. Besides, our work uses the resulting coding scheme in graph construction.

Performance evaluation The datasets used are Extended Yale, PF01, and PIE. In order to quantify the classification performance of the graph construction methods we proceed as follows. First, we reduce the dimension of the images by applying PCA such that it preserves 99% of the data variance then the graph for the whole data set is constructed using the method in hand. The dataset was then separated into labeled and unlabeled sets (three different percentages of labeled samples were used). Finally, the labels of the unlabeled samples were estimated by the Gaussian Random Field (GRF) [6] label propagation algorithm.

For each percentage of labeled samples, ten different splits were used and the percentage of the correct label was calculated over the unlabeled part. The average rates over the ten splits are reported in Table 1. The compared graph construction methods are: KNN graph, LLE graph, LHC graph and our proposed adaptive LHC method. For the KNN and LLE methods, different K values ranging from 5 to 50 were tested. For LHC and our proposed method, we varied the size of local bases K from 5 to 50 and non-local bases K' from 50 to 200. The experiments show that the proposed method can outperform competing methods.

Table 1: Recognition rate (%) for different graph construction methods on three face datasets.

Dataset \ Method	Label percentage	KNN	LLE	LHC	Adaptive LHC
Ext Yale	15%	80.55	85.11	90.47	91.95
	23%	82.73	86.24	92.12	93.76
	31%	83.25	87.56	93.07	94.66
PF01	30%	43.21	49.11	60.24	64.97
	50%	48.49	56.00	67.39	72.22
	70%	51.33	62.9	71.87	76.17
PIE	30%	43.47	51.08	61.05	74.63
	50%	51.89	60.66	70.04	81.05
	70%	54.91	65.28	72.99	82.86

References

- [1] W. Xiang, J. Wang, and M. Long: *Local Hybrid Coding for Image Classification*. International Conference on Pattern Recognition, (2014) 3744–3749.
- [2] M. Beklin and P. Niyogi : *Laplacian Eigenmaps for Dimensionality Reduction and Data Representation*. Neural Computing, **15(6)** (2003) 1373–1396.
- [3] C. Cortes and M. Mohri : *On transductive regression*. In NIPS. (2006) 305–312.
- [4] S. Roweis and L. Saul: *Nonlinear dimensionality reduction by locally linear embedding*. Science. **290(5500)** (2000) 2323–2326.
- [5] F. Dornaika and A. Bosaghzadeh: *Adaptive graph construction using Data self-representativeness for pattern classification*. Information Sciences, 325, (2015) 118–139.
- [6] X. Zhu, Z. Ghahramani, and J. Lafferty: *Semi-supervised learning using gaussian fields and harmonic functions*. ICML, (2003) 912–919.

Information Dynamics of Machine Learning

Francisco J. Valverde-Albacete, Carmen Peláez-Moreno

*Departamento de Teoría de la Señal y de las Comunicaciones.
Universidad Carlos III de Madrid, 28911 Leganés, Spain
{fva,carmen}@tsc.uc3m.es*

Motivation. A mechanism to measure the entropic efficiency of classifiers was presented in [2] that can be used, among other applications, to measure the effects of imbalanced datasets on multiple-class classification [3]. Consider the contingency matrix of a classifier C on a task T . Properly normalized, this matrix can be conceived as a joint distribution P_{XY} between the random variable of input labels (X, P_X) and that of output labels (Y, P_Y) . Then, we may write the following balance equation between the entropies of X and Y :

$$H_{U_X \cdot U_Y} = \Delta H_{P_X \cdot P_Y} + 2 * MI_{P_{XY}} + VI_{P_{XY}} \quad 0 \leq \Delta H_{P_X \cdot P_Y}, MI_{P_{XY}}, VI_{P_{XY}} \leq H_{U_X \cdot U_Y} \quad (1)$$

where the bounds are easily obtained from distributional considerations [2]. Figure 1 depicts this decomposition showing the three crucial regions:

- The *divergence with respect to uniformity*, $\Delta H_{P_X \cdot P_Y}$, between the joint distribution where P_X and P_Y are independent and the uniform distributions with the same cardinality of events as P_X and P_Y , $\Delta H_{P_X \cdot P_Y} = H_{U_X \cdot U_Y} - H_{P_X \cdot P_Y}$, that shows how imbalanced the data and results are.
- The *mutual information*, $MI_{P_{XY}}$, quantifies the force of the stochastic binding between P_X and P_Y , $MI_{P_{XY}} = H_{P_X \cdot P_Y} - H_{P_{XY}}$.
- The *variation of information*, $VI_{P_{XY}}$, embodies the residual entropy, not used in binding the variables, $VI_{P_{XY}} = H_{P_{X|Y}} + H_{P_{Y|X}}$.

If we normalize (1) by the overall entropy $H_{U_X \cdot U_Y}$ we obtain

$$1 = \Delta' H_{P_X \cdot P_Y} + 2 * MI'_{P_{XY}} + VI'_{P_{XY}} \quad 0 \leq \Delta' H_{P_X \cdot P_Y}, MI'_{P_{XY}}, VI'_{P_{XY}} \leq 1 \quad (2)$$

Equation (2) is the 2-simplex in normalized $\Delta' H_{P_X \cdot P_Y} \times 2MI'_{P_{XY}} \times VI'_{P_{XY}}$ space. Each joint distribution P_{XY} can be characterized by its *joint entropy fractions*, $F(P_{XY}) = [\Delta' H_{P_{XY}}, 2 * MI'_{P_{XY}}, VI'_{P_{XY}}]$, whose projection onto the plane with director vector $(1, 1, 1)$ is its *de Finetti (entropy) diagram*, so every binary distribution shows as a point in the triangle. This has been used in several occasions and was further developed in [3]. Recently, it has been generalized to evaluate the information content of datasets [4]: we next generalize it for multilabel classification.

CPM & FVA have been partially supported by the Spanish Government-MinECo projects TEC2014-53390-P and TEC2014-61729-EXP.

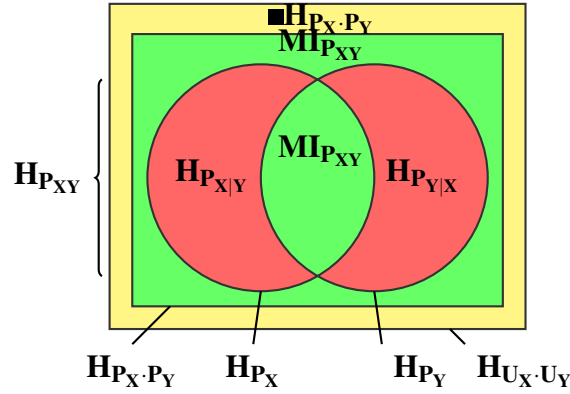


Figure 1: **Extended entropy diagram related to a bivariate distribution, from [2].**

Results. Let $\bar{X} = \{X_i \mid 1 \leq i \leq n\}$ be a set of discrete random variables with joint multivariate distribution $P_{\bar{X}} = P_{X_1 \dots X_n}$, and their corresponding marginals $P_{X_i}(x_i) = \sum_{j \neq i} P_{\bar{X}}(\bar{x})$ where $\bar{x} = x_1 \dots x_n$ is a tuple of n elements. And likewise for $\bar{Y} = \{Y_j \mid 1 \leq j \leq m\}$, with $P_{\bar{Y}} = P_{Y_1 \dots Y_m}$ and their marginals P_{Y_j} . Furthermore let $P_{\bar{X}\bar{Y}}$ be the joint distribution of the $(n+m)$ -length tuples $\bar{X}\bar{Y}$. $\bar{X}$ are the features and $\bar{Y}$ the (multi)labels or classes.

Then the following decompositions hold:

$$\begin{aligned} H_{U_{\bar{X}\bar{Y}}} &= \Delta H_{P_{\bar{X}} \cdot P_{\bar{Y}}} + 2 * M_{P_{\bar{X}\bar{Y}}} + V I_{P_{\bar{X}\bar{Y}}} \\ H_{U_{\bar{X}}} &= \Delta H_{P_{\bar{X}}} + M_{P_{\bar{X}\bar{Y}}} + H_{P_{\bar{X}|\bar{Y}}} & H_{U_{\bar{Y}}} &= \Delta H_{P_{\bar{Y}}} + M_{P_{\bar{X}\bar{Y}}} + H_{P_{\bar{Y}|\bar{X}}} \end{aligned}$$

where $M_{P_{\bar{X}\bar{Y}}} = H_{P_{\bar{X}}} - H_{P_{\bar{X}|\bar{Y}}} = H_{P_{\bar{Y}}} - H_{P_{\bar{Y}|\bar{X}}}$ represents the *dual total correlation* between the sequences of variables or *binding information* [1]. Overall, these balance equations provide a gross estimate of how information from the features reflects on the labels and viceversa: combining them with the classification triangle of (2) should provide a powerful tool for all types of classification tasks.

References

- [1] T. S. Han. Nonnegative entropy measures of multivariate symmetric correlations. *Information and Control*, 36(2):133–156, Feb. 1978.
- [2] F. J. Valverde-Albacete and C. Peláez-Moreno. Two information-theoretic tools to assess the performance of multi-class classifiers. *Pattern Recognition Letters*, 31(12):1665–1671, 2010.
- [3] F. J. Valverde-Albacete and C. Peláez-Moreno. 100% classification accuracy considered harmful: the normalized information transfer factor explains the accuracy paradox. *PLOS ONE*, pages 1–10, january 2014. doi: 10.1371/journal.pone.0084217.
- [4] F. J. Valverde Albacete and C. Peláez-Moreno. The multivariate entropy triangle and applications. In *Hybrid Artificial Intelligence Systems (HAIS 2016), Proceedings*, pages 1–12, Seville (Spain), April 2016. Springer.

Una aproximación difusa del vecino más cercano para clasificación multi-instancia

Pedro Villar¹, Rosana Montes¹, Ana María Sánchez¹, Francisco Herrera²

¹ *Dpto. Lenguajes y Sistemas Informáticos. Universidad de Granada*
pvillarc@ugr.es, rosana@ugr.es, amlopez@ugr.es

² *Dpto. Ciencias de la Computación e Inteligencia Artificial. Universidad de Granada*
herrera@decsai.ugr.es

En clasificación multi-instancia (MIC en sus siglas en inglés) [1], los datos son colecciones de instancias (llamadas bolsas). Las instancias son similares a los ejemplos que se utilizan en problemas de clasificación tradicionales (mono-instancia) y cada bolsa puede tener un número distinto de instancias. La principal característica de la MIC es que las instancias no están etiquetadas, lo están las bolsas en su conjunto. Por tanto, los conjuntos de datos son más complejos y es necesario utilizar algoritmos de aprendizaje específicos para clasificar nuevas bolsas. Los algoritmos para MIC se suelen agrupar en función del nivel de información usado para clasificar nuevas bolsas [2], en clasificadores a nivel de instancia y a nivel de bolsa. En estos últimos, la información utilizada para determinar la clase es de la bolsa en conjunto, no de las instancias individuales y un ejemplo es Citation-KNN [3], que constituye una adaptación del conocido algoritmo del vecino más cercano para MIC. En dicho método, aparte de los K-vecinos de una bolsa b , se usan los C-citadores de la misma, que son las bolsas que contienen a b en su lista de C-vecinos. Es necesario un valor de distancia para determinar cuáles son los vecinos más cercanos. En el caso de MIC, las bolsas son conjuntos de instancias y existen varias medidas de la distancia entre dos bolsas, una de las más utilizadas es la distancia mínima de Hausdorff [3], cuya expresión es: $h_1(A, B) = \min_{a \in A} \min_{b \in B} \|a - b\| = \min_{b \in B} \min_{a \in A} \|b - a\| = h_1(B, A)$

Los conjuntos difusos se han utilizado bastante en aplicaciones de aprendizaje automático tradicionales (mono-instancia). Una de las primeras propuestas de utilización de los conjuntos difusos para mejorar el algoritmo del vecino más cercano es el algoritmo FuzzyKNN [4]. En este trabajo proponemos una extensión difusa de Citation-KNN, que denominamos FuzzyCitation-KNN, incorporando algunas ideas propuestas en el diseño del algoritmo FuzzyKNN. Nuestro método consta de dos fases:

- Una fase preliminar, donde se calcula un valor de pertenencia a cada clase (entre $[0, 1]$) para todas las bolsas del conjunto de entrenamiento, teniendo en cuenta los k_{init} vecinos más cercanos:

$$\mu_c(x_i) = \begin{cases} 0.51 + (v_c/k_{init}) \times 0.49 & \text{if } c = \omega \\ (v_c/k_{init}) \times 0.49 & \text{otherwise} \end{cases} \quad (1)$$

Este trabajo está soportado por el Ministerio de Ciencia e Innovación en el marco del proyecto TIN2014-57251-P y por la Junta de Andalucía en el marco del proyecto P11-TIC-7765.

Donde v_c es el número de vecinos de clase c y ω es la clase original de la bolsa x_i . De esta manera, los ejemplos situados en el "centro" de la clase mantienen sus valores *crisp* originales y las bolsas situadas en la frontera entre clases dividen su valor de pertenencia entre ellas.

- Regla de clasificación. Para clasificar una nueva bolsa b , se calculan sus K -vecinos y sus C -citadores de acuerdo a la distancia mínima de Hausdorff. A continuación, se calcula un grado de pertenencia de b a cada una de las clases (el mayor determinaría la clase) donde cada vecino y cada citador contribuye con su μ_c anteriormente calculado, ponderado por la distancia que le separa de b :

$$\mu_c(b) = \frac{\sum_{i=1}^K \mu_c(x_i)(1/\|b - x_i\|^{2/(m-1)}) + \sum_{j=1}^{n_c(b)} \mu_c(x_j)(1/\|b - x_j\|^{2/(m-1)})}{\sum_{i=1}^K (1/\|b - x_i\|^{2/(m-1)}) + \sum_{j=1}^{n_c(b)} (1/\|b - x_j\|^{2/(m-1)})}$$

En la ecuación anterior, $n_c(b)$ es el número de citadores de b y el parámetro m determina cuánto pondera la distancia: con valores altos la distancia influye poco y con valores cercanos a uno tienen más peso los vecinos/citadores más cercanos.

Para el estudio experimental hemos utilizado 10 conjuntos de datos MIC y el método "leave-one-out". La tabla 1 muestra la media de resultados obtenida. Se puede observar que nuestra propuesta mejora a citation-KNN para todos los valores de K (excepto $K = 0$). Considerando el mejor K para cada método, hemos realizado test estadísticos no paramétricos que indican que hay diferencias significativas entre ambos métodos.

Tabla 1: Porcentaje de acierto. $C = K + 2$, $k_{init} = 3$, $m = 2$

Algoritmo	$K = 0$	$K = 1$	$K = 2$	$K = 3$	$K = 4$	$K = 5$	$K = 6$	$K = 7$
Citation-KNN	68.41	68.41	68.41	69.04	68.73	69.11	68.87	69.03
FuzzyCitation-KNN	67.92	69.79	69.41	70.92	71.51	70.74	71.51	71.60

Referencias

- [1] T. Dietterich, R. Lathrop, and T. Lozano-Pérez: *Solving the multiple instance problem with axis-parallel rectangles*. Artificial intelligence 89:1 (1997) 31–71.
- [2] J. Amores: *Multiple instance classification: Review, taxonomy and comparative study*. Artificial Intelligence 201 (2013) 81-105.
- [3] J. Wang and J. Zucker: *Solving multiple-instance problem: A lazy learning approach*. Proc. 17th Int. Conf. on Machine Learning (2000) 1119–1125.
- [4] J. Keller, M. Gray, and J. Givens: *A fuzzy k-nearest neighbor algorithm*. IEEE Transactions on Systems, Man and Cybernetics 4 (1985) 580-585.

Automatic Simplification of Numerical Expressions using Fuzzy Logic

Susana Bautista¹, Pablo Gervás², Raquel Hervás¹

¹ *Facultad de Informática, Universidad Complutense de Madrid*
subautis@fdi.ucm.es, raquelhb@fdi.ucm.es

² *Instituto de Tecnología del Conocimiento, Universidad Complutense de Madrid*
pgervas@sip.ucm.es

We live in an “Information Technology Society”, where there is a tendency to digitalize all kinds of information. The way in which information is written or presented can exclude many people whose level of reading skills makes them have problems in reading comprehension. There are several factors by which these skills can be affected, such as having had limited access to training, having social problems or having some cognitive disability. In addition, there are specific groups such as the deaf, the autistic, people with language disorders such as aphasia or dyslexia, people who are learning another language or the elderly, who have specific problems with reading.

A first solution to this problem is the manual simplification of information manually to adapt to the difficulties of target users to whom it is directed. There are several initiatives designed to develop the manual processing of text simplification following the European guidelines established by the IFLA [4]. However, manual simplification is too slow and tedious to be efficient in producing the desired material. Therefore, various attempts to automate part of this simplification process have been launched, focusing on the different transformations that can be applied in the process of text simplification.

The aim of simplifying texts is to transform a text to make it easier to understand for certain target users. These tasks mainly aimed at syntactic and lexical constructions that can be applied to the original text to generate a simplified version. In order to perform these tasks researchers must identify what causes this difficulty in specific readers [3], [5].

We focus on a case of information that can cause problems of understanding for many people and creates difficulties for readers: numerical information [1], [2]. Many times, we access information that is represented in the form of numerical expressions such as economic data, statistical data or demographic data. In addition we can find numerical information on a recipe, a news article or a report.

In our work we consider the expressions which represent quantities to be *numerical expressions*; 52.3%, 0.48 or 3489. These expressions might be simplified as *more than 50%*, *around 1/2* or *almost 3500*. In the process of simplification there is a loss of precision to achieve the simplified version of numerical expressions. The modifiers *more than* or *around* are used in these cases to solve this loss of precision. Furthermore, we consider different strategies such as changing the mathematical representation of the expression or rounding the original quantity. For example, 25% can be transformed by 1/4 or 74.7% by 75%.

This paper has been partially supported by the project Prosecco (600653), funded by the European Commission under FP7, the ICT theme, and the Future Emerging Technologies (FET) program.

In this kind of strategies there is a decision to choose a modifier depending on the loss of numerical precision. We need to translate the numerical information lost into linguistic representation using words. We work with a range of values of the original numerical expressions and we want to obtain discrete values represented by linguistic modifiers to simplify the expressions. We want to study applying techniques of fuzzy logic in order to achieve the simplified versions of the original numerical expressions where there is a loss of precision.

The decision to use a modifier or not and what kind of modifier to select is based on the difference between the value of the original quantity of the numerical expression and the value of the simplified expression in each strategy. The kind of modifier depend on the rounded value of the original expression and the distance calculated by the loss of precision. These features should be considered in a set of rules in fuzzy logic in order to define what kind of modifiers could be used in order to help to generate the simplified version or the original numerical expression.

With this goal, the main of this work is to carry out the automatic simplification of numerical expressions present in the text using fuzzy logic to improve the current approaches. Our work is based on the conclusions achieved by empirical studies [1], [2] developed with experts. We consider real important to include fuzzy logic to redefine the rules of simplification applied in our systems and evaluate the output in real experiments with users. The way in which information is presented can cause reading and comprehension problems for many people. The adaptation of information is not an easy process but clearly necessary.

References

- [1] S. Bautista, R. Hervás, P. Gervás, R. Power, and S. Williams. A System for the Simplification of Numerical Expressions at Different Levels of Understandability. In *Proceedings of the Workshop on Natural Language Processing for Improving Textual Accessibility (NLP4ITA)*, 2013.
- [2] S. Bautista and H. Saggion. Making Numerical Information more Accessible: Implementation of a Numerical Expressions Simplification Component for Spanish. *ITL-International Journal of Applied Linguistics. Special Issue on Readability and Text Simplification*. Peeters Publishers, Belgium, 165(2):299–323, 2014.
- [3] J. Carroll, G. Minnen, Y. Canning, S. Devlin, and J. Tait. Practical Simplification of English Newspaper Text to Assist Aphasic Readers. In *Proceedings of the Workshop on Integrating Artificial Intelligence and Assistive Technology (AAAI)*, pages 7–10, Madison, Wisconsin, 1998.
- [4] G. Freyhoff, G. Hess, L. Kerr, E. Menzel, B. Tronbacke, and K. V. Veken. European Guidelines for the Production of Easy-to-Read Information for People with Learning Disability. Technical Report ISLMH, 1998.
- [5] A. Siddharthan and M. Angrosh. Hybrid text simplification using synchronous dependency grammars with hand-written and automatically harvested rules. In *Proceedings of the 14th Conference of the European Chapter of the Association for Computational Linguistics (EACL 2014)*, Gothenburg, Sweden, 2014.

Computing protoforms from high-rate sensor streams

J. Medina¹, M. Espinilla¹, S. Campana², L. Martinez¹

¹ Department of Computer Science. University of Jaen, Spain

jmquero@ujaen.es, mestevéz@ujaen.es, luis.martinez@ujaen.es

² Universidad Nacional Abierta y a Distancia - UNAD, Colombia

sixto.campana@unad.edu.co

In the context of the Internet of Things (IoT) many sensors have been enhanced with respect to their ordinary role to be proactive and collaborative with other devices. This fact provokes a description of the world and the human activity as we have never collected [1].

The data streams provided by the set of sensors is continuously increasing. Generally, these data streams are just analyzed and partially stored. In order to interpret the information of sensor streams is critical developing linguistic descriptions on natural language in a comprehensiveness and interpretability way, depending of the application domain.

In this contribution, we introduces a methodology to compute linguistic descriptions of high-rate sensor streams that considers relative quantifiers on a temporal component by means of *protoforms* [2], which are an useful knowledge model for summarization [3], quantification [4] or in time series [5]. In this work, we propose protoforms in the shape of s^j is $Q_s^j V_r^j T_k$, to represent the linguistic terms V_k^j which describe the sensor stream s^j with the relative quantifier Q_s^j in a temporal terms T_k .

An example of linguistic description computed by the proposed methodology from a data stream of a pulsometer is as follows: *the most of heart rate measurements are high recently*.

Methodology Each sensor data stream s^j is represented as a set of measures $s^j = \{m_i^j\}$, where each measure is composed by $m_i^j = \{v_i^j, t_i^j\}$; where v_i^j represents a value and t_i^j represents the related timestamp of v_i^j provided by the sensor. For each sensor data stream s^j , we describe a fuzzy linguistic term set $V^j = \{V_r^j; r = 1, \dots, n\}$ which represents the linguistic terms of the sensor values. Secondly, we associate a set of fuzzy linguistic temporal terms $T^j = \{T_k^j; k = 1, \dots, m\}$ with the temporal component of the data stream. For sake of simplicity, we write $V_r^j(v_i^j)$ and $T_k^j(t_i^j)$ instead of $\mu_{V_r^j}(v_i^j)$ and $\mu_{T_k^j}(t_i^j)$.

Our methodology proposes to compute the degree of relevance of a linguistic value term $V_k^j(v_i^j)$ in a temporal term $T_k^j(\Delta t_i^j)$ in the sensor streams, being $\Delta t_i^j = |t_0 - t_i^j|$ and t_0 the current time. To do so, we propose the use of aggregation operations [6] based on uninorms using t-norm t-conorm [7], respectively:

$$V_r^j \cap T_k^j(s^j) = \bigcup_{m_i^j \in s^j} V_r^j(v_i^j) \cap T_k^j(t_i^j) \in [0, 1] \quad (1)$$

The Eq. (1) provides the degree of the protoform instances in the shape of s^j is $V_r^j T_k^j$, for example, *heart rate measurements are low recently*. In order to introduce relative quantifiers in the protoform s^j is $Q_s^j V_r^j T_k^j$, each quantifier Q_s^j is described by a transformation function $\mu_{Q_s^j} : [0, 1] \rightarrow [0, 1]$, translating the degree from Eq. (1) to a quantified aggregated degree $Q_s^j(s^j) = (V_r^j \cap T_k^j(s^j))$.

We highlight that protoforms require an accurate assessment to express properly their semantic and granularity. For example, to aggregate measures of high sample rates where data are frequent and spaced, we propose the fuzzy weighted average [8], which represents the frequency degree of the linguist term V_r^j weighted by the term T_k^j in the sensor stream.

$$Q_s^j(s^j) = V_r^j \cap T_k^j(s^j) = \frac{\sum_{m_i^j \in s^j} (V_r^j(v_i^j) \cdot T_k^j(\Delta t_i^j))}{\sum_{m_i^j \in s^j} T_k^j(\Delta t_i^j)} \in [0, 1] \quad (2)$$

References

- [1] Bogue, R. (2014). Towards the trillion sensors market. *Sensor Review*, 34(2), 137-142.
- [2] Zadeh, L. A. (2006). Generalized theory of uncertainty (GTU) principal concepts and ideas. *Computational Statistics & Data Analysis*, 51(1), 15-46.
- [3] Kacprzyk, J., & Zadrony, S. (2005). Linguistic database summaries and their protoforms: towards natural language based knowledge discovery tools. *Information Sciences*, 173(4), 281-304.
- [4] Zadeh, L. A. (1983). A computational approach to fuzzy quantifiers in natural languages. *Computers & Mathematics with applications*, 9(1), 149-184.
- [5] Sanchez, A. B. N. M. D., & Trivino, G. (2015). Fuzzy Knowledge Representation for Linguistic Description of Time Series. 16th World Congress of the International Fuzzy Systems Association - 9th Conference of the European Society for Fuzzy Logic and Technology, Gijn (Spain), Atlantis Press, pp. 1346-1353.
- [6] Medina, J., Espinilla, M., Moya, F. & Nugent, C. (2015). Activity recognition by means of rule-based inference engine based on fuzzy linguistic approach. The 13th Scandinavian Conference on Artificial Intelligence. Sweden.
- [7] Gupta, M. M., & Qi, J. (1991). Theory of T-norms and fuzzy inference methods. *Fuzzy sets and systems*, 40(3), 431-450.
- [8] Dong, W. M., & Wong, F. S. (1987). Fuzzy weighted averages and implementation of the extension principle. *Fuzzy sets and systems*, 21(2), 183-199.

Una aproximación a la generación de expresiones de referencia con propiedades difusas

Nicolás Marín, Gustavo Rivas-Gervilla, Daniel Sánchez

Departamento de Ciencias de la Computación e I.A.

Universidad de Granada, 18071 - Granada, España

nicm@decsai.ugr.es, g.r.gervilla@gmail.com, daniel@decsai.ugr.es

Los sistemas de generación automática de descripciones lingüísticas de imágenes, que podemos llamar sistemas *image-to-text*, tienen como objetivo generar texto que describa la escena representada en la imagen tratando de simular la forma en la que lo haría el ser humano [4, 2]. Este tipo de sistemas han atraído mucha atención en la intersección de los campos de Visión Artificial y Generación de Lenguaje Natural (GLN) [6].

Nuestro enfoque para la construcción de sistemas *image-to-text* se basa en los siguientes elementos [3]:

- Primero hay que modelar los conceptos que se van a utilizar en la descripción. Muchos de estos conceptos, como los relacionados con características de la imagen como el color, la textura, la definición de regiones y su forma (entre otras) son difusos por naturaleza y su definición depende del usuario o, en un sentido más amplio, del contexto. Por tanto, su modelado puede beneficiarse de la bien conocida capacidad de los conjuntos difusos para afrontar la representación de este tipo de conceptos mediante etiquetas lingüísticas [1].
- En segundo lugar, hay que establecer una forma de relacionar los datos de la imagen con dichos conceptos. Una técnica que se ha empleado para ello en el ámbito de la generación de lenguaje natural es la construcción de grafos que representen distintos elementos de la imagen y relaciones entre los mismos en base a esos conceptos.
- Por último, es necesario determinar medidas de calidad y desarrollar técnicas algorítmicas de generación basadas en dichas medidas de forma que, a partir de la exploración del grafo anteriormente mencionado, encuentren una buena expresión lingüística de la información relevante para el usuario que hay en la imagen.

En este trabajo presentamos algunos avances en relación con los dos últimos puntos: con respecto al segundo punto, hemos desarrollado un nuevo modelo de representación formal que llamamos *red de información* y con respecto al tercero nos hemos centrado en el desarrollo de medidas y técnicas heurísticas para elaborar expresiones de referencia (*referring expressions*) [5].

Este trabajo ha sido desarrollado en el marco del proyecto TIN2014-58227-P *Descripción lingüística de información visual mediante técnicas de minería de datos y computación flexible*. Este proyecto está financiado por el Ministerio de Economía y Competitividad y el Fondo Europeo de Desarrollo Regional - FEDER.

En una red de información, los nodos representan regiones en la imagen y tienen como propiedades conceptos de interés como los que hemos enunciado en el primer punto. Los arcos sirven para representar relaciones entre las regiones. Tanto las propiedades como las relaciones pueden ser difusas. Mediante el uso de este tipo de redes se puede establecer una correspondencia entre los datos de la imagen y los conceptos difusos que son la base para elaborar las descripciones. El uso de redes de información resulta adecuado como mecanismo de representación del conocimiento en tareas de descripción lingüística, entre ellas, aquellas destinadas a la elaboración de expresiones de referencia.

La generación de expresiones de referencia es una de la tareas más importantes dentro del ámbito de la generación de lenguaje natural. Una expresión de referencia es un sintagma nominal cuyo objetivo comunicativo es señalar sin ambigüedad un objeto, en nuestro caso dentro de una escena. Debe por tanto ser una descripción precisa del objeto en cuestión al tiempo que no debe serlo para el resto de los objetos que aparecen [5]. El reto consiste en elegir un subconjunto adecuado de propiedades del objeto que lo identifiquen de forma única, y hacerlo de la misma forma en la que lo haría un ser humano en una situación similar. Nuestra aportación en relación con este problema es doble: por un lado, hemos estudiado medidas adecuadas de la calidad de un sintagma nominal como expresión de referencia en un determinado contexto, teniendo en cuenta que las propiedades que la constituyen pueden ser difusas, medidas muy relacionadas con el concepto de especificidad de conjuntos difusos; por otro lado, hemos explorado la aplicación de técnicas heurísticas voraces en la construcción de expresiones de referencia sobre redes de información.

Referencias

- [1] Rita Castillo-Ortega, Jesús Chamorro-Martínez, Nicolás Marín, Daniel Sánchez, and José M. Soto-Hidalgo. Describing images via linguistic features and hierarchical segmentation. In *FUZZ-IEEE 2010*, pages 1–8, 2010.
- [2] Girish Kulkarni, Visruth Premraj, Vicente Ordonez, Sagnik Dhar, Siming Li, Yejin Choi, Alexander C. Berg, and Tamara L. Berg. Baby talk: Understanding and generating simple image descriptions. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 35(12):2891–2903, 2013.
- [3] Nicolás Marín and Daniel Sánchez. On generating linguistic descriptions of time series. *Fuzzy Sets and Systems*, 285:6–30, 2016.
- [4] Margaret Mitchell, Jesse Dodge, Amit Goyal, Kota Yamaguchi, Karl Stratos, Xufeng Han, Alyssa Mensch, Alex Berg, Xufeng Han, Tamara Berg, and Hal Daume III. Midge: Generating Image Descriptions From Computer Vision Detections. In *EACL 2012*, pages 747–756, Avignon, France, 2012. Association for Computational Linguistics.
- [5] Ehud Reiter and Robert Dale. A fast algorithm for the generation of referring expressions. In *COLING’92*, pages 232–238, 1992.
- [6] Ehud Reiter and Robert Dale. *Building Natural Language Generation Systems*. Cambridge University Press, Cambridge, UK, 2000.

Perspective-based Character Description in Interactive 3D Environments

Gonzalo Méndez¹, Raquel Hervás², Pablo Gervás¹, Ricardo de la Rosa², Daniel Ruiz²

¹ *Instituto de Tecnología del Conocimiento - Universidad Complutense de Madrid*
{gmendez,pgervas}@ucm.es

² *Facultad de Informática - Universidad Complutense de Madrid*
raquelhb@fdi.ucm.es, {ricarosa,danruiz}@ucm.es

Interactive narrative generation in 3D environments can benefit from the fact that a lot of information is implicitly and explicitly available in the environment to generate perspective based descriptions of people and objects that change according to the user's point of view, a problem that has not been addressed in general in the field of referring expression generation [5]. One of the difficulties when generating these descriptions is deciding when we must change the referring expression we use to describe the element's situation, even more if we take into account that not everybody refers to an element in the same manner.

A survey was carried out to investigate how spatial properties and points of view affect the composition of descriptions used when referring to another person. With the results of this survey, and using the graphical engine Unity 3D, an algorithm has been developed to generate descriptions of characters in real time that change according to the user's position within the environment. Further details about this work can be read in [7].

The developed algorithm was implemented in a game where the user had to find the character that was being described, for which he could move around the environment and the provided description changed accordingly. The content of the description is based on the character's physical appearance, its position within the environment, and its situation with respect to other relevant characters and objects that are present in the environment: (1) *attribute-based description*, which refers to the static characteristics of an individual and its environment, and they cannot change during the simulation. This part of the description is generated using classical referring expression generation algorithms, such as *Greedy* or *Incremental* [5]; (2) *perspective-based description*, which has to be generated in real time according to the situation of the user relative to the situation of the described person.

Each description can be composed of sub-descriptions, which contain partial information of a full description. Every sub-description is generated according to the data that is obtained from the scene. There are four possible types of sub-descriptions: (1) *meta description*: the static description generated by the classical algorithms; (2) *description of the static points*: this description contains the information of reference points scattered all over the scene; (3) *description of the visibility*: it contains the information about the visibility of the described person from the user's point of view. For example, the fact that the described person is behind a column; (4) *distance between the described person and the user*: it contains a textual explanation of the distance between the described person and the user.

This paper has been partially supported by the project ConCreTe (611733), funded by the European Commission under FP7, the ICT theme, and the Future Emerging Technologies (FET) program.

The last three sub-descriptions are generated and stored until the algorithm checks certain flags and the information needed for the description is generated accordingly. First it checks the distance between the user and the reference points previously put into the scene. If the distance to the user is greater than a predefined constant, the generated description contains information about the proximity of the described person to that point. Otherwise, the algorithm checks if the described user is on the camera area. If not, the generated description must contain the positional references of the described person to the user: it indicates if the described person is to the left, right or behind the user. If the described person is within the camera area, it checks the absolute distance between the described person and the user. It indicates the user whether he is near, not far or far from the described person. By mixing sub-descriptions and flag checking, the perspective-based descriptions are generated. As a result, in each case a specific order and sub-description appear in the generated description (e.g. *The boy in the black shirt sitting near the window. The described person is behind another person. He is not far*).

The proposed solution to generate descriptions is not currently based on the use of fuzzy logic, although it can benefit from it to generate descriptions of spatial relationships. Although we have not been able to find an approach of this kind in the reviewed literature, there have been some efforts to solve similar problems in the fields of image analysis and computer vision, as described in [2], whose authors have subsequently developed several methods to describe spatial relations between objects [1, 3]. Other authors have tackled the problem of automatic scene descriptions in image analysis using fuzzy rule-based systems [4] and fuzzy sets [6], through the use of histograms of angles and forces and a dictionary of labels.

References

- [1] I. Bloch. Fuzzy relative position between objects in image processing: new definition and properties based on a morphological approach. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 7(2):99 – 133, April 1999.
- [2] I. Bloch and A. Ralescu. Directional relative position between objects in image processing: a comparison between fuzzy approaches. *Pattern Recognition*, 36(7), 2003.
- [3] C. Hudelot, J. Atif, and I. Bloch. Fuzzy spatial relation ontology for image interpretation. *Fuzzy Sets and Systems*, 159(15):19291951, August 2008.
- [4] J. M. Keller and X. Wang. A fuzzy rule-based approach to scene description involving spatial relationships. *Comp. Vision and Image Understanding*, 80(1):21–41, Oct. 2000.
- [5] E. Krahmer and K. van Deemter. Computational generation of referring expressions: A survey. *Computational Linguistics*, 38(1):173–218, March 2012.
- [6] P. Matsakis, J.M. Keller, L. Wendling, J. Marjamaa, and O. Sjahputera. Linguistic description of relative positions in images. *IEEE Transactions on Systems, Man and Cybernetics*, 31(4):573–88, 2001.
- [7] R. de la Rosa and D. Ruiz. Perspective-based referring expression generation in a 3D environment. BSc. Project, Facultad de Informática - UCM, June 2015.

Using Fuzzy Logic to Model Character Affinities for Story Plot Generation

Gonzalo Méndez¹, Pablo Gervás¹, Carlos León²

¹ *Instituto de Tecnología del Conocimiento - Universidad Complutense de Madrid*
gmendez@fdi.ucm.es, pgervas@sip.ucm.es

² *Facultad de Informática - Universidad Complutense de Madrid*
cleon@fdi.ucm.es

In the book *The Thirty-Six Dramatic Situations* [4] Polti explores the assertion made by Gozzi (author of *Turandot*) saying that there can only be thirty-six tragic situations. At the end of the book, he begins his conclusions by saying that, to obtain the nuances of the situations, the first thing he did was to “*enumerate the ties of friendship or kinship between the characters*”. A century before that, Goethe had already proposed his theory of *elective affinities* to depict human relations, specially marriages, and he showed how affinities between characters can be represented by a topological chart [6]. This evidences that the affinity between characters is an important factor to take into account when generating stories, and one that can help us to maintain the necessary narrative tension to keep the reader interested in the story. One of the most relevant research works on the subject is *Thespian* [5]. The authors describe the use of affinity to model social interaction which affects how characters can behave towards each other. Affinity is affected by other factors, such as social obligations and characters goals.

We have developed a storytelling system where we use a numeric representation that allows us to use common arithmetic operators to modify the degree of affinity between characters [3]. The main limitation of this approach is that it is difficult to calibrate the model and interpret what is happening in the simulation. To overcome the limit, we have opted for a representation similar to the fuzzy concepts proposed in [7], an approach that has already been used by other authors to model cognitive architectures [1, 2]. We have modeled four levels of affinity according to four different affinity kinds: foe (no affinity), indifferent (slight affinity), friend (medium affinity) and mate (high affinity). These four levels of affinity overlap on their limits, which allows for relationships not to change constantly when moving around the limits of two different levels. An additional aspect of affinity is the lack of symmetry. Given two characters, their mutual affinity is likely to have different values and it may even be situated in different levels, with the exception of mates: *character A* considers *character B* as its mate only if *character B* considers *character A* as its mate, too. However, if they are not mates, *character A* may think *character B* is a friend, while *character B* may think *character A* is a foe.

There are two ways in which the affinity value can change. The first one is by lack of interaction, in which case the affinity value moves towards the indifferent level. The second

This paper has been partially supported by the project WHIM (611560), funded by the European Commission under FP7, the ICT theme, and the Future Emerging Technologies (FET) program.

one is through interactions among characters, which perform according a hand-written rule set. Each affinity level defines a set of valid interactions. For instance, a character may only propose to carry out friend actions when dealing with a friend. Characters ignore proposals that do not correspond to their perceived affinity level, and receiving such proposals may penalize the affinity with the character proposing them. The exception to this rule are foes, who carry out what they intend to do irrespective of what the other character may want. When receiving a proposal, a character may decide to either accept or reject it. If the proposal is accepted, both characters increase their mutual affinity. If it is rejected, the sender will penalize its affinity with the receiver.

The described model has been implemented by means of a multi-agent system which contains two types of agents: a *Director Agent*, which is in charge of setting up the execution environment and creating the characters; and *Character Agents*, one for each character of the story. They are the ones that interact to generate the story. Currently, the story consists of a set of interactions that make the affinity between characters change accordingly. Each Character Agent is endowed with three different behaviours independent from each other: one to interact with its mate, another one to interact with its friends and the last one to interact with its foes; no interaction is initiated with indifferent characters.

We intend to endow characters with personality traits and emotions, in order to complement the affinity model and give characters the possibility to make decisions in a more cognitive way. We plan to use an approach similar to the one described in [2] to model emotions, so that it can be easily integrated with the present model. This work line is specially relevant, since it will allow characters to make consistent decisions according to the personality and state of mind, instead of random ones, and it will also allow them to behave in a different way from each other.

References

- [1] M. El-Nasr, J. Yen, and T. R. Ioerger. Flame - fuzzy logic adaptive model of emotions. *Autonomous Agents and Multi-agent Systems*, 3(3):219–257, 2000.
- [2] R. Imbert and A. de Antonio. An emotional architecture for virtual characters. In *International Conference on Virtual Storytelling*, pages 63–72, 2005.
- [3] G. Méndez, P. Gervás, and C. León. *On the Use of Character Affinities for Story Plot Generation*, volume 416 of *Advances in Intelligent Systems and Computing*, chapter 15, pages 211–225. Springer, 2016.
- [4] G. Polti. *The Thirty-Six Dramatic Situations*. 1917.
- [5] M. Si, S. C. Marsella, and D. V. Pynadath. Thespian: Modeling socially normative behavior in a decision-theoretic framework. In *Intelligent Virtual Agents*, volume 4133 of *Lecture Notes in Computer Science*, pages 369–382. Springer, 2006.
- [6] J. W. von Goethe. *Elective Affinities / Kindred by Choice*. 1809.
- [7] L. A. Zadeh. A computational approach to fuzzy quantifiers in natural languages. *Computers & Mathematics with Applications*, 9(1):149–184, 1983.

Generación de lenguaje natural a partir de secuencias de datos utilizando modelado espacio-temporal

Juan Moreno-García¹, Luis Jimenez-Linares¹, Luis Rodriguez-Benitez¹

¹ *Departamento de Tecnologías y Sistemas de Información, Uni. de Castilla-La Mancha*
 juan.moreno@uclm.es, luis.jimenez@uclm.es, luis.rodriguez@uclm.es

Hoy en día se utilizan series de tiempo en gran cantidad de situaciones. Los consumidores de esta series necesitan obtener la información más relevante de la misma y que se represente en un formato apropiado. Por ello se está trabajando mucho en describir las mediante lenguaje natural. Por este motivo, en este trabajo pretendemos mostrar algunas propuestas para procesar las series de una manera sencilla, rápida y que puede ser utilizada en tiempo real y generar una descripción de la misma en lenguaje natural. El propuesta se fundamente en la lógica difusa que es muy apropiada para este tipo de problemas. Las series de tiempo representan una muestra continua y secuencial en el tiempo tomada de un origen. Se pueden distinguir dos tipos de series (expresado en inglés): (1) Time Series (TS); (2) Multivariate Time Series (MTS). Ambos tipos de series tienen en común que los datos tienen una relación temporal, es decir, son secuenciales en el tiempo. Además, también existe una relación espacial en las MTS, con lo que esta relación espacio-temporal puede ser utilizada para modelar la entrada y obtener la información deseada. Así, se puede hacer un tratamiento secuencial en el tiempo (y el espacio si procede) y agrupar en intervalos “situaciones similares” que posteriormente serán tratadas mediante operadores de la lógica difusa para la obtención de las descripciones lingüísticas. Esta idea puede ser utilizada en el dominio espacio y tiempo.

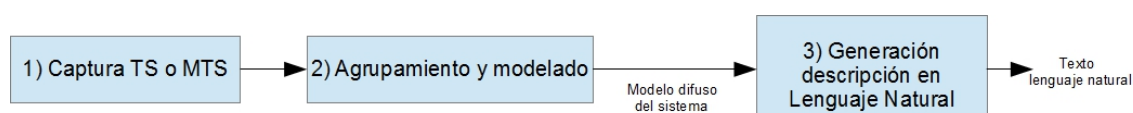


Figura 1: Esquema de propuesta.

Nuestro grupo lleva tiempo trabajando en este tipo de técnicas que han sido refinadas con la experiencia y aplicadas a diferentes problemas, como movimientos humanos [1, 2], deporte [3], vídeo en formato comprimido (MPEG) [4]. Nuestros trabajos siguen un esquema común (Figura 1). Su funcionamiento consiste en representar la entrada representada como una TS o una MTS, a continuación, se modela mediante técnicas de lógica difusa realizando un agrupamiento de los puntos consecutivos como conjuntos difusos o etiquetas lingüísticas (pasos 1 y 2). Posteriormente, utilizando el modelo difuso del sistema, se generan las descripciones lingüísticas (paso 3). A continuación se detallarán los trabajos más destacados en esta línea.

En [1] se presenta cómo obtener un modelo difuso de un sistema dinámico modelado como MTS. El modelo creado se denomina *Temporal Fuzzy Model* y obtiene una cadena

Este trabajo está bajo el proyecto TIN2015-64776-C3-3-R del Ministerio de Economía y Competitividad.

formada por estados y transiciones difusas que representan al sistema de forma secuencial. Su método de inferencia se diferencia principalmente de los tradicionales en que es secuencial, se salta de un estado a otro cuando la transición tiene más grado de pertenencia a la entrada que es estado actual. Es un modelo que obtiene un bajo error y es interpretable. La experimentación se realiza sobre la marcha humana.

En [4] se presenta un sistema para reconocer acciones en secuencias de vídeo. Se realiza un proceso de segmentación y seguimiento utilizando los datos de entradas codificados en MPEG y que son fuzzificados antes de ser procesados. Para el reconocimiento de las acciones se utiliza una máquina de Mealy. Es un método rápido y que utiliza modelos similares a los descritos en este trabajo.

En [2] se presenta un método para generar descripciones de MTS utilizando un modelo denominado *Temporal Fuzzy Model (TFM)* como representación intermedia. Los TFM agrupan la información temporal consecutiva de la MTS en etiquetas lingüísticas definidas previamente y tienen una apariencia similar a los modelos difusos tradicionales. Su método de inferencia es secuencial, comienza con una regla activa, la primera, y va saltando a la siguiente cuando esta última tiene más grado de pertenencia que la anterior. El trabajo también presenta una primera aproximación a cómo utilizar los operadores difusos para generar las descripciones lingüísticas.

En [3] se detalla un sistema para automáticamente generar descripciones en lenguaje natural desde una MTS. La MTS se modela mediante un TFM, del que a continuación se obtienen las tendencias y los puntos de interés (máximos y mínimos) que son representados en otra estructura más compleja y que recoge toda la información espacio-temporal del sistema. Posteriormente se utiliza un modelo que refleja los eventos interesantes para el experto del campo de aplicación. Las descripciones obtenidas reflejan en lenguaje natural la ocurrencia de los eventos, luego se puede generar la información interesante para el experto. La experimentación del método propuesto se realiza en el deporte.

Referencias

- [1] J. Moreno-Garcia, J.J. Castro-Schez, L. Jimenez: *A Fuzzy Inductive Algorithm for Modeling Dynamical Systems in a Comprehensible Way*. IEEE transactions on Fuzzy Systems. **15(4)** (2007) 652–672.
- [2] J. Moreno-Garcia, L. Rodriguez-Benitez, J. Giralt, E. del Castillo: *The generation of qualitative descriptions of multivariate time series using fuzzy logic*. Applied Soft Computing. **23** (2014) 546–555.
- [3] J. Moreno-Garcia, J. Abián-Vicén, L. Jimenez-Linares, L. Rodriguez-Benitez: *Description of multivariate time series by means of trends characterization in the fuzzy domain*. Fuzzy Sets and Systems. **285** (2016) 118–139.
- [4] L. Rodriguez-Benitez, C.J. Solana-Cipres, J. Moreno-Garcia, L. Jimenez-Linares: *Approximate reasoning and finite state machines to the detection of actions in video sequences*. International Journal of Approximate Reasoning. **52(4)** (2011) 526–540.

A Contribution to the Study of Paraphrasing using Fuzzy Logic and WordNet

M. Pereira-Fariña¹, Chris Reed²

¹ *Centro Singular de Investigación en Tecnoloxías da Información (CiTIUS), Universidade de Santiago de Compostela, Santiago de Compostela, Spain*

`martin.pereira@usc.es`

² *Centre for Argument Technology (ARG-Tech), University of Dundee, Dundee, Scotland (United Kingdom)*

`c.a.reed@dundee.ac.uk`

Paraphrasing is usually defined as expressing one thing in other words. For instance, let us consider the following statements (<http://www.aifdb.org/search>; ArgMap 2134): (1) Anne's misunderstood the *nature* of housing benefit; and, (2) Anne's *completely* misunderstood the *point* of housing benefit. Although they differ in three words (*completely*, *point* and *nature*), they express approximately the same message. This phenomenon has been deeply studied by computational linguistics, appearing different methods devoted to recognize, generate or extract paraphrases [1].

G. Lakoff [2] introduced the concept of *hedge*, which, from our point of view, plays a essential role in paraphrasing. It is defined as a word or phrases that are capable of transforming the meaning of predicates with the objective of making them fuzzier or less fuzzy. He distinguishes four types of hedges, which we sort according to their *fuzzifier* capability as follows: Strictly speaking $\prec$ Technically speaking $\prec$ Loosely speaking $\prec$ Regular speaking. For instance, in (2), the term “completely” is a hedge that denotes *strictly speaking*; i.e., the utterance must be interpreted in a rigorous way.

According to Lakoff, the semantic of a concept have four different meaning components: (i) definitional, (ii) primary, (iii) secondary and (iv) characteristic though incidental. Components (i), (ii) and (iii) are enough to confer conceptual category membership while criterion iv) non. These components are used for defining the meaning of the hedges; e.g., *technically speaking* means that exclusively *definitional* component is considered while *strictly speaking* consider both *definitional* and *primary* ones. Nonetheless, how to describe the extension of these meaning components is an open question. For instance, which are the items for these meaning components in concepts such as *misunderstand*, *nature* or *point*? In this work, we propose to use WordNet [3] an fuzzy sets for defining them.

Our approach is supported on Lexical Field Theory [4]; i.e., a structuralist conception of word meaning. It introduces the concept of lexical or semantic field, which denotes the set of semantically related items whose meanings are interdependent among them in a given domain. Thus, the study of words is addressed attending to their relations with other words in the same lexical field.

This work was supported in part by the European Regional Development Fund (ERDF/FEDER) under the projects CN2012/151 and GRC2014/030 of the Galician Ministry of Education, and the Spanish Ministry for Economy and Competitiveness under grant TIN2014-56633-C3-1-R.

Thus, the different relationships described in WordNet (synset, hypernyms, etc.) provide us the information for defining the lexical field of a concept. In our proposal, each meaning component is a fuzzy set which kernel elements are: for (i), the synset; for (ii), the direct hyponyms; for (iii), the direct hypernyms and, for (iv), other relationships (meronym, holonym, etc.). For support elements, their degrees of membership are calculated using a semantic similarity measure [5].

Thus, we define the four different hedges proposed by Lakoff using fuzzy connectives (in this case, the T-norm *minimum*). This provide us a method for recognizing paraphrases (and eventually generate them) assuming that two statements are paraphrases if they are two compatible alternative ways for expressing the same meaning according to the given hedges. Table 1 summarizes the representation process for *nature* (1) and *point* (2) using Rada measure [5] and the measure used to obtain the paraphrasing assessment (Tversky's measure [5]), obtaining 0.42 for statements (1) and (2). This result may seem too low at a first view; however, a deep analysis of the statements, lead us to consider that in (1), technically speaking, *nature* refers to the main characteristics of housing benefits, while, in statement (2), strictly speaking, *point* refers to the relevant characteristics of housing benefits in Ann's case.

Table 1: Semantic field of *point* and *nature*.

TERMS	CATEGORY	ITEMS
Nature (gloss 1)	Def	nature/1, quality/0.5, attribute/0.33, characteristic/0.33, abstraction/0.25, entity/0.2
	Prim	The same as Def because <i>nature</i> does not have hyponyms.
	Sec	quality/1, nature/0.5, attribute/0.5, characteristic/0.5, abstraction/0.33, entity/0.25
	Other	characteristic/1, quality/0.5, nature/0.33, attribute/0.33, abstraction/0.25, entity/0.2
Point (gloss 23)	Def	selling point/1, point/0.5, characteristic/0.33, quality/0.25, attribute/0.2, abstraction/0.16, entity/0.14
	Prim	point/1, selling point/0.5, characteristic/0.5, quality/0.33, attribute/0.25, abstraction/0.2, entity/0.16
	Sec	point/0.5, selling point/0.5, characteristic/1, quality/0.5, attribute/0.33, abstraction/0.25, entity/0.2
	Other	The same as Sec because <i>point</i> does not have meronyms or holonyms
Technically speaking (nature):		nature/1, quality/0.5, attribute/0.33, characteristic/0.33, abstraction/0.25, entity/0.2
Strictly speaking (point):		point/0.5, selling point/0.5, characteristic/0.33, quality/0.35, attribute/0.2, abstraction/0.16, entity/0.14
$k = 0.5, sim_{Tversky}(nature, point) =$		$\frac{\sum -countmin(nature, point)}{\sum -countmin(nature, point) + 0.5(\frac{\sum -countmin(nature)}{\sum -countmin(point)}) + 0.5(\frac{\sum -countmin(point)}{\sum -countmin(nature)})} = 0.42$

References

- [1] I. Androutsopoulos and P. Malakasiotis. A survey of paraphrasing and textual entailment methods. *Journal of Artificial Intelligence Research*, 38:135 – 187, 2010.
- [2] G. Lakoff. Hedges: A study in meaning criteria and the logic of fuzzy concepts. *Journal of Philosophical Logic*, 2:458–508, 1973.
- [3] G. A. Miller. Wordnet: A lexical database for english. *Communications of the ACM*, 38(11):39–41, 1995.
- [4] A. Lehrer. *Semantic Fields and Lexical Structure*. North-Holland, 1974.
- [5] M. A. Hadj Taieb, M. Ben Aouicha, and A. Ben Hamadou. Ontology-based approach for measuring semantic similarity. *Engineering Applications of Artificial Intelligence*, 36:238–261, 2014.

Fuzzy sets and natural language generation. What to make out of it?

Alejandro Ramos¹, Alberto Bugarín¹ Senén Barro¹

¹ *Centro Singular de Investigación en Tecnoloxías da Información (CiTIUS),
Universidade de Santiago de Compostela*
{alejandro.ramos,alberto.bugarin.diz,senen.barro}@usc.es

The creation of systems which are able to automatically produce texts as if they were written by humans is a problem which has been actively researched in the last three decades by the discipline of natural language generation (NLG). Many NLG systems emerged over this period for very different purposes (e.g. narrate stories, dialog in interactive systems, data summarization, and text summarization, among others) and application domains (health, environmental information systems, industry, project management, education...).

At the same time, the use of fuzzy sets as a tool for obtaining meaningful linguistic information from data which also supports uncertainty management allowed the emergence of the linguistic description of data subfield (primarily based on the concept of fuzzy quantified sentence), a discipline with a similar purpose to that of NLG approaches which generate texts from numeric data (also known as data-to-text approaches): to provide an understandable interface between the data and the human users in the form of information expressed in terms of natural language. However, one is aimed at the production of actual texts (data-to-text), while linguistic description of data remains at a more conceptual level.

Both NLG and fuzzy fields were developed in a parallel and independent way, totally unaware of each other until recent times, when interest in a potential relationship between these disciplines was raised by researchers from the linguistic description of data subfield [1]. This was mainly motivated by the need for a means to use effectively the fuzzy techniques developed in linguistic description in practical application domains. As a result, many researchers in this discipline have proposed practical cases where text generation has also been considered.

While it may seem that linguistic description of data requires NLG to be viable in applied use cases, and that this interest is not given in the opposite direction, it would be unfair to state such thing. Actually, one of the open challenges in NLG is the management of imprecision and uncertainty (or vagueness, which is also a very frequent word in this domain). In this regard, fuzzy sets and derived applications such as linguistic description of data are intuitively appropriate for this task. It is not clear why such approaches haven't been explored thoroughly yet, but this unawareness may be partially explained by the opposition between the traditional theoretical nature of the fuzzy field, more focused on its logical aspects, and the more applied nature of NLG, focused on linguistic problems of a more empirical weight.

Nowadays, the interest on this relationship has increased and work has been made from both sides on this respect [2, 3, 4, 5, 6]. Even so, it is not clear yet what form (or how many forms) this relationship will take, as natural language generation involves several tasks which can potentially be imbued with fuzzy techniques. In this regard, and taking into account the different perspectives and backgrounds from both disciplines, we will:

- Discuss how researchers in fuzzy sets (and especially in linguistic description of data) can use NLG effectively and the implications of the techniques they usually employ from an NLG point of view. In this regard, we will focus on: *i)* Exploring where fuzzy quantified statements can be placed in terms of NLG tasks, *ii)* Providing illustrative examples and use cases taken from the literature.
- Likewise, we will also discuss how fuzzy sets can be used for including uncertainty management in several tasks involved in standard NLG systems. More specifically, we will show some examples on how fuzzy set related techniques can be used for: *i)* content determination, *ii)* lexicalization, *iii)* referring expression generation, *iv)* aggregation.
- We will also address the differences between approaches in linguistic summarization of data and NLG regarding evaluation tasks.
- The previous topics will serve to open up the discussion on another important issue involved in the cross-fertilization of both fields, namely the usage of empirical techniques to define fuzzy linguistic terms and expressions in real application contexts.

References

- [1] J. Kacprzyk and S. Zadrozny. Computing with words is an implementable paradigm: Fuzzy queries, linguistic data summaries, and natural-language generation. *Fuzzy Systems, IEEE Transactions on*, 18(3):461–472, June 2010.
- [2] Nicolás Marín and Daniel Sánchez. On generating linguistic descriptions of time series. *Fuzzy Sets and Systems*, 285:6 – 30, 2016. Special Issue on Linguistic Description of Time Series.
- [3] A. Ramos-Soto, A. Bugarín, and S. Barro. On the role of linguistic descriptions of data in the building of natural language generation systems. *Fuzzy Sets and Systems*, 285:31–51, 2016.
- [4] Nicolás Marín and Daniel Sánchez. Fuzzy sets and systems + natural language generation: A step forward in the linguistic description of time series. *Fuzzy Sets and Systems*, 285:1 – 5, 2016. Special Issue on Linguistic Description of Time Series.
- [5] A. Bugarin, N. Marin, D. Sanchez, and G. Trivino. Aspects of quality evaluation in linguistic descriptions of data. In *Fuzzy Systems (FUZZ-IEEE), 2015 IEEE International Conference on*, pages 1–8, Aug 2015.
- [6] Albert Gatt and François Portet. Multilingual generation of uncertain temporal expressions from data: A study of a possibilistic formalism and its consistency with human subjective evaluations. *Fuzzy Sets and Systems*, 285:73 – 93, 2016. Special Issue on Linguistic Description of Time Series.



ISBN: 978-84-608-7831-5



9 788460 878315