

ESTYLF 2025



LIBRO DE ABSTRACTS

XXIII

Congreso Español sobre Tecnologías y Lógica Fuzzy
2-5 Septiembre 2025 • Costa Ballena, Cádiz

Editores:

M. Eugenia Cornejo, Luis Jiménez Linares, Jesús Medina,
Juan Moreno García, Manuel Ojeda-Aciego

Editores Asociados:

Roberto G. Aragón, Fernando Chacón-Gómez, David Lobo,
Francisco J. Ocaña-Alcázar, Eloísa Ramírez-Poussa

Libro de abstracts de ESTYLF 2025

© M. Eugenia Cornejo, Luis Jiménez Linares, Jesús Medina,
Juan Moreno García, Manuel Ojeda-Aciego, Editores
Roberto G. Aragón, Fernando Chacón-Gómez, David Lobo,
Francisco J. Ocaña-Alcázar, Eloísa Ramírez-Poussa, Editores
Asociados

1ª Edición

Primera publicación: 2025

ISBN: 978-84-09-75858-6

Publicado e impreso por:

Universidad de Cádiz (Dept. Matemáticas), España.

Comité de Programa

Javier Albusac	Universidad de Castilla-La Mancha
Jesús Alcalá-Fernández	Universidad de Granada
Cristina Alcalde	Universidad del País Vasco, UPV/EHU
Jose M. Alonso	Universidad de Santiago de Compostela
Sergio Alonso	Universidad de Granada
Edurne Barrenechea	Universidad de Navarra
Senén Barro	Universidad de Santiago de Compostela
Fernando Bobillo	Universidad de Zaragoza
Alberto Bugarín-Diz	Universidad de Santiago de Compostela
Humberto Bustince	Universidad Pública de Navarra
Francisco Javier Cabrerizo	Universidad de Granada
José M. Cadenas	Universidad de Murcia
Tomas Calvo	Universidad de Alcalá
Pablo Carmona	Universidad de Extremadura
Juan Luis Castro	Universidad de Granada
Pablo Cordero	Universidad de Málaga
Óscar Cordón	Universidad de Granada
Maria Eugenia Cornejo Piñero	Universidad de Cádiz
Susana Cubillo	Universidad Politécnica de Madrid
Rocío De Andrés	Universidad de Salamanca
Maria José del Jesús	Universidad de Jaén
Miguel Delgado	Universidad de Granada
Susana Díaz	Universidad de Oviedo
Jorge Elorza	Universidad de Navarra
Javier Fernández	Universidad Pública de Navarra
Mikel Galar	Universidad Pública de Navarra
José Luis García Lapresta	Universidad de Valladolid
Luis Garmendia	Universidad Complutense de Madrid
Lluís Godo	Artificial Intelligence Research Institute, IIIA – CSIC
Daniel Gómez	Universidad Complutense de Madrid
Antonio González	Universidad de Granada
Luis Jiménez Linares	Universidad de Castilla-La Mancha
María Teresa Lamata	Universidad de Granada
Bonifacio Llamazares	Universidad de Valladolid
David Lobo Palacios	Universidad de Cádiz
Carlos López Molina	Universidad Publica de Navarra
Nicolás Miguel Madrid Labrador	Universidad de Cádiz
Luis Magdalena	Universidad Politécnica de Madrid
María J. Martín Bautista	Universidad de Granada
Luis Martínez	Universidad de Jaén
Sebastià Massanet	Universidad de las Islas Baleares

Francisco Mata	Universidad de Jaén
Jesús Medina	Universidad de Cádiz
Jose M. Molina	Universidad Carlos III de Madrid
Javier Montero	Universidad Complutense de Madrid
Susana Montes	Universidad de Oviedo
Juan Moreno García	Universidad de Castilla-La Mancha
Manuel Ojeda-Aciego	Universidad de Málaga
Jose Ángel Olivas	Universidad de Castilla-La Mancha
Antonio Peregrín	Universidad de Huelva
Héctor Pomares	Universidad de Granada
Ana Pradera	Universidad Rey Juan Carlos
Eloísa Ramírez Poussa	Universidad de Cádiz
Jordi Recasens	Universitat Politècnica de Catalunya
Juan Vicente Riera	Universidad de las Islas Baleares
Adolfo Rodríguez	Universidad de León
Rosa Rodríguez	Universidad de Jaén
Francisco P. Romero	Universidad de Castilla La Mancha
Daniel Sánchez	Universidad de Granada
Luciano Sánchez	Universidad de Oviedo
José Antonio Sanz Delgado	Universidad Pública de Navarra
Jesús Serrano-Guerrero	Universidad de Castilla-La Mancha
Miguel Ángel Sicilia	Universidad de Alcalá
Vicenç Torra	Universidad de Skövde
Aida Valls	Universitat Rovira i Virgili
José Luis Verdegay	Universidad de Granada

Programa general del

XXIII Congreso Español sobre Tecnologías y Lógica Fuzzy

2-5 septiembre, 2025. Costa Ballena, España



Universidad
de Cádiz

Departamento
de Matemáticas

DIGFORASP



EUSFLAT



EUROPEAN SOCIETY
FOR FUZZY LOGIC
AND TECHNOLOGY



MARTES 2	
Lugar: Hotel Barceló Costa Ballena	
12:50	Apertura del registro
14:00–15:30	Almuerzo
15:30–16:50	Sesión 1: Toma de decisiones Moderador: Luis Martínez
	<i>Construyendo Conjuntos Difusos Personalizados mediante el Método de la Baraja de Cartas</i> Diego García-Zamora, Bapi Dutta, José Rui Figueira, Luis Martínez
	<i>Un marco basado en la Lógica difusa para el apoyo en la toma de decisiones mediante una variante de modelos TSK</i> David Muñoz Valero, Juan Moreno García, Julio Alberto López Gómez, Enrique Adrian Villarrubia Martín
	<i>An overview of Fuzzychain: a fuzzy logic-based consensus algorithm for Blockchain</i> Bruno Ramos-Cruz, Javier Andreu-Perez, Francisco J. Quesada-Real, Luis Martínez
	<i>Enfoque de inferencia difusa en la toma de decisiones para medir indicadores emocionales</i> Miguel Andrés, Matilde Santos, Victoria López, Diego Riofrío Luzcando
	<i>Three-Way Group Decision-Making With Personalized Numerical Scale of Comparative Linguistic Expression: An Application to Traditional Chinese Medicine</i> Yaya Liu, Lina Zhu, Rosa M. Rodríguez, Yiyu Yao, Luis Martínez
	<i>Marco de toma de decisiones multicriterio con expresiones lingüísticas flexibles y conjuntos nube aproximados multigranulares: aplicación en evaluación de proyectos de Internet de la Energía</i> Han Wang, Yanbing Ju, Carlos Porcel
	<i>Un Modelo Difuso de Consenso Para Problemas de Toma de Decisiones en Grupo Basado en Computación Granular y Blockchain</i> Juan Carlos González-Quesada, Ignacio Javier Pérez, Francisco Mata, Ramón Alberto Carrasco, Enrique Herrera Viedma, Francisco Javier Cabrerizo
	<i>Un nuevo modelo de consenso que unifica y coordina modelos generativos de lenguaje y experiencia humana</i> F.J. Cabrerizo, M.J. Cobo, J. Bernabe Moreno, E. Herrera Viedma, I.J. Perez
16:50–17:00	Descanso
17:00–17:40	Sesión 2: Lógica Difusa en Deep Learning Moderador: Juan Moreno
	<i>Pooling over-time no-simétrico mediante funciones de pseudo-grouping en redes neuronales convolucionales</i> M. Ferrero Jaurrieta, L. De Miguel, C. Lopez Molina, H.Bustince, A. Cruz, B. Bedregal, R. Paiva, Z. Takáč
	<i>Pooling global de características en Redes Neuronales Convolucionales mediante funciones de agregación basadas en desviaciones moderados</i> Iosu Rodríguez Martínez, Jana Špirková, Mikel Ferrero Jaurrieta, Humberto Bustince, Francisco Javier Fernández
	<i>Modelado de Perfiles Difusos para la Interpretación de Redes Neuronales en la Recomendación de Vinos</i> Arturo Fernández-Mora, Juan Moreno-García, David Muñoz-Valero, Luis Martínez
	<i>Generalizing a construction of non-strong fuzzy metrics from metrics and studying their induced topology</i> Olga Grigorenko, Juan-José Miñana, Simona Talia
17:40–19:00	Sesión de pósteres Moderador: Jesús Medina
20:30	Bienvenida a los participantes

MIÉRCOLES 3 Lugar: Hotel Barceló Costa Ballena	
8:50–9:00	Inauguración
9:00–11:00	Sesión 3: Aplicaciones de la Lógica Difusa Moderador: Macarena Espinilla
	<i>On the use of Aggregation Operators to improve Human Identification using Dental Records</i> Antonio D. Villegas-Yeguas, Guillermo R-García, Tzipi Kahana, Oscar Ibáñez, Oscar Cordón
	<i>Generando explicaciones fiables por medio de un modelo de lenguaje enriquecido por sistemas basados en reglas borrosas</i> Pablo Miguel Pérez Ferreiro, Alejandro Catala, Alberto Bugarín Diz, Jose M. Alonso Moral
	<i>Fuzzy processing applied to improve multimodal sensor data fusion to discover frequent behavioral patterns for smart healthcare</i> Carlos Fernández-Basso, David Díaz-Jiménez, José L. López Ruiz, Macarena Espinilla
	<i>A discrete Median filter for RGB images</i> Pablo Flores-Vidal, Daniel Gómez, Luis Magdalena
	<i>Modelado Interpretable de Datos de Glucosa</i> Juan F. Gaitán-Guerrero, José L. López, Macarena Espinilla Estévez, David Díaz-Jiménez, Carmen Martínez-Cruz
	<i>Sistemas Difusos Lingüísticos que Evolucionan para Regresión</i> Javier Martín Moreno, Francisco Alfredo Márquez Hernández, Antonio Peregrín Rubio
	<i>Estudio comparativo de funciones de agregación inspiradas en Choquet para la reducción de dimensionalidad en imágenes con ruido</i> Joaquin Arellano, Humberto Bustince, Francisco Javier Fernández, Carlos Guerra, Cedric Marco-Detchart
	<i>Integración de datos clínicos y modelos de lenguaje para resúmenes lingüísticos de historias médicas</i> Andrea Morales Garzón, Raúl Martínez Alonso, Hossam El Amraoui Leghzali, María J. Martín Bautista, Carlos Fernández Basso
	<i>APIA: Sistema Borroso para la estimación del riesgo de incumplimiento en la ejecución de contratos públicos</i> Andrés Montoro-Montarroso, Guadalupe Plaza, Jose A. Olivas, Jesus Serrano-Guerrero, Jared D.T. Guerrero-Sosa, Francisco P. Romero
	<i>Consistencia semántica y patrones de estilo en el tiempo para la detección de bots en redes sociales</i> Salvador López-Joya, Jose A. Díaz-García, M. Dolores Ruiz, María J. Martín-Bautista
	<i>Metodología híbrida para el aprendizaje y la evaluación de modelos de lenguaje bajo esquemas de incertidumbre difusa</i> Juan F. Gaitán-Guerrero, Carmen Martínez-Cruz, Macarena Espinilla, David Díaz-Jiménez, José L López
	<i>Knowledge base tuning by particle swarm optimization for centralized wireless sensor network clustering algorithm</i> Juan-Luis Herreros Bódalo, Jose Enrique Muñoz-Exposito, Antonio Jesús Yuste Delgado, Alicia Triviño Cabrera, Macarena Espinilla-Estévez, David Díaz Jiménez, Juan Carlos Cuevas Martínez
11:00–12:00	Descanso para pósteres y café

12:00–14:00	Sesión 4: Fundamentos de la Lógica Difusa I Moderador: Manuel Ojeda
	<i>Uso de las dCF-integrales en sistemas de clasificación basados en reglas difusas para abordar problemas no balanceados</i> José Antonio Sanz, Mikel Sesma-Sara
	<i>An order-oriented approach to scoring hesitant fuzzy elements</i> Luis Merino, Gabriel Navarro, Carlos Salvatierra, Evangelina Santos
	<i>On the non-falsity and threshold preserving variants of MTL logics</i> Francesc Esteva, Joan Gispert, Lluís Godó
	<i>Una Propuesta de Función Generadora de Funciones Opuestas para Conjuntos Difusos Pareados</i> Javier Bonilla
	<i>Ecuaciones bipolares de relaciones difusas con negación suelo</i> David Lobo Palacios, Pablo López Molina
	<i>Funciones de agregación inspiradas en la integral Choquet</i> Humberto Bustince, Javier Fernández, Julio Lafuente, Graçaliz Dimuro
	<i>Sobre funciones de implicación basadas en órdenes admisibles en el conjunto de números borrosos discretos</i> M. González-Hidalgo, S. Massanet, A. Mir, J. V. Riera, L. De Miguel
	<i>Random choice of referring expressions with fuzzy properties based on user preferences</i> Nicolás Marín, Gustavo Rivas-Gervilla, Daniel Sánchez
	<i>Estudio de la monotonía de las medidas de similitud basadas en incrustaciones</i> Sara Suárez-Fernández, Agustina Bouchet, Susana Díaz-Vazquez, Susana Montes
	<i>Converxidad en conjuntos difusos typical hesitant aplicada a la toma de decisiones</i> Pedro Huidobro, Pedro Alonso, Vladimir Janiš, Susana Montes
	<i>Eficiencia normalizada de conjuntos de reglas de decisión</i> Fernando Chacón-Gómez, M. Eugenia Cornejo, Jesús Medina
	<i>Caracterización del rendimiento de jugadores de balonmano mediante lógica difusa y algoritmos bioinspirados</i> Eusebio Angulo Sánchez-Herrera, Francisco P. Romero, Julio Alberto López-Gómez
14:00–15:10	Almuerzo
15:10–16:10	Sesión para pósteres y café
16:10–20:30	Excursión
20:30	Noche temática. Noche ibicenca

JUEVES 4		
Lugar: Hotel Barceló Costa Ballena		
9:00–9:45	Sesión Plenaria – Juan Bernabé Moreno. Director de IBM Research Europe para Irlanda y Reino Unido Título: Cómo la IA Generativa está cambiando nuestra manera de hacer ciencia Moderadora: María José del Jesús.	
9:45–10:00	Descanso	
10:00–11:20	Mesa redonda 60 años de lógica difusa. ¿Hacia dónde debe ir? En el marco de la IA Alberto J. Bugarín, María José del Jesús, José Ángel Olivas Moderador: Javier Montero	
11:20–11:30	Descanso	
11:30–12:15	Sesión Plenaria – Iñaki Pinillos Resano. Ex-Director de Nasertic. Consejero y asesor independiente Título: La IA en la industria. Del algoritmo al producto Moderador: Humberto Bustince	
12:15–12:30	Descanso	
12:30–14:00	Mesa redonda Industria e Inteligencia Artificial Empresas: Airbus, Navantia, Grupo Energético de Puerto Real Moderador: Jesús Medina	
14:00–15:30	Almuerzo	
15:30–17:00	Tu tesis en 5 minutos Moderadores: Sara Suárez, Óscar Cerdón	Sesión de proyectos Moderadora: María Dolores Ruiz
17:00–18:00	Reuniones grupos Moderador: David Lobo	Sesión de pósteres y café
20:30	Cena de gala	

VIERNES 5 Lugar: Hotel Barceló Costa Ballena	
9:30–10:15	Sesión Plenaria – Ricardo Baeza Yates. KTH, Sweden; UPF, Spain; and Univ. of Chile Título: IA responsable: riesgos y recomendaciones Moderador: Alberto Bugarín Diz
10:15–10:25	Descanso
10:25–11:45	Sesión 5: Fundamentos de la Lógica Difusa II Moderador: María Eugenia Cornejo
	<i>On direct systems of implications with graded attributes</i> Manuel Ojeda-Hernández, Domingo López-Rodríguez
	<i>Teorías, modelos y bases de implicaciones de atributos en retículos de conceptos multi-adjuntos</i> M. Eugenia Cornejo, Jesús Medina, Francisco José Ocaña
	<i>Suma de contextos basada en enlaces multiadjuntos</i> Roberto G. Aragón, Jesús Medina, Samuel Molina-Ruiz
	<i>Contextos formales fibrados</i> Víctor Ramos-González, Joaquín Borrego-Díaz, Fco. Félix Lara-Martín
	<i>Definición de sistemas de inferencia basados en la f-inclusión</i> Carolina Díaz-Montarroso, Nicolás Madrid, Eloísa Ramírez-Poussa
	<i>Influencia probabilística del número de votantes en el muestreo mediante el modelo de Mallows</i> Mario Villar, Noelia Rico, Irene Díaz
	<i>Representativeness systems and coverage indexes</i> Inmaculada Gutiérrez, J. Tinguaro Rodríguez, Xabier González, Daniel Gómez, Javier Montero, Humberto Bustince
	<i>Salvando ratones: ChenseNet121, una arquitectura de aprendizaje profundo para la estimación de toxicidad aguda</i> Enol Junquera Álvarez, Beatriz Remeseiro, Ferdinando Febbraio, Irene Díaz
11:45–12:10	Descanso para pósteres y café
12:10–13:30	Mesa redonda 60 años de lógica difusa. ¿Hacia dónde debe ir? En el marco de las matemáticas Humberto Bustince, Inés Couso, Luis Magdalena, Moderador: Manuel Ojeda Aciego
13:30–16:00	Clausura

Índice general

Construyendo Conjuntos Difusos Personalizados mediante el Método de la Baraja de Cartas.....	1
<i>Diego García-Zamora, Babi Dutta, José Rui Figueira y Luis Martínez</i>	
Un marco basado en la Lógica difusa para el apoyo en la toma de decisiones mediante una variante de modelos TSK.....	3
<i>David Muñoz-Valero, Juan Moreno-García, Julio Alberto López-Gómez y Enrique Adrian Villarubia-Martin</i>	
An overview of Fuzzychain: a fuzzy logic-based consensus algorithm for Blockchain	4
<i>Bruno Ramos-Cruz, Javier Andreu-Perez, Francisco J. Quesada-Real y Luis Martinez</i>	
Enfoque de inferencia difusa en la toma de decisiones para medir indicadores emocionales.....	7
<i>Miguel Andrés, Matilde Santos, Victoria López y Diego Riofrío-Luzcando</i>	
Three-Way Group Decision-Making With Personalized Numerical Scale of Comparative Linguistic Expression: An Application to Traditional Chinese Medicine	9
<i>Yaya Liu, Lina Zhu, Rosa M. Rodríguez, Yiyu Yao y Luis Martínez</i>	
Marco de toma de decisiones multicriterio con expresiones lingüísticas exibles y conjuntos nube aproximados multigranulares: aplicación en evaluación de proyectos de Internet de la Energía.....	11
<i>Han Wang, Yanbing Ju y Carlos Porcel</i>	
Un Modelo Difuso de Consenso Para Problemas de Toma de Decisiones en Grupo Basado en Computación Granular y Blockchain	13
<i>Juan Carlos González-Quesada, Ignacio Javier Pérez, Francisco Mata, Ramón Alberto Carrasco, Enrique Herrera-Viedma y Francisco Javier Cabrerizo</i>	
Un nuevo modelo de consenso que unifica y coordina modelos generativos de lenguaje y experiencia humana	15

<i>F.J. Cabrerizo, M.J. Cobo, J. Bernabe-Moreno, E. Herrera-Viedma e I.J. Perez</i>	
Salvando ratones: ChenseNet121, una arquitectura de aprendizaje profundo para la estimación de toxicidad aguda	17
<i>Enol Junquera Álvarez, Beatriz Remeseiro, Ferdinando Febbraio e Irene Díaz</i>	
Pooling over-time no-simétrico mediante funciones de pseudo-grouping en redes neuronales convolucionales	19
<i>M. Ferrero-Jaurrieta, L. De Miguel, C. Lopez-Molina, H. Bustince, A. Cruz, B. Bedregal, R. Paiva y Z. Takáč</i>	
Generando explicaciones ables por medio de un modelo de lenguaje enriquecido por sistemas basados en reglas borrosas	21
<i>Pablo Miguel Perez-Ferreiro, Alejandro Catala, Alberto Bugarín-Diz y Jose M. Alonso-Moral</i>	
Pooling global de características en Redes Neuronales Convolucionales mediante funciones de agregación basadas en desviaciones moderados	25
<i>Iosu Rodriguez-Martinez, Jana Špírková, Mikel Ferrero-Jaurrieta, Humberto Bustince y Francisco Javier Fernández</i>	
Modelado de Perles Difusos para la Interpretación de Redes Neuronales en la Recomendación de Vinos	25
<i>Arturo Fernández-Mora, Juan Moreno-Garcia, David Muñoz-Valero y Luis Martínez</i>	
On the use of Aggregation Operators to improve Human Identification using Dental Records	27
<i>Antonio D. Villegas-Yeguas, Guillermo R-García, Tzipi Kahana, Oscar Ibáñez y Oscar Cordón</i>	
Caracterización del rendimiento de jugadores de balonmano mediante Lógica difusa y algoritmos bioinspirados	30
<i>Eusebio Angulo Sánchez-Herrera, Francisco P. Romero y Julio Alberto-Gómez</i>	
Fuzzy processing applied to improve multimodal sensor data fusion to discover frequent behavioral patterns for smart healthcare	32
<i>Carlos Fernández-Basso, David Díaz-Jiménez, José L. López Ruiz y Macarena Espinilla</i>	
A discrete Median lter for RGB images	34
<i>Pablo Flores-Vidal, Daniel Gómez y Luis Magdalena</i>	
Modelado Interpretable de Datos de Glucosa	36
<i>Juan F. Gaitán-Guerrero, José L. López, Macarena Espinilla Estévez, David Díaz-Jiménez y Carmen Martínez-Cruz</i>	
Sistemas Difusos Lingüísticos que Evolucionan para Regresión	38
<i>Javier Martín Moreno, Francisco Alfredo Márquez Hernández y Antonio Peregrín Rubio</i>	

Estudio comparativo de funciones de agregación inspiradas en Choquet para la reducción de dimensionalidad en imágenes con ruido	40
<i>Joaquín Arellano, Humberto Bustince, Francisco Javier Fernández, Carlos Guerra y Cedric Marco-Detchart</i>	
Integración de datos clínicos y modelos de lenguaje para resúmenes lingüísticos de historias médicas	41
<i>Andrea Morales-Garzón, Raúl Martínez-Alonso, Hossam El Amraoui y Maria J. Martin-Bautista</i>	
APIA: Sistema Borroso para la estimación del riesgo de incumplimiento en la ejecución de contratos públicos	43
<i>Andrés Montoro-Montarroso, Guadalupe Plaza, Jose A. Olivas, Jesus Serrano, Jared D.T. Guerrero-Sosa y Francisco P. Romero</i>	
Consistencia semántica y patrones de estilo en el tiempo para la detección de bots en redes sociales	45
<i>Salvador Lopez-Joya, Jose A. Diaz-Garcia, M. Dolores Ruiz y Maria J. Martin-Bautista</i>	
Metodología híbrida para el aprendizaje y la evaluación de modelos de lenguaje bajo esquemas de incertidumbre difusa	47
<i>Juan F. Gaitán-Guerrero, Carmén Martínez-Cruz, Macarena Espinilla, David Díaz-Jiménez y José L. López</i>	
Knowledge base tuning by particle swarm optimization for centralized wireless sensor network clustering algorithm	49
<i>Juan-Luis Herreros-Bodalo, Jose-Enrique Muñoz-Exposito, Antonio-Jesús Yuste-Delgado, Alicia Triviño-Cabrera, Macarena Espinilla-Estevez, David Diaz-Jimenez y Juan Carlos Cuevas Martínez</i>	
Uso de las dCF-integrales en sistemas de clasificación basados en reglas difusas para abordar problemas no balanceados	51
<i>José Antonio Sanz y Mikel Sesma-Sara</i>	
An order-oriented approach to scoring hesitant fuzzy elements	53
<i>Luis Merino, Gabriel Navarro, Carlos Salvatierra y Evangelina Santos</i>	
On the non-falsity and threshold preserving variants of MTL logics	55
<i>Francesc Esteva, Joan Gispert y Lluís Godo</i>	
Una Propuesta de Función Generadora de Funciones Opuestas para Conjuntos Difusos Pareados	57
<i>Javier Bonilla</i>	
Representativeness systems and coverage indexes	59
<i>Inmaculada Gutiérrez, J. Tinguaro Rodríguez, Xabier González, Daniel Gómez, Javier Montero y Humberto Bustince</i>	
Funciones de agregación inspiradas en la integral Choquet	61
<i>Humberto Bustince, Javier Fernández, Julio Lafuente y Graçaliz Dimuro</i>	

Sobre funciones de implicación basadas en órdenes admisibles en el conjunto de números borrosos discretos	63
<i>M. González-Hidalgo, S. Massanet, A. Mir, J. V. Riera y L. De Miguel</i>	
Random choice of referring expressions with fuzzy properties based on user preferences	65
<i>Nicolás Marín, Gustavo Rivas-Gervilla y Daniel Sánchez</i>	
Estudio de la monotonía de las medidas de similitud basadas en incrustaciones	67
<i>Sara Suarez-Fernandez, Agustina Bouchet, Susana Diaz-Vazquez y Susana Montes</i>	
Convexidad en conjuntos difusos typical hesitant aplicada a la toma de decisiones	69
<i>Pedro Huidobro, Pedro Alonso, Vladimir Janiš y Susana Montes</i>	
Eficiencia normalizada de conjuntos de reglas de decisión	71
<i>Fernando Chacón-Gómez, M. Eugenia Cornejo y Jesús Medina</i>	
Ecuaciones bipolares de relaciones difusas con negación suelo	73
<i>David Lobo Palacios y Pablo López Molina</i>	
On direct systems of implications with graded attributes	75
<i>Manuel Ojeda-Hernández y Domingo López-Rodríguez</i>	
Teorías, modelos y bases de implicaciones de atributos en retículos de conceptos multiadjuntos	77
<i>M. Eugenia Cornejo, Jesús Medina y Francisco José Ocaña</i>	
Suma de contextos basada en enlaces multiadjuntos	79
<i>Roberto G. Aragón, Jesús Medina y Samuel Molina-Ruiz</i>	
Contextos formales fibrados	81
<i>Víctor Ramos-González, Joaquín Borrego-Díaz y Fco. Félix Lara-Martín</i>	
Definición de sistemas de inferencia basados en la f-inclusión	83
<i>Carolina Díaz-Montarroso, Nicolás Madrid y Eloísa Ramírez-Poussa</i>	
Influencia probabilística del número de votantes en el muestreo mediante el modelo de Mallows	85
<i>Mario Villar, Noelia Rico e Irene Díaz</i>	
Generalizing a construction of non-strong fuzzy metrics from metrics and studying their induced topology	87
<i>Olga Grigorenko, Juan-José Miñana y Simona Talia</i>	

Construyendo Conjuntos Difusos Personalizados mediante el Método de la Baraja de Cartas

Diego García-Zamora*, Bapi Dutta†, José Rui Figueira‡, y Luis Martínez†

*Departamento de Matemáticas, Universidad de Jaén, España. Email: dgzamora@ujaen.es

†Departamento de Informática, Universidad de Jaén, España. Emails: bdutta@ujaen.es, martin@ujaen.es

‡CEGIS, Instituto Superior Técnico, Universidade de Lisboa, Portugal. Email: figueira@tecnico.ulisboa.pt

Abstract—Este trabajo resume el trabajo *The deck of cards method to build interpretable fuzzy sets in decision-making* [1], en el que se presenta una metodología socio-técnica para construir conjuntos difusos personalizados en problemas de decisión multicriterio. Se propone un proceso co-constructivo entre el analista y el decisor basado en el Método de la Baraja de Cartas (Deck of Cards Method), permitiendo generar funciones de pertenencia denominadas DoC-MFs que reflejan la semántica individual del decisor incluso con escalas heterogéneas. Se detalla la metodología en tres fases: construcción de una función de valor, identificación del soporte y núcleo, y modelado de los tramos laterales de la función de pertenencia.

Index Terms—Toma de decisiones multicriterio, lógica difusa, funciones de pertenencia, método de la baraja de cartas, semántica individualizada.

I. INTRODUCCIÓN

La teoría de conjuntos difusos, introducida por Zadeh [2], ha sido ampliamente utilizada en la toma de decisiones bajo incertidumbre, especialmente en contextos donde las preferencias u opiniones de expertos son vagas o imprecisas. En esta teoría, las funciones de pertenencia constituyen el elemento central para representar grados de satisfacción respecto a los conceptos difusos [3], [4]. No obstante, una crítica recurrente es que estas funciones suelen definirse sin la participación activa del decisor, recurriendo a suposiciones arbitrarias o formas predefinidas como funciones paramétricas [5].

Esto conlleva una pérdida de precisión y puede comprometer tanto la validez semántica del modelo como los resultados obtenidos. Diversos trabajos han resaltado la necesidad de enfoques que permitan capturar la semántica individualizada de los decisores [6], [7], integrando tanto información cuantitativa como cualitativa [8].

En este contexto, proponemos una metodología que permite construir funciones de pertenencia personalizadas mediante el *Método de la Baraja de Cartas* [9], una técnica de comparación por pares utilizada con éxito en decisión multicriterio [1]. La metodología presentada se enmarca en una perspectiva socio-técnica, ya que promueve la co-construcción de la función de pertenencia entre el analista y el decisor, garantizando que la estructura resultante refleje de forma explícita las opiniones y ambigüedades propios del decisor.

II. METODOLOGÍA PROPUESTA

La metodología consta de tres etapas principales, todas ellas guiadas por la interacción estructurada entre analista y

decisor. Como resultado, obtendremos las llamadas Deck-of-Cards Membership Functions (DoC-MF).

A. Etapa 1: Construcción de la escala de valor

Se parte de una escala original, que puede ser numérica, ordinal o lingüística. Utilizando el Método de la Baraja de Cartas, se solicita al decisor que ordene los niveles de rendimiento (o etiquetas) para cada criterio en función de su preferencia, colocando tarjetas en blanco entre ellos para expresar la intensidad relativa entre niveles. A partir de estas comparaciones, se construye una escala cardinal de valor normalizada en el intervalo $[0,1]$, donde 0 representa el nivel menos preferido y 1 el más preferido.

B. Etapa 2: Identificación del soporte y núcleo de la función de pertenencia

Para cada etiqueta construiremos un conjunto difuso que la modele (por ejemplo, “alto rendimiento”). En primer lugar se determina el intervalo de valores que conforma el núcleo, es decir, aquellos con pertenencia total (valor 1). Luego, se identifican los extremos del soporte: valores para los cuales la pertenencia es cero. Esta información se recoge a través de la consulta directa al decisor o mediante interpretación de las funciones de valor previas.

C. Etapa 3: Construcción de los tramos laterales

Los tramos izquierdo y derecho de la función de pertenencia, que conectan el soporte con el núcleo, se construyen empleando nuevamente el Método de la Baraja de Cartas. Se presentan pares de valores al decisor y se le solicita indicar el número de tarjetas en blanco que representan su distancia semántica en términos de pertenencia al concepto difuso. De este modo, se obtiene una representación lineal a trozos de la función que respeta la percepción del decisor y permite operaciones posteriores (como agregación o difusión).

III. REGLAS COMPUTACIONALES

Una vez construidas las funciones de pertenencia, se pueden aplicar diferentes operaciones computacionales clave en la toma de decisión difusa. Estas operaciones se basan en las propiedades matemáticas de los números difusos lineales por tramos y aseguran la compatibilidad con modelos existentes de agregación y ordenación.

A. Agregación

Las DoC-MF pueden combinarse mediante operadores típicos de la teoría de la decisión difusa, como la media ponderada difusa, o el operador OWA (Ordered Weighted Averaging) [10]. Dado que las funciones son representadas de forma explícita mediante segmentos lineales, estas operaciones pueden implementarse de forma eficiente a través de los α -cortes [1].

B. Comparación y Ordenación

El uso de DoC-MFs permite aplicar reglas de ordenación para seleccionar o priorizar alternativas. Se pueden emplear métricas de distancia entre funciones difusas, reglas basadas en *ranking indices*, como el centroide, el valor esperado, o los órdenes admisibles [11].

IV. EJEMPLO DE CONSTRUCCIÓN DE UNA DoC-MF

A continuación se ilustra cómo construir una función de pertenencia difusa (DoC-MF) para el término lingüístico *acceptable* en el contexto de un cierto criterio.

A. Construcción de la escala de valor

El conjunto de niveles lingüísticos es: md (muy deficiente), d (deficiente), a (aceptable), b (bueno) y mb (muy bueno). Aplicando el Método de la Baraja de Cartas, el decisor proporciona el número de tarjetas en blanco entre los niveles adyacentes:

md [7] d [11] a [6] b [8] mb

A partir de esta información, el número total de unidades es: $r = 7 + 11 + 6 + 8 = 32$ y el valor de cada unidad es $\gamma = \frac{1}{32} = 0.03125$. Así, los valores normalizados asociados a cada nivel son:

- $v(md) = 0.00$
- $v(d) = 7 \cdot \gamma = 0.21875$
- $v(a) = 18 \cdot \gamma = 0.5625$
- $v(b) = 24 \cdot \gamma = 0.75$
- $v(mb) = 1.00$

B. Identificación del núcleo y soporte

A través de la interacción con el decisor, se establece el núcleo de la función de pertenencia para *acceptable* en el intervalo [0.5, 0.6]. El soporte completo se establece en [0.45, 0.80].

C. Modelado del tramo izquierdo y derecho

Al preguntar al decisor sobre el tramo izquierdo de la función de pertenencia, tiene claro que sigue una tendencia lineal. Sin embargo, al preguntar por el tramo derecho, siente ciertos cambios de tendencia. Esta situación se resuelve nuevamente mediante el Método de la Baraja asignando tarjetas en blanco entre los puntos clave del intervalo [0.60, 0.80]. Se obtienen los siguientes puntos de ruptura: (0.60, 1.00), (0.65, 0.81), (0.70, 0.45), (0.75, 0.18) y (0.80, 1.00).

D. Función de pertenencia resultante

La función DoC-MF para *acceptable* queda definida como:

$$\mu(x) = \begin{cases} 0 & \text{if } x \notin [0.45, 0.80] \\ r_{(0.45, 0), (0.50, 1.00)}(x) & \text{if } x \in [0.45, 0.50] \\ 1 & \text{if } x \in [0.50, 0.60] \\ r_{(0.60, 1.00), (0.65, 0.81)}(x) & \text{if } x \in [0.60, 0.65] \\ r_{(0.65, 0.81), (0.70, 0.45)}(x) & \text{if } x \in [0.65, 0.70] \\ r_{(0.70, 0.45), (0.75, 0.18)}(x) & \text{if } x \in [0.70, 0.75] \\ r_{(0.75, 0.18), (0.80, 0.00)}(x) & \text{if } x \in [0.75, 0.80] \end{cases}$$

donde $r_{(a,b),(c,d)}$ es la recta que une los puntos (a, b) y (c, d) .

Obtenemos así una función continua, lineal por tramos, que refleja la percepción del decisor sobre el término *acceptable*.

V. CONCLUSIONES

La metodología propuesta permite la construcción de funciones de pertenencia personalizadas mediante un proceso participativo. Al capturar la semántica individual del decisor, las DoC-MFs facilitan decisiones más precisas y transparentes en contextos complejos. Además, las funciones de pertenencia obtenidas se pueden comparar y agregar eficientemente [1].

AGRADECIMIENTOS

José Rui Figueira está financiado por fondos Portugueses mediante FCT - Foundation for Science and Technology, I.P., con los proyectos UIDB/00097/2025 and UIDP/00097/2025 (CEGIST).

REFERENCES

- [1] D. García-Zamora, B. Dutta, J. R. Figueira, and L. Martínez, "The deck of cards method to build interpretable fuzzy sets in decision-making," *European Journal of Operational Research*, vol. 319, no. 1, pp. 246–262, 2024.
- [2] L. A. Zadeh, "Fuzzy sets," *Information Control*, vol. 8, no. 3, pp. 338–353, 1965.
- [3] R. E. Bellman and L. A. Zadeh, "Decision-making in a fuzzy environment," *Management Science*, vol. 17, no. 4, pp. B–141, 1970.
- [4] D. J. Dubois, *Fuzzy sets and systems: theory and applications*. Academic press, 1980, vol. 144.
- [5] L. Martínez, R. M. Rodríguez, and F. Herrera, *The 2-tuple Linguistic Model*. Springer International Publishing, 2015.
- [6] T. Bilgiç and I. B. Turksen, "Elicitation of membership functions: How far can theory take us?" in *Proceedings of 6th International Fuzzy Systems Conference*, vol. 3. IEEE, 1997, pp. 1321–1325.
- [7] J.-L. Chameau and J. C. Santamarina, "Membership functions I: Comparing methods of measurement," *International Journal of Approximate Reasoning*, vol. 1, no. 3, pp. 287–301, 1987.
- [8] F. Herrera and L. Martinez, "An approach for combining linguistic and numerical information based on the 2-tuple fuzzy linguistic representation model in decision-making," *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, vol. 8, no. 05, pp. 539–562, 2000.
- [9] S. Corrente, J. R. Figueira, and S. Greco, "Pairwise comparison tables within the deck of cards method in multiple criteria decision aiding," *European Journal of Operational Research*, vol. 291, no. 2, pp. 738–756, 2021.
- [10] D. García-Zamora, A. Cruz, F. Neres, R. H. Santiago, A. F. Roldán López de Hierro, R. Paiva, G. P. Dimuro, L. Martínez, B. Bedregal, and H. Bustince, "Admissible owa operators for fuzzy numbers," *Fuzzy Sets and Systems*, vol. 480, p. 108863, Mar. 2024.
- [11] N. Zumelzu, B. Bedregal, E. Mansilla, H. Bustince, and R. Díaz, "Admissible orders on fuzzy numbers," *IEEE Transactions on Fuzzy Systems*, vol. 30, no. 11, pp. 4788–4799, 2022.

Un marco basado en la lógica difusa para el apoyo en la toma de decisiones mediante una variante de modelos TSK

David Muñoz-Valero

*Departamento de Tecnologías y Sistemas de Información
Escuela de Ingeniería Industrial y Aeroespacial
Universidad de Castilla-La Mancha
Toledo, España
david.munoz@uclm.es*

Juan Moreno-García

*Departamento de Tecnologías y Sistemas de Información
Escuela de Ingeniería Industrial y Aeroespacial
Universidad de Castilla-La Mancha
Toledo, España
juan.moreno@uclm.es*

Julio Alberto López-Gómez

*Departamento de Tecnologías y Sistemas de Información
Escuela Superior de Informática
Universidad de Castilla-La Mancha
Ciudad Real, España
julioalberto.lopez@uclm.es*

Enrique Adrián Villarrubia-Martin

*Departamento de Tecnologías y Sistemas de Información
Escuela Superior de Informática
Universidad de Castilla-La Mancha
Ciudad Real, España
enrique.villarrubia@uclm.es*

Abstract—Este trabajo tiene como objetivo abordar la generación automática de modelos de decisión difusos basados en el conocimiento de expertos. Para lograrlo, se propone una variante del Sistema de Inferencia Difusa de Takagi-Sugeno-Kang (TSK FIS) y un marco metodológico capaz de generar modelos de decisión mediante un conjunto de parámetros definidos por expertos. El marco introduce un mecanismo de coocurrencia de variables que captura interacciones, permitiendo a los expertos definir intuitivamente las relaciones entre variables. Este mecanismo garantiza que el experto no necesite un conocimiento profundo de la lógica difusa. Para asegurar el correcto funcionamiento, también se ha modificado el mecanismo de inferencia, permitiendo que las reglas de decisión produzcan resultados robustos e interpretables mediante la obtención acumulativa del valor de salida.

Para probar la propuesta, se ha aplicado a un problema de selección de armas en videojuegos por parte de los jugadores según su perfil. Los resultados experimentales han validado la capacidad del sistema para generar modelos de decisión que se adaptan a los perfiles y objetivos específicos de los usuarios, manteniendo tanto la precisión como la interpretabilidad en la toma de decisiones.

Dependiendo del número de variables y su coocurrencia, el mecanismo puede generar un gran número de reglas. Para mitigar este efecto, los trabajos futuros se centrarán en simplificar el modelo mediante la fusión de reglas y la reducción de variables, actualizando los consecuentes en el proceso mientras se preservan la precisión, la robustez y la claridad. Además, cuando se disponga de conjuntos de datos, se desarrollará un método

de inducción completamente automático para estos modelos de decisión, con el fin de ajustar los parámetros en función de un conjunto de datos de entrada.

Este trabajo es un resumen que describe el artículo “A Framework for the Automatic Generation of Fuzzy Decision Models Based on Expert Knowledge to Support the Decision-Making Process”, que se encuentra en fase de Revisión Menor en *Applied Soft Computing*. Se puede ver en https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5009719.

Index Terms—KD-DSS, Fuzzy rules, Decision Models, TSK FIS

Este trabajo ha sido financiado por los proyectos PID2020-112967GB-C32 y PID2020-112967GB-C33, subvencionados por MCIN/AEI/10.13039/501100011033, por el FEDER “Una manera de hacer Europa” y por el Vicerrectorado de Investigación de la Universidad de Castilla-La Mancha. Fue realizado mientras Enrique Adrián Villarrubia-Martin era investigador predoctoral en la Universidad de Castilla-La Mancha, financiado por el Fondo Social Europeo Plus (FSE+).

An overview of Fuzzychain: a fuzzy logic-based consensus algorithm for Blockchain

1st Bruno Ramos-Cruz

Computer Science Department
University of Jaen
Jaen, Spain
brcruz@ujaen.es

2nd Javier Andreu-Perez

Centre for Computational Intelligence
School of Computer Science and Electronic Engineering
University of Essex
Colchester, United Kingdom
j.andreu-perez@essex.ac.uk

3rd Francisco J. Quesada-Real

Computer Science Department
University of Jaen
Jaen, Spain
fqreal@ujaen.es

4th Luis Martínez

Computer Science Department
University of Jaen
Jaen, Spain
martin@ujaen.es

I. INTRODUCTION

Establishing trustworthiness in digital transactions is still a significant challenge. Blockchain, a decentralized and secure peer-to-peer network, addresses this issue through immutable and transparent transaction records, which are validated by network nodes using consensus algorithms such as Proof of Work (PoW) [1], [2], Proof of Stake (PoS) [3], and Delegated Proof of Stake (DPoS) [4], among others. PoW favors high-computation nodes, leading to potential centralization, while PoS—based on staked assets—struggles with subjective stake perceptions and power imbalances. Both mechanisms lack equitable validator diversification, which is crucial for achieving true decentralization and ensuring network security. In [5] we have proposed Fuzzychain, a novel consensus mechanism that leverages fuzzy set theory to improve stake-based validation by using nodes stake and reputation, fostering greater inclusivity, equity, and fairness in the selection of validators. The **main contributions of Fuzzychain** are the following:

- ⊕ An innovative consensus algorithm for enabling stake selection through the soft association of a validator with multiple fuzzy groups, rather than relying solely on the conventional mean crisp value of the stake, as typically used in PoS.
- ⊕ A soft categorisation and reputation-based validator selection approach that ensures validators are chosen based on both stake and reputation rather than solely on stake.
- ⊕ A function that models the behaviour of reputation is incorporated to enhance the consensus algorithm in the selection phase.
- ⊕ A novel penalization mechanism where validators who perform incorrect validations face a rapid decrease in their reputation, while reputation recovery is harder and gradual.

II. RELATED WORK

Reputation-based mechanisms have emerged as promising solutions to mitigate centralization issues in consensus

algorithms. For example, [6] proposes a vote-based reputation leader selection mechanism tailored for consortium blockchains, while [7] introduces a reputation-weighted asynchronous BFT protocol that incorporates filtering techniques. [8] presents a hierarchical delegated Proof of Reputation (hDPoR) model using multi-weight subjective logic, and [9] enhances PoS by employing deep reinforcement learning for dynamic reputation modeling. Variants of Proof of Reputation (PoR), such as the tripartite evolutionary game model in [10] and the Fair PoR framework based on variational neural networks in [11], aim to achieve fairer validator selection. Despite these advancements, most of these approaches rely on rigid, “crisp” reputation metrics, which limit their ability to capture the nature of stakeholder trust and perceptions.

III. FUZZYCHAIN: EQUITABLE CONSENSUS ALGORITHM

Our proposal introduces an equitable consensus algorithm based on the proof of stake, reputation and fuzzy set theory to validate and verify block transactions, thus giving rise to the Fuzzychain. Moreover, our method can be adapted and extended to other consensus mechanisms that rely on stake-based validator selection, such as DPoS, hybrid PoS/PoW system, and other approaches that incorporate stake and reputation in the validator selection process. The proposed equitable consensus algorithm is composed by three phases: scaling, selection and voting phase.

Scaling Phase: In this phase, the procedure involves scaling a participant’s stake, denoted by the numerical value x , into a set of fuzzy sets $T = \{T_1, T_2, T_3, \dots, T_{n-2}, T_{n-1}, T_n\}$ through the application of a membership function $\mu_T(x)$, as illustrated in Figure 1. The universe of discourse $X = [l, r]$ is partitioned into n uniformly spaced and distributed type-1 fuzzy membership functions (T1-MFs), where each T1-MF corresponds to a fuzzy set $T_i \in T$. Subsequently, the numerical value x of a participant’s stake is evaluated through the membership function $\mu_{T_i}(x)$, which scale it to the appropriate

fuzzy set T_i within the defined universe of discourse $[l, r]$, corresponding to the linguistic variable Participant's stake. The degree of membership is determined by identifying the highest membership degree function (HMDF). At this point, all stake values x_i in the participant set have been scaled into fuzzy sets T_i , each assigned a linguistic label (Very Low, Low, Moderate, among others), using the appropriate fuzzy membership functions.

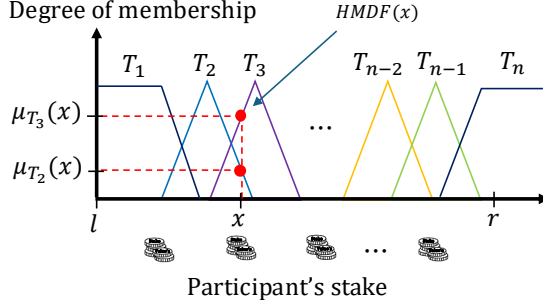


Fig. 1. The items x in the set of Participant's stakes are scaled through the fuzzy membership function $\mu_{T_i}(x)$ into fuzzy sets T_i .

Selection phase: This phase aims to choose participants t_i from each fuzzy set T_i . To achieve it, the phase involves two selection algorithms (i) validator random selection and (ii) validator selection according to the reputation. The former is an algorithm that randomly chooses the participants, and the latter is then used to select participants based on their reputation from each fuzzy set.

Each participant t_i entering a specific fuzzy set starts with an initial reputation set to 1 (maximum reputation). Subsequently, their reputation may be maintained or decreased depending upon their performance, i.e., success validation (SV) or unsuccessful validation ($UnSV$) in their verification tasks. To model the behaviour of the reputation when the validator's reputation differs from 1, the following function is defined.

Definition 1: Let η be the decrease rate and let $\frac{\eta}{l}$, $l \in \mathbb{N}$ be the increase rate. If $rep(t_i, j)$ is the reputation for the validator t_i in the round j , then $rep(t_i, j + 1)$ for the round $j + 1$ is defined by

$$rep(t_i, j + 1) = \begin{cases} 1 & \text{if } t_i \text{ is a } SV \\ & \text{and } rep(t_i, j) = 1 \\ rep(t_i, j) + \frac{\eta}{l} & \text{if } t_i \text{ is a } SV \\ rep(t_i, j) - \eta & \text{if } t_i \text{ is a } UnSV \end{cases}$$

where $0 \leq rep(t_i, j + 1) \leq 1$.

Voting phase: Once the validators have been selected from each fuzzy set T_i at the selection phase, every one of them individually engages in the validation process and subsequently indicates whether the block should be accepted or rejected. This phase involves a voting mechanism employed to determine the block's validity, wherein a consensus is reached based on the majority of participants' decisions. If the majority indicates acceptance, the block is accepted; otherwise, it is

TABLE I
THIS TABLE SHOWS THE COMPARISON BETWEEN PoW, PoS, DPOS, PBFT, PoR AND FUZZYCHAIN.

Metrics	Consensus algorithms					
	PoW	PoS	DPOS	PBFT	PoR	Fuzzychain
Gini coefficient	0.5992	0.4934	0.4126	0.5147	0.3728	0.1720
Skewness	1.8855	1.5243	0.9630	1.5569	0.5201	0.2243
Kurtosis	3.3206	2.8247	1.3253	3.0123	0.6985	-1.7489

rejected. The sample spaces of the voting mechanism is $\Omega = \{accepted, rejected\}$; there are no other possibilities. The validators who won the vote will be considered successful, that is if the majority indicates that the block is accepted or rejected.

IV. RESULTS

This proposal includes two experiments. Experiment 1 focuses exclusively on evaluating the diversity of selection of Fuzzychain showed in Figure 2. Experiment 2 is designed to compare the proposed algorithm with other established consensus mechanisms. Table I presents a comparative analysis of PoW, PoS, DPOS, PBFT, PoR, and Fuzzychain. The Gini coefficient is used to measure inequality within a distribution, such as income levels. A Gini coefficient of 0 indicates perfect equality, where all participants hold equal shares, while a coefficient of 1 reflects maximum inequality, with all value concentrated in a single participant. Skewness values near zero indicate a more symmetrical and balanced distribution. Furthermore, lower kurtosis suggests a reduced likelihood of extreme values, indicating a more stable and evenly distributed system.

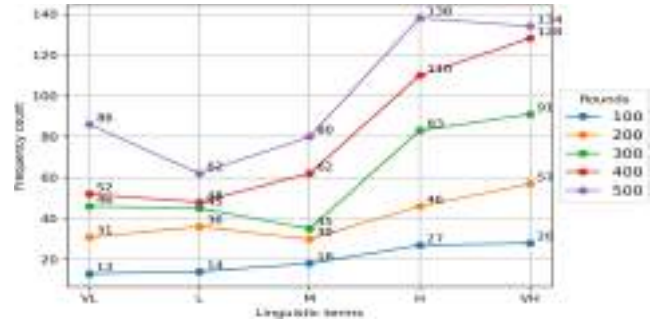


Fig. 2. The figure despite five plots corresponds to the selection of validators for each linguistic label in 100, 200, 300, 400 and 500 rounds..

V. CONCLUSION

This study introduces Fuzzychain, a novel consensus algorithm that leverages fuzzy logic to model the stakes and reputation of the validators. By doing so, it promotes more equitable stake distribution, enhances transaction integrity, and improves security by preventing the concentration of control. Integrated with a reputation-based reward system, Fuzzychain prioritizes reputable validators, thereby fostering greater inclusivity, equity, and fairness in validator selection.

REFERENCES

- [1] C. Dwork and M. Naor, “Pricing via processing or combatting junk mail,” in *Advances in Cryptology - CRYPTO '92, 12th Annual International Cryptology Conference, Santa Barbara, California, USA, August 16-20, 1992, Proceedings* (E. F. Brickell, ed.), vol. 740 of *Lecture Notes in Computer Science*, pp. 139–147, Springer, 1992.
- [2] S. Nakamoto, “Bitcoin: A peer-to-peer electronic cash system,” *Decentralized business review*, 2008.
- [3] Ethereum, “Proof-of-stake (pos).” [Online], 2023. Last accessed 2024-01-15, <https://ethereum.org/developers/docs/consensus-mechanisms/pos>.
- [4] D. Larimer, “Delegated proof-of-stake (dpos),” *Bitshare whitepaper*, vol. 81, p. 85, 2014.
- [5] B. Ramos-Cruz, J. Andreu-Pérez, F. J. Quesada, and L. Martínez, “Fuzzychain: An equitable consensus mechanism for blockchain networks,” *Journal of Network and Computer Applications*, p. 104204, 2025, in press.
- [6] M. Hussain, A. Mehmood, M. A. Khan, R. Khan, and J. Lloret, “Reputation-based leader selection consensus algorithm with rewards for blockchain technology,” *COMPUTERS*, vol. 14, JAN 2025.
- [7] G. Jia, Z. Shen, H. Sun, J. Xin, and D. Wang, “Rwa-bft: Reputation-weighted asynchronous bft for large-scale iot,” *SENSORS*, vol. 25, JAN 2025.
- [8] Z. Hu, B. Liu, A. Shen, and J. Luo, “Blockchain-based resource allocation mechanism for the internet of vehicles: Balancing efficiency and security,” *IEEE TRANSACTIONS ON NETWORK AND SERVICE MANAGEMENT*, vol. 21, pp. 3971–3987, AUG 2024.
- [9] X. Hao, P. L. Yeoh, C. She, B. Vucetic, and Y. Li, “Secure deep reinforcement learning for dynamic resource allocation in wireless mec networks,” *IEEE TRANSACTIONS ON COMMUNICATIONS*, vol. 72, pp. 1414–1427, MAR 2024.
- [10] C. Gan, X. Xiao, Y. Zhang, Q. Zhu, J. Bi, D. K. Jain, and A. Saini, “An asynchronous federated learning-assisted data sharing method for medical blockchain,” *APPLIED INTELLIGENCE*, vol. 55, JAN 2025.
- [11] C. P. Singh, R. Yamaganti, and L. S. Umrao, “Variational onsager neural networks-based fair proof-of-reputation consensus for blockchain with transaction prioritisation for smart cities,” *JOURNAL OF SUPERCOMPUTING*, vol. 81, JAN 2025.

Enfoque de inferencia difusa en la toma de decisiones para medir indicadores emocionales

1st Miguel Andrés

dept. Métodos Cuantitativos
CUNEF Universidad
miguel.andres@cunef.edu
Facultad de Informática
Universidad Complutense
de Madrid
Madrid, Spain
maherrero@ucm.es

2nd Matilde Santos

Inst. de Tecno. del Conocimiento
Universidad Complutense
de Madrid
Madrid, Spain
msantos@ucm.es

3rd Victoria López

dept. Métodos Cuantitativos
CUNEF Universidad
Madrid, Spain
victoria.lopez@cunef.edu

4th Diego Riofrío-Luzcando

dept. Métodos Cuantitativos
CUNEF Universidad
Madrid, Spain
diego.riofrío@cunef.edu

Resumen—Se presenta un sistema de inferencia basado en lógica difusa para la toma de decisiones en un entorno de videojuego (de recolección de frutas), diseñado para evaluar la motivación en adolescentes sanos y con trastornos emocionales. El sistema presenta un escenario en el que el jugador debe tomar una serie de decisiones a partir de las que se calcula un indicador emocional (estrés, apatía, ansiedad). Para ello se utiliza un conjunto de variables iniciales cuya actualización determina el perfil de juego tras cada acción. Estas variables definidas mediante lógica borrosa (energía, tiempo, premios conseguidos y distancia) permiten definir un conjunto de reglas de inferencia que modelan las acciones en el juego relacionando estas acciones con una evaluación de resultados. El sistema se ha implementado en lenguaje Python utilizando la biblioteca *scikit-fuzzy*, que permite la derivación de decisiones ponderadas en una escala continua de 0 a 10. El sistema borroso se integra en el desarrollo del videojuego para el estudio de problemas de comportamiento en adolescentes con trastornos emocionales.

Index Terms—Lógica difusa, toma de decisiones, motivación, inferencia borrosa, trastornos emocionales, videojuegos.

I. INTRODUCTION

Los sistemas de lógica difusa han demostrado ser eficaces para representar un razonamiento humano aproximado, especialmente en contextos donde la información es incierta, imprecisa o subjetiva [2]–[4]. Este estudio presenta un problema de toma de decisiones en un entorno de videojuegos desarrollado específicamente para medir la motivación en adolescentes con trastornos emocionales [1]. Los indicadores utilizados para medir la motivación se basan en las acciones del usuario durante el juego, ponderando la decisión tomada en cada momento por el jugador. El videojuego se ha desarrollado de forma específica para el estudio y se describe en [5]. Básicamente se trata de recolectar el máximo número de frutas posible en un tiempo y con una energía limitados. Las frutas (objetivo) están distribuidas de forma que el jugador tiene que tomar una decisión ponderable y penalizable. En el artículo [5] también se propone una solución difusa basada en la lógica que evalúa múltiples condiciones y produce una decisión ponderada sobre el curso de la acción. El objetivo es medir las emociones relacionadas con la motivación, como

la ansiedad y el estrés, y sus consecuencias en la toma de decisiones a lo largo de un juego.

II. FUNDAMENTOS DEL SISTEMA

El uso de dispositivos inteligentes para monitorizar la recogida de datos de comportamiento humano ha sido probado con éxito (ver [6]) en el proyecto ACERTA [7] del que forma parte este trabajo. Las variables utilizadas se describen en la Tabla I, incluyendo tanto las variables de entrada como de salida, sus rangos de valores, y las etiquetas lingüísticas que representan sus niveles. Asumiremos funciones trapezoidales al ser una buena primera elección para las pruebas de concepto.

Tabla I
RELACIÓN DE VARIABLES DEL SISTEMA

Tipo	Variable	Rango	Etiquetas
Entrada	Energía	0–100 %	Baja, Media, Alta
Entrada	Tiempo restante	0–10	Bajo, Medio, Alto
Entrada	Valor de la fruta	1–10	Bajo, Medio, Alto
Entrada	Distancia	0–10	Corta, Media, Larga
Salida	Decisión	0–10	No, Quizas, Yes

La salida numérica se interpreta de la siguiente manera:

- 0–3: No recolectar
- 4–6: Recolección dudosa
- 7–10: Recolectar

Los rangos de pertenencia trapezoidal para cada variable se definen en Tabla II.

III. OBJETIVO E IMPLEMENTACIÓN

Durante el tiempo de juego, el jugador debe tomar la decisión de desplazarse de un punto a otro en función de la energía, el tiempo restante, de la distancia al objetivo y su valor (valor de la fruta).

III-A. Motor de lógica difusa de Python

El sistema difuso se implementa usando la librería *scikit-fuzzy* que permite la creación de variables difusas, funciones de pertenencia y conjuntos de reglas. Se ejecuta un

Tabla II
RANGOS DE LAS FUNCIONES DE PERTENENCIA

Variable	Etiqueta	Rango (a, b, c)
Energía (E)	Baja	(0, 0, 40)
	Media	(20, 50, 80)
	Alta	(60, 100, 100)
Tiempo restante (T)	Bajo	(0, 0, 4)
	Medio	(2.5, 8)
	Alto	(6, 10, 10)
Valor de la fruta (V)	Bajo	(0, 0, 4)
	Medio	(2.5, 8)
	Alto	(6, 10, 10)
Distance (D)	Corta	(0, 0, 3)
	Media	(2, 5, 8)
	Larga	(6, 10, 10)
	Quizás	(3, 5, 7)
Decisión (R)	No	(0, 0, 4)
	Quizás	(3, 5, 7)
	Yes	(6, 10, 10)

caso de prueba estableciendo valores de entrada específicos y computando la decisión de salida. Se definen las siguientes reglas para modelizar un comportamiento esperado de un jugador emocionalmente no sintomático:

- si $E \wedge T$ son bajos $\wedge D$ es larga, $\Rightarrow R = No$
- si E es media $\wedge V$ es alta $\wedge D$ es corta, $\Rightarrow R = Yes$
- Si $E \wedge T$ son altos $\wedge V$ es medio, $\Rightarrow R = Yes$
- Si E es media $\wedge T$ es bajo $\wedge V$ es bajo, $\Rightarrow R = No$
- Si E es baja $\wedge D$ es media $\wedge V$ es alta, $\Rightarrow R = Maybe$
- Si E es alta $\wedge D$ es larga $\wedge V$ es alta, $\Rightarrow R = Yes$
- Si $E \wedge T$ son medios $\wedge V$ es medio, $\Rightarrow R = Maybe$

Esta especificación, permite la implementación del modelo y la simulación de casos. La Tabla III muestra un ejemplo de simulación.

Tabla III
VALORES DE ENTRADA PARA LA SIMULACIÓN

Variable	Valor	Etiqueta
Energía	55	Media
Tiempo restante	6	Medio-Alto
Valor de la fruta	7	Medio-Alto
Distancia	2	Baja-Media

IV. RESULTADOS

Al ejecutar el código se obtienen los resultados de la Figura 1, donde se discretiza la decisión de recoger la fruta (yes o no) entre los valores 0 a 10. Sobre la base del valor resultante, ligeramente por encima de 6, el sistema sitúa la decisión dentro del rango *Maybe*, lo que sugiere que en el juego simulado, la recolección de la fruta parece beneficiosa pero con cierta incertidumbre. Por lo tanto, el jugador puede decidir recogerlo si no hay riesgos o si no hay mejores opciones..

Si revisamos los resultados por las variables de energía (E), tiempo (T) y valor de la fruta (V) se puede ver cómo cada variable de entrada influye en la decisión final mediante su grado de pertenencia en las respectivas funciones difusas. La energía (55) se sitúa dentro del rango “media” aportando una base aceptable pero no óptima para la acción. El valor de la fruta (7) se ubica entre las categorías “medio” y “alto”, lo que

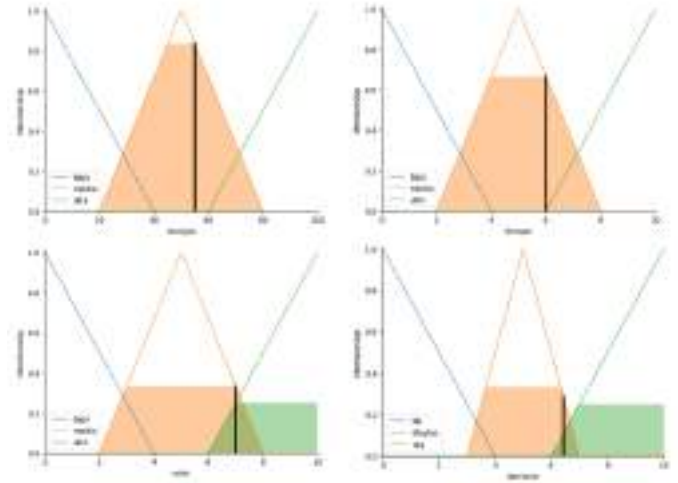


Figura 1. Resultados de la simulación.

incrementa su atractivo y refuerza parcialmente la decisión de recolectar. Por su parte, el tiempo restante (6) se encuentra en una zona limítrofe entre “medio” y “alto”, lo cual contribuye a una sensación de oportunidad, aunque no sin reservas.

V. CONCLUSIONES

Se presenta un sistema de toma de decisiones de lógica difusa cuyo propósito es evaluar la motivación a través del uso un videojuego y variables de entrada definidas con términos lingüísticos. Un conjunto de reglas de inferencia modela un comportamiento esperado, resultando en una decisión de salida ponderada en una escala de 0 a 10. Se muestra un caso de prueba simulado, con valores específicos, que retorna un valor ligeramente superior a 6, situando la decisión en el rango “Maybe”. Esto demuestra la capacidad del sistema para asistir en decisiones bajo incertidumbre. Se concluye que el sistema propuesto es una herramienta eficaz, adaptable y transparente para analizar comportamientos del jugador en escenarios inciertos, siendo integrable en el estudio de problemas de comportamiento en adolescentes con trastornos emocionales.

REFERENCIAS

- [1] P. Llamocca, P., V. López, and M. Čukić, “The proposition for bipolar depression forecasting based on wearable data collection”, in *Frontiers in Physiology*, vol.12, 2022.
- [2] V. López, M. Santos, and J. Montero, “Fuzzy specification in real estate market decision making,” *Int. J. of Computational Intelligence Systems*, vol. 3, no. 1, pp. 8–20, 2010.
- [3] V. F. Reyna, S. Edelson, B. Hayes, and D. Garavito, “Supporting health and medical decision making: findings and insights from fuzzy-trace theory,” *Medical Decision Making*, vol. 42, no. 6, pp. 741–754, 2022.
- [4] D. Rad et al., “A fuzzy logic modelling approach on psychological data,” *J. of Intelligent & Fuzzy Systems*, vol. 43, no. 2, pp. 1727–1737, 2022.
- [5] M. Andrés-Herrero, V. Lopez, M. Santos, D. Urgelés, and M. Favieres, “Design of a video game adapted to the study of motivation in young people with emotional disorders” in *Human Interaction and Emerging Technologies (IHET-AI 2025): Artificial Intelligence and Future Applications AHFE*, vol. 161, pp. 419–427, 2025.
- [6] P. Llamocca, C. Guevara, Y. Jiménez and V. López., “Smartwatches should facilitate the aggregation of mental health data” in *SIAM Society for industrial and applied mathematics SIAM News Blog*, 2024
- [7] Available 2025: <http://www.computational-intelligence.com>

Three-Way Group Decision-Making With Personalized Numerical Scale of Comparative Linguistic Expression: An Application to Traditional Chinese Medicine

1st Yaya Liu

School of Business

University of Shanghai for Science and Technology

Shanghai, China

yayaliu@usst.edu.cn

2nd Lina Zhu

School of Business

University of Shanghai for Science and Technology

Shanghai, China

222251075@st.usst.edu.cn

3rd Rosa M. Rodríguez

Department of Computer Science

University of Jaén

Jaén, Spain

rmrodrig@ujaen.es

4th Yiyu Yao

Department of Computer Science

University of Regina

Regina, Canada

yiyu.yao@uregina.ca

5th Luis Martínez

Department of Computer Science

University of Jaén

Jaén, Spain

martin@ujaen.es

Abstract—This is a summary of our article published in *IEEE Transaction on Fuzzy Systems* [1]. This article proposes a novel multiattribute three-way group decision making method dealing with comparative linguistic expressions (CLEs) and a personalized numerical scale computation model to handle different interpretations of CLEs.

Index Terms—Comparative linguistic expression, computing with words, group decision making, personalized numerical scale, three-way decision

I. INTRODUCTION

Real-world decision-making problems are defined in changing contexts in which uncertainty and vagueness are quite common. In these situations, the use of linguistic information has provided successful results such as, in medical diagnosis, risk evaluation, or strategic planning. To address the cognitive alignment between human judgment and computational models, recent research has turned to linguistic group decision-making techniques and the broader framework of Computing with Words (CWW). One such approach is Three-Way Decision (3WD) theory [2], which extends classical binary decision frameworks by introducing a third option: *noncommitment*. This structure provides a more nuanced mechanism for reasoning under uncertainty, acknowledging that sometimes the most rational action is to delay decision-making until further information becomes available.

Despite its potential, traditional 3WD models have several limitations. Most notably, they deal with single linguistic terms to model experts' preferences, and fail to capture the inherent

subjectivity of linguistic assessments -what one expert considers "moderate" may be perceived by another as "severe". To overcome these limitations, this article proposes a novel Three-Way Group Decision-Making (3WGDM) model that deals with Comparative Linguistic Expressions (CLEs) [3] and Personalized Numerical Scales (PNS) [4].

CLEs enhance expressiveness by enabling experts to provide evaluations such as "greater than moderate" or "between mild and severe", which are more natural in human reasoning than single linguistic terms. However, these expressions can be interpreted differently by different experts. To cope with these situations the PNS introduces a method for quantifying CLEs in a way that reflects individual cognitive differences. This allows linguistic terms to be mapped to numerical values based on personal interpretation, thus preserving subjectivity while enabling computation.

Additionally, the framework addresses experts' influence by embedding social trust dynamics into the decision process. In traditional GDM, experts' opinions are often weighted, since the 3WGDM model proposes a 3WD-based social network where experts' weights are determined by trust relationships using in-degree and out-degree centrality indices. Experts that are widely trusted (high in-degree) receive greater influence, whereas those who distribute trust liberally (high out-degree) are weighted less. Moreover, the model introduces a third trust category, *noncommitment*, aligning again with 3WD's core philosophy.

The CLE-3WGDM model consists of the following main phases:

- 1) **Preclassification based on predecision:** Experts provide preliminary judgments based on their own observa-

This work is supported partially by the research project Andalucía Excellence Research Program under Grant ProyExcel-00257.

tion on alternatives, that is called predecisions, and they have influence over the final decisions. Based on such predecisions, alternatives are classified into categories. The way to obtain the categories from predecisions can be various, and the number of categories can be also different.

- 2) **CLEs to hesitant fuzzy linguistic term sets (HFLTSSs) transformation:** Experts elicitate their assessments on alternatives with regards to attributes by using CLEs. These linguistic expressions are transformed into HFLTSSs by a transformation function [3], to facilitate the CWW processes.
- 3) **PNS computation:** An optimization model is used to derive personalized numerical values for the linguistic terms, based on each expert's predecisions and classification tendencies. Thus, the PNS of CLEs can be obtained from the PNS of the linguistic terms.
- 4) **Attribute weight computation:** A new optimization model is introduced to compute the attributes weights. This model uses the entropy of the HFLTSS to measure the uncertainty contained in the linguistic terms, and the distance measure of the HFLTSS to obtain the difference of CLEs.
- 5) **Experts weight computation:** Expert weights are determined from a 3WD-informed social trust network. It can be graphically presented, in which experts are shown by points, and trust relationships are reflected by edges. Therefore, the number of edges is used to distinguish trust attitudes among experts, where arrows represent the trust direction among experts. The expert who has a higher in-degree is more important in the social trust network and the expert who has a higher out-degree is less important. There are 3 kinds of trust relationships between experts:
 - *Case A:* absolute trust.
 - *Case B:* not sure trust or not, delay to determine the trust situation.
 - *Case C:* absolute distrust.

This classification follows the idea of 3WD, where case *A* is related to the strategy of "accept", case *B* is related to "noncommitment" and case *C* is related to "reject".

- 6) **Loss function transformation:** The CLEs are translated into relative loss functions, accounting for PNS. To do so, the method presented by Jia and Liu [5] is extended to generate the relative loss functions from the assessments regarding the attributes.
- 7) **Final decision-making:** Using an extended TODIM method [6] to deal with CLEs based on PNS, the conditional probabilities are obtained. Finally, decisions are made according to 3WD's tripartite classification.

A case study based on traditional chinese medicine is used to show the effectiveness and robustness of the proposed 3WGDM method. It consists of nine patients that show symptoms potentially indicative of "deficiency-cold of spleen and stomach". Five experts assessed the patients based on

4 attributes: fatigue, appetite, cold intolerance, and loose stools. CLEs are used to express the symptom severity, and experts' trust relationships form a social network for influence weighting. The model successfully classified patients into three categories: (i) those recommended for treatment, (ii) those not recommended, and (iii) those requiring further observation. Importantly, the results were consistent with known diagnostic standards and expert consensus, demonstrating the effectiveness of the model.

A comparative analysis has been also introduced to compare the proposed model with other existing ones taking into account the following features:

- How determinate attributes weights.
- How determinate experts weights.
- How to compute conditional probability.

Moreover, the proposed 3WGDM model is the only one that deal with linguistic expressions more flexible than single linguistic terms and considers personalized semantics in comparison with the others.

II. CONCLUSIONS

This approach introduces a novel 3WGDM framework that integrates CLEs and PNS, expanding the capacity of existing 3WD models to handle uncertainty and linguistic semantics. By embedding personal semantic interpretation and social trust dynamics into the decision process, the framework significantly improves both the flexibility and realism of group evaluations. Its application in traditional chinese medicine highlights its capacity to replicate complex expert reasoning and adapt to high-uncertainty domains. The inclusion of a social trust network further refines expert influence, ensuring that decision outputs are not only technically sound but also socially grounded.

REFERENCES

- [1] Y. Liu, L. Zhu, R. M. Rodríguez, Y. Yao, and L. Martínez, "Three-way group decision-making with personalized numerical scale of comparative linguistic expression: An application to traditional chinese medicine," *IEEE Transactions on Fuzzy Systems*, vol. 32, pp. 4352–4363, 2024.
- [2] Y. Yao, "Three-way decisions with probabilistic rough sets," *Information Sciences*, vol. 180, no. 3, pp. 341–353, 2010.
- [3] R. M. Rodríguez, L. Martínez, and F. Herrera, "Hesitant fuzzy linguistic term sets for decision making," *IEEE Transactions on Fuzzy Systems*, vol. 20, no. 1, pp. 109–119, 2012.
- [4] C.-C. Li, Y. Dong, F. Herrera, E. Herrera-Viedma, and L. Martínez, "Personalized individual semantics in computing with words for supporting linguistic group decision making. an application on consensus reaching," *Information Fusion*, vol. 33, pp. 29–40, 2017.
- [5] F. Jia and P. Liu, "A novel three-way decision model under multiple-criteria environment," *Information Sciences*, vol. 471, pp. 29–51, 2019.
- [6] L. Gomes and M. Lima, "TODIM: Basics and application to multicriteria ranking of projects with environmental impacts," *Foundations of Computing and Decision Sciences*, vol. 16, pp. 113–127, 1991.

Marco de toma de decisiones multicriterio con expresiones lingüísticas flexibles y conjuntos nube aproximados multigranulares: aplicación en evaluación de proyectos de Internet de la Energía

1st Han Wang
Dept. Ciencias de la Computacion
e I.A., DaSCI
Universidad de Granada
Granada, Espana
wangh936@163.com

2nd Yanbing Ju
School of Management
Beijing Institute of Technology
Beijing, China
juyb@bit.edu.cn

3th Carlos Porcel
Dept. Ciencias de la Computacion
e I.A., DaSCI
Universidad de Granada
Granada, Espana
cporcel@decsai.ugr.es

Abstract—Con el avance de la economía circular y los objetivos de neutralidad de carbono, proyectos de Internet de la Energía (IE) desempeñan un papel fundamental en la promoción de la transformación ecológica de los sistemas industriales. Teniendo en cuenta la incertidumbre, imprecisión y subjetividad involucradas en la evaluación de estos proyectos, este trabajo propone un nuevo marco de Toma de Decisiones Multicriterio (MCDM) que integra Expresiones Lingüísticas Flexibles (FLEs), modelos de nube y Conjuntos Nube Aproximados Multigranulares (MGCRS). Este enfoque convierte la información lingüística discreta en datos continuos mediante modelos de nube y permite una clasificación estable a través de aproximaciones multigranulares. Para validar la precisión, robustez y aplicabilidad del método, se usa un caso de estudio con cuatro proyectos. Los resultados muestran que la propuesta proporciona ayuda en la evaluación y selección de proyectos energéticos, siendo especialmente útil para escenarios complejos como la economía circular, la gestión energética y la toma de decisiones sostenible.

Index Terms—Toma de decisiones multicriterio, Economía circular, Internet de la Energía.

I. INTRODUCCIÓN

En el contexto de neutralidad de carbono y desarrollo sostenible, los sistemas industriales están en proceso de transición hacia una economía circular. La IE desempeña un papel vital en esta transformación. Al integrar redes inteligentes, energía distribuida y análisis de datos en tiempo real, la IE ofrece soluciones energéticas verdes que optimizan los flujos de energía y reducen el desperdicio. El concepto de IE está alineado con los objetivos estratégicos de innovación ecológica, eficiencia energética y transformación digital. Iniciativas como la Interconexión Energética Global (GEI) y los marcos estratégicos introducidos por organismos internacionales como el Instituto de Ingenieros Eléctricos y Electrónicos (IEEE) y el Instituto Nacional de Estándares y Tecnología (NIST), han establecido una base sólida para la estandarización y la implementación práctica. Países como China, Japón y los estados miembros de la Unión Europea están impulsando

activamente proyectos de IE como componentes integrales de sus estrategias de desarrollo bajas en carbono.

Sin embargo, aunque los sistemas técnicos y de estandarización de la IE se han desarrollado adecuadamente, la evaluación sistemática de esos proyectos sigue siendo un desafío. Los enfoques tradicionales a menudo resultan insuficientes, debido a la complejidad de estos sistemas, que involucran varias partes interesadas, distintas tecnologías y objetivos heterogéneos. La subjetividad de los juicios de los expertos, combinada con la ambigüedad en las evaluaciones lingüísticas, complica aún más la toma de decisiones. Para facilitar estas evaluaciones se han usado con éxito las FLEs, especialmente en escenarios complejos como la evaluación de proyectos de IE, donde a menudo no se dispone de entradas numéricas precisas o no son apropiadas. Sin embargo, la investigación existente sobre FLEs se centra principalmente en la normalización simbólica o el mapeo cualitativo de términos, lo que tiende a ignorar la continuidad de la semántica lingüística y limita el potencial para un análisis cuantitativo más profundo. Por otro lado, los métodos clásicos de Conjuntos Aproximados Multigranulares (MGRS) suelen estar diseñados para entornos de datos discretos y se enfrentan a desafíos significativos al tratar con la naturaleza difusa, continua o probabilística de la información que típicamente se encuentra en las evaluaciones de proyectos de IE.

Para superar estas limitaciones, en [1] se propone un marco MCDM integrado que combina FLEs y MGCRS. Este enfoque híbrido permite la transformación de términos lingüísticos discretos en representaciones continuas mediante modelos de nube y facilita una toma de decisiones estable y flexible con varias granularidades. Así ayuda a completar un vacío metodológico importante en la evaluación de proyectos energéticos y la planificación de políticas en el contexto de la economía circular.

II. DESCRIPCIÓN DE LA PROPUESTA

El marco propuesto se compone de tres módulos clave que se describen a continuación:

1. Se construye un sistema de indicadores para evaluar proyectos de IE en sus dimensiones de madurez tecnológica, integración de energía verde y beneficios socioeconómicos. Con el fin de garantizar su relevancia y exhaustividad, estas dimensiones se desglosan en diez criterios específicos, seleccionados a partir de una revisión de la literatura y consultas con expertos. El marco proporciona la base para recopilar juicios de expertos y asegurar la coherencia en las comparaciones entre proyectos.

2. Se emplea una escala lingüística flexible de 7 niveles para reflejar las opiniones de los expertos de una manera más natural e interpretable. Estos términos lingüísticos son difusos e inciertos, por lo que no son adecuados para un procesamiento numérico directo. Para abordarlo, el marco introduce un mecanismo de transformación basado en modelos de nube, donde cada término lingüístico se mapea a un modelo de nube caracterizado por los parámetros: valor esperado, entropía e hiperentropía. Esto permite una representación continua y probabilística de las opiniones de los expertos, preservando eficazmente la incertidumbre e imprecisión inherentes a evaluaciones lingüísticas.

3. Para realizar el razonamiento y clasificación se utiliza el modelo MGCRS. Este método opera con varios niveles de granularidad ofreciendo la flexibilidad de simular diferentes perspectivas cognitivas. Para cada nivel se establecen conjuntos de aproximación superior e inferior bajo actitudes tanto optimistas como pesimistas, capturando así una amplia gama de comportamientos de toma de decisiones. A continuación se calculan los valores esperados a partir de los modelos de nube bajo cada escenario de aproximación, y se obtiene una clasificación global.

Este método integrado y multinivel permite la fusión de evaluaciones de expertos subjetivas y maneja eficazmente tanto datos lingüísticos discretos como continuos. Al combinar el modelado lingüístico con el razonamiento granular, el marco es capaz de ofrecer resultados de evaluación robustos y adaptables, lo que lo hace especialmente adecuado para entornos de toma de decisiones complejos e inciertos, como los que caracterizan a los proyectos de IE.

III. CASE DE ESTUDIO Y RESULTADOS

Se llevó a cabo un caso de estudio con cuatro proyectos de IE en la región Beijing-Tianjin-Hebei de China: Proyecto de demostración de energía inteligente en Binhai (Tianjin), zona de demostración de energía renovable en Zhangjiakou, proyecto piloto de IE en la ciudad de la ciencia del futuro de Beijing y sistema energético integrado del parque industrial de Tangshan.

Los expertos evaluaron estos proyectos en dimensiones como la madurez tecnológica, la integración de energía verde y los beneficios globales. Las evaluaciones se procesaron mediante el marco propuesto para establecer un modelo de clasificación comprensivo. Los resultados muestran que el

proyecto de Beijing ocupa el primer lugar debido a su fortaleza en integración de múltiples fuentes de energía y gestión inteligente, seguido por los de Zhangjiakou y Tangshan. El proyecto de Tianjin ocupa el último lugar debido a su integración y flexibilidad relativamente limitadas.

Para evaluar la robustez del modelo, los resultados se compararon con los obtenidos mediante métodos clásicos de MCDM como TOPSIS y VIKOR. La alta consistencia entre los métodos confirma la fiabilidad y generalización de la propuesta. Además, el análisis de sensibilidad demostró que los cambios en la granularidad lingüística o en la configuración de los pesos tienen un impacto mínimo en los resultados, lo que valida aún más la robustez del modelo en entornos de evaluación complejos y difusos. Por todo ello se confirma la validez científica y el valor práctico de la propuesta para la formulación de políticas y la priorización de proyectos.

IV. CONCLUSIONES Y TRABAJOS FUTUROS

Este artículo presenta un novedoso marco de MCDM para la evaluación de proyectos de IE, integrando expresiones lingüísticas flexibles, modelado de nubes y teoría de conjuntos aproximados multigranulares. El método aborda eficazmente los desafíos relacionados con la incertidumbre, la información incompleta y las evaluaciones subjetivas, transformando entradas lingüísticas discretas en datos en la nube continuos y elaborando razonamientos aproximados a través de varios niveles de granularidad.

Los resultados de los casos de uso demuestran el rigor científico, la viabilidad y la robustez del método. No solo captura las preferencias de los expertos y las características de los proyectos, sino que también muestra un alto rendimiento en términos de consistencia y sensibilidad, lo que lo hace realmente aplicable. En comparación con los enfoques tradicionales, este método mejora la integración de información cualitativa y cuantitativa, y es especialmente adecuado para evaluar proyectos complejos en el contexto de economía circular, transición energética y desarrollos bajos en carbono.

En trabajos futuros se podría extender con evaluaciones dinámicas o con múltiples etapas, incorporando métricas relacionadas con la sostenibilidad y la gobernanza o modelos de precios del carbono, e integrar algoritmos inteligentes para mejorar la agregación de las opiniones de los expertos.

V. AGRADECIMIENTOS

Este trabajo es parte del proyecto PID2022-139297OB-I00 financiado por MICIU/AEI/10.13039/501100011033 y por el Fondo Europeo de Desarrollo Regional (FEDER)/Unión Europea, y C-ING-165-UGR23 financiado por la Consejería de Universidad, Investigación e Innovación y por el Programa FEDER de Andalucía 2021-2027.

REFERENCES

- [1] H. Wang, Y. Ju, and C. Porcel, 'Energy internet project evaluation in circular economy practices: A novel multi-criteria decision-making framework with flexible linguistic expressions based on multi-granularity cloud-rough set', *Computers and Industrial Engineering*, vol. 201, no. 110890, 2025.

Un Modelo Difuso de Consenso Para Problemas de Toma de Decisiones en Grupo Basado en Computación Granular y *Blockchain*

Juan Carlos González-Quesada
dept. Ciencias de la Computación e IA
Universidad de Granada
Granada, España
juancarlosqg@ugr.es

Ignacio Javier Pérez
dept. Ciencias de la Computación e IA
Universidad de Granada
Granada, España
ijperez@decsai.ugr.es

Francisco Mata
dept. de Informática
Universidad de Jaén
Jaén, España
fmata@ujaen.es

Ramón Alberto Carrasco
dept. de Marketing
Universidad Complutense de Madrid
Madrid, España
ramoncar@ucm.es

Enrique Herrera-Viedma
dept. Ciencias de la Computación e IA
Universidad de Granada
Granada, España
viedma@decsai.ugr.es

Francisco Javier Cabrerizo
dept. Ciencias de la Computación e IA
Universidad de Granada
Granada, España
cabrerizo@decsai.ugr.es

Abstract—Este trabajo es un resumen de un artículo en desarrollo en el que se propone un nuevo modelo de consenso basado en computación granular que incorpora el uso de la tecnología *blockchain* para fomentar la confianza entre los participantes que intervienen en problemas de toma de decisión en grupo definidos en ambientes difusos. El modelo introduce una estructura de comunicación segura y eficiente basada en *blockchain* que está diseñada para minimizar las interacciones, generalmente sesgadas, entre los participantes. El modelo incorpora además un mecanismo de creación de confianza, basado también en *blockchain*, que anima a los participantes a reconsiderar y, potencialmente, modificar sus opiniones. A diferencia de modelos de consenso anteriores, no se basa en la propagación de la confianza, mejorando así la eficiencia computacional en el establecimiento de esta entre los participantes. Esto permite llegar a un consenso de forma más rápida. Los resultados obtenidos corroboran la validez, eficacia y aplicabilidad práctica de este nuevo modelo de consenso.

Index Terms—*Blockchain*, confianza, consenso, computación granular, toma de decisión en grupo difusa.

I. PROPUESTA

Para mejorar los modelos de consenso existentes en la literatura que utilizan el concepto de granularidad de la información, este estudio presenta un modelo de consenso granular que, por un lado, considera la voluntad de los participantes en un problema de toma de decisiones en grupo de aceptar o rechazar las modificaciones propuestas y, por otro, incorpora un mecanismo de creación de confianza basado en la tecnología *blockchain* [1]. Asumiendo relaciones de preferencia difusas para modelar las opiniones de los participantes, este objetivo general da lugar a las siguientes contribuciones específicas:

- Implementación de una estructura de comunicación eficiente y segura a través de la utilización de la tecnología *blockchain*. Dentro de esta estructura de comunicación, la identidad de cada participante, así como sus respectivas opiniones, permanecen confidenciales para los demás mediante el empleo de contratos inteligentes. Esto ayuda a salvaguardar a los participantes contra las «pérdidas» y, en consecuencia, fomenta la colaboración para alcanzar un acuerdo. Esto permite eliminar las influencias internas en el grupo, como, por ejemplo, el pensamiento y la presión del grupo [2], que podrían afectar a las opiniones individuales. Este marco permite también modelar las interacciones entre los participantes sin tener en cuenta la similitud de sus opiniones.
- Desarrollo de un mecanismo de creación de confianza dentro de la estructura de comunicación anterior que no depende de un proceso de propagación de la confianza, sino que también se basa en la tecnología *blockchain*. Al eliminar la restricción de requerir opiniones similares para generar confianza, el mecanismo de creación de confianza propuesto mejora significativamente la confianza entre los participantes y los guía más eficazmente hacia una solución de consenso.

La Figura 1 muestra el esquema del modelo propuesto. Los participantes comparten sus opiniones iniciales mediante relaciones de preferencia difusas. A continuación, se calculan los niveles iniciales de confianza, una matriz social y la importancia asignada a cada participante. Posteriormente, de las relaciones de preferencia difusas dadas por los participantes se derivan las medidas de consenso y de consistencia. En los casos en los que no se satisface un cierto criterio de parada, el moderador transmite a cada participante las modificaciones

Esta publicación es parte del proyecto de I+D+i PID2022-139297OB-I00, financiado por MICIU/AEI/10.13039/501100011033 y por FEDER/UE. También se ha realizado gracias a la ayuda PREP2022-000352 financiada por MICIU/AEI/10.13039/501100011033.

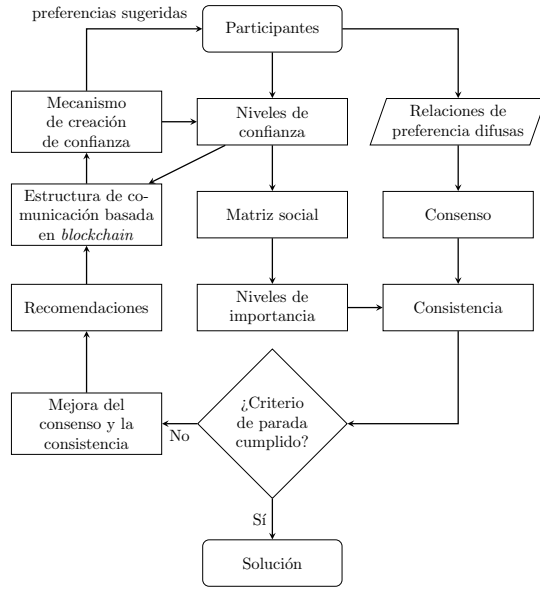


Fig. 1. Esquema del modelo de consenso propuesto.

necesarias en sus opiniones. Para ello, se asigna un nivel de granularidad de la información para generar los nuevos grados de preferencia que deben proporcionar los participantes para mejorar el consenso y la consistencia con la mínima pérdida de información posible. A continuación, se encomienda a los participantes la tarea de decidir si aceptan o rechazan estos cambios recomendados, influidos por factores como su confianza en el grupo, así como sus motivos externos e internos de traición. Es importante recordar que los individuos comunican sus decisiones de forma anónima a través de una estructura de comunicación segura basada en la tecnología *blockchain* y sus contratos inteligentes asociados. Se incluye también un mecanismo de creación de confianza, basado en la tecnología *blockchain*, que aumenta significativamente la confianza entre los participantes, guiándolos de manera más eficaz hacia una solución consensuada. Este proceso iterativo continúa hasta que el grupo alcanza un umbral mínimo de satisfacción predefinido o se supera un número máximo de rondas de negociación, los cuales componen el criterio de parada. Cuando se cumple, el proceso finaliza y se obtiene la solución al problema según las relaciones de preferencia difusas dadas por los participantes.

Para ilustrar la esencia del modelo propuesto, así como su flexibilidad y rendimiento, se han realizado diversos estudios experimentales y se ha efectuado un análisis comparativo. Los estudios experimentales realizados han corroborado la validez de la propuesta, al tiempo que han demostrado la eficacia de esta a la hora de fomentar el consenso y mejorar la consistencia de los participantes en sus opiniones, así como que el modelo propuesto es escalable y puede tratar eficazmente problemas de toma de decisiones en grupo de gran escala. Si se observa la Tabla I, se comprueba que el modelo de consenso basado en computación granular propuesto presenta una clara ventaja

TABLE I
NÚMERO MEDIO DE RONDAS DE NEGOCIACIÓN PARA ALCANZAR EL
CONSENSO Y LA CONSISTENCIA

#Participantes	20	40	60	80	100
Modelo propuesto	9.9	10.7	11.2	13.3	16.1
Modelo granular de [4]	13.2	18.9	29.4	45.6	76.5
Modelo granular de [3]	15.7	23.7	38.9	61.5	98.8

sobre los modelos granulares propuestos en [3] y [4], ya que, en media, necesita un menor número de rondas de negociación para llegar a un consenso y mejorar la consistencia. Aunque el modelo presentado en [4] se basa en una distribución no uniforme de la granularidad de la información—lo cual se ha demostrado que obtiene mejor resultados que una distribución uniforme en términos de obtención de consenso—el modelo propuesto reduce de forma efectiva las rondas de negociación gracias a la implementación de mecanismos basados en la tecnología *blockchain*. Adicionalmente, la robustez del modelo propuesto se incrementa con un mayor número de participantes, posibilitando converger de forma más rápida a los niveles de consenso y consistencia deseados. Por el contrario, los modelos granulares existentes encuentran dificultades para alcanzar esta convergencia.

Como extensión de este modelo de consenso, se propone su integración con el paradigma de la inteligencia artificial explicable [6]. En la mayoría de los modelos de toma de decisiones en grupo, el mecanismo que subyace detrás del proceso de decisión no está claro, lo que da lugar a que se vea como una «caja negra», donde es imposible saber el razonamiento seguido para llegar a un cierto resultado (decisión). Por tanto, para construir modelos de toma de decisiones en grupo cuyos resultados sean comprensibles por sus usuarios, el paradigma de la inteligencia artificial explicable debería integrarse en estos modelos. Esto no solo informaría a los participantes sobre la mejor decisión, sino también sobre las razones y los criterios empleados para llegar a esa conclusión.

REFERENCES

- [1] X. Yi, X. Yang, A. Kelarev, K. Y. Lam, and Z. Tari, *Blockchain Foundations and Applications*. Springer International Publishing, 2022.
- [2] I. L. Janis, *Victims of Groupthink: A Psychological Study of Foreign-Policy Decisions and Fiascoes*. Boston, MA, USA: Houghton Mifflin Company, 1972.
- [3] F.J. Cabrerizo, R. Ureña, W. Pedrycz, and E. Herrera-Viedma, “Building consensus in group decision making with an allocation of information granularity,” *Fuzzy Sets Syst.*, vol. 255, pp. 115–127, Nov. 2014.
- [4] J. C. González-Quesada, A. Velykorusova, A. Banaitis, A. Kaklauskas, and F. J. Cabrerizo, “Non-uniform allocation of information granularity to improve consistency and consensus in multi-criteria group decision-making: Application to building refurbishment,” *Eng. Appl. Artif. Intell.*, vol. 130, Art. no. 107737, Apr. 2024.
- [5] J. C. González-Quesada, I. J. Pérez, E. Herrera-Viedma, and F. J. Cabrerizo, “A study of different protocols of distribution of information granularity to build consensus in fuzzy group decision-making,” in 57th Hawaii Int. Conf. Syst. Sci., Honolulu, USA, 2024, pp. 1754–1763.
- [6] A. Barredo Arrieta, N. Díaz-Rodríguez, Javier Del Ser, A. Bannetot, S. Tabik, A. Barbado, S. García, S. Gil-López, D. Molina, R. Benjamins, R. Chatila, and F. Herrera, “Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI,” *Inf. Fusion*, vol. 58, pp. 82–115, Jun. 2020.

Un nuevo modelo de consenso que unifica y coordina modelos generativos de lenguaje y experiencia humana

1st F.J. Cabrerizo

dept. de Ciencias de la Computación e IA
Universidad de Granada
Granada, España
cabrerizo@decsai.ugr.es

2nd M.J. Cobo

dept. de Ciencias de la Computación e IA
Universidad de Granada
Granada, España
mjcobo@decsai.ugr.es

3rd J. Bernabe-Moreno

IBM Research
Dublin, Irlanda
Juan.Bernabe-Moreno@ibm.com

4th E. Herrera-Viedma

dept. de Ciencias de la Computación e IA
Universidad de Granada
Granada, España
viedma@decsai.ugr.es

5th I.J. Perez

dept. de Ciencias de la Computación e IA
Universidad de Granada
Granada, España
ijperez@decsai.ugr.es

Abstract—Los modelos tradicionales de toma de decisiones en grupo han utilizado con frecuencia herramientas de inteligencia artificial, como la lógica difusa, para actuar como moderadores entre múltiples expertos humanos encargados de seleccionar la mejor opción entre un conjunto de alternativas. Estos modelos facilitaban el consenso al interpretar e integrar diversas opiniones expertas. Sin embargo, la aparición de modelos generativos de inteligencia artificial introduce un nuevo desafío: cómo integrar estos modelos avanzados como participantes activos en el grupo de decisión, junto a los expertos humanos. Este artículo propone un nuevo marco híbrido de consenso que integra modelos generativos de lenguaje de gran escala como contribuyentes clave en el proceso de negociación. Aprovechando las fortalezas únicas tanto de la experiencia humana como de los conocimientos impulsados por la inteligencia artificial, nuestro modelo tiene como objetivo mejorar la solidez y la eficiencia de la toma de decisiones en grupo. Exploramos métodos para integrar eficazmente los modelos generativos, abordando los posibles sesgos y garantizando una colaboración coherente entre participantes humanos y de IA.

Este artículo ha sido presentado en “The Hawaii International Conference on System Sciences (HICSS)” que tuvo lugar en Hawaii, Estados Unidos, en Enero de 2025 [1]. Dicha conferencia está clasificada como rating A en el GII-GRIN-SCIE (GGS) Conference Rating.

I. INTRODUCCIÓN

La toma de decisiones en grupo (TDG) es un proceso complejo que implica integrar opiniones y conocimientos diversos para alcanzar un consenso.

Los enfoques tradicionales de toma de decisiones en grupo a menudo presentan dificultades para gestionar eficazmente

las distintas perspectivas y garantizar que todas las voces sean adecuadamente consideradas, lo que puede derivar en conflictos y resultados no óptimos.

La lógica difusa es una herramienta valiosa de la inteligencia artificial (IA) para apoyar la toma de decisiones en grupo, al modelar el concepto de consenso suave. Este enfoque permite también representar los conceptos de incertidumbre y de verdad parcial inherentes a las opiniones y juicios humanos. Los sistemas basados en lógica difusa pueden agregar estas aportaciones matizadas para facilitar un proceso de consenso más flexible e inclusivo, sirviendo de puente entre las reglas de decisión estrictas y la complejidad del razonamiento humano.

La aparición de los modelos generativos de lenguaje de gran escala (LLMs) ofrece una nueva dimensión para mejorar los marcos de toma de decisiones en grupo. Por un lado, estas nuevas herramientas de IA pueden emplearse como complemento a los modelos tradicionales de toma de decisiones en grupo, enriqueciéndolos con capacidades como el procesamiento del lenguaje natural y el análisis de sentimientos, que permiten un tratamiento más profundo de las preferencias de los expertos. Algunos autores utilizan análisis de sentimientos y modelos generativos para analizar el comportamiento de los expertos durante el debate, midiendo el nivel de positividad y agresividad de sus comentarios, y empleándolo posteriormente para ajustar el peso de sus valoraciones.

Por otro lado, gracias a sus avanzadas capacidades de comprensión y generación del lenguaje natural, los modelos generativos de lenguaje pueden actuar como participantes sofisticados en el proceso de consenso. Al integrar los LLMs como contribuyentes activos, es posible crear un modelo híbrido en el que tanto los expertos humanos como estos agentes artificiales colaboren activamente en la toma de decisiones.

Esta contribución ha sido financiada por el proyecto PID2022-139297OB-I00, financiado por MICIU/AEI/10.13039/501100011033 y por FEDER/UE. Además, forma parte del proyecto C-ING-165-UGR23, cofinanciado por la Consejería de Universidad, Investigación e Innovación y por la Unión Europea en el marco del Programa FEDER de Andalucía 2021-2027.

II. RAZONAMIENTO Y DEBATE CON LLMs

Los modelos preentrenados de gran tamaño exhiben distintas capacidades debido a la diversidad de sus datos de entrenamiento y variaciones en su arquitectura. Por ejemplo, la técnica del “*Mixture of Experts*”, un método popular de aprendizaje mediante ensamblado, entrena múltiples modelos pequeños y especializados para aumentar la robustez y la precisión de los resultados.

En el caso de uso de los modelos de lenguaje como expertos, la técnica “*Self-Consistency*” genera múltiples caminos de razonamiento utilizando “*Chain-of-Thought*” (CoT) y selecciona la respuesta más consistente como salida final. Otros autores proponen “*LLM-Blender*”, un método para clasificar y combinar salidas de diferentes modelos.

Los avances en LLMs han catalizado el desarrollo de técnicas sofisticadas de “*prompting*” y “*fine tuning*” para abordar problemas de razonamiento. Aunque obtener razonamientos de un único agente resulta prometedor, está limitado por la falta de diversidad en los puntos de vista.

A diferencia de estos enfoques, otros autores analizan la comunicación mediante explicaciones entre diferentes agentes LLM y su capacidad para debatir y persuadirse mutuamente con el fin de mejorar el razonamiento colectivo. De esta forma, se presenta un marco llamado ReConcile, diseñado para potenciar las capacidades de razonamiento de los LLMs mediante una “*mesa redonda*” entre diferentes agentes LLM. ver figura 1.

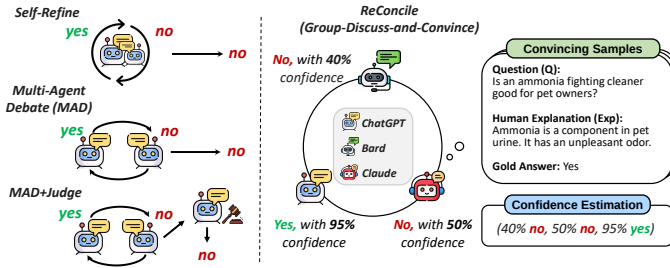


Fig. 1. Resumen de los trabajos previos

III. ARQUITECTURA DEL MARCO Y FLUJO DEL PROCESO

A. Arquitectura del marco

Como orquestador principal de la arquitectura propuesta, se considera el uso de LangChain. Este framework de código abierto nos proporciona una infraestructura flexible y escalable que simplifica el despliegue y la coordinación de múltiples modelos de IA. Su arquitectura modular permite la integración fluida de diferentes LLMs, facilitando la experimentación con distintos conjuntos de modelos para determinar las combinaciones más efectivas.

Además, el sólido soporte de LangChain para el procesamiento y la comprensión del lenguaje natural permite desarrollar protocolos de comunicación sofisticados entre agentes de IA y participantes humanos. Esto incrementa la coherencia y efectividad general del proceso de toma de decisiones. Al

aprovechar las ventajas de LangChain, nuestro sistema puede combinar eficazmente las fortalezas tanto de la inteligencia artificial como del conocimiento humano, garantizando un enfoque más robusto y dinámico para intentar garantizar el consenso.

B. Flujo del proceso

- 1: **Entrada:** Conjunto de alternativas $\{Alt_1, Alt_2, \dots, Alt_n\}$ y criterios $\{Crit_1, Crit_2, \dots, Crit_n\}$
- 2: **Salida:** Alternativas finales seleccionadas $\{Alt_{selected}\}$
- 3: **Revisión inicial:**
- 4: **for** cada alternativa Alt_i **do**
- 5: **Expertos humanos:** Asignar la alternativa Alt_i a un subconjunto de expertos humanos
- 6: Cada experto revisa Alt_i y proporciona sus puntuaciones iniciales y comentarios adicionales
- 7: **LLMs:** Usar LangChain para revisar Alt_i conectando con diferentes modelos preentrenados mediante APIs
- 8: Los LLMs proporcionan una valoración preliminar teniendo en cuenta los distintos criterios
- 9: **end for**
- 10: **Procesamiento preliminar de preferencias:**
- 11: Combinar las puntuaciones y comentarios de los expertos humanos y los LLMs
- 12: Normalizar puntuaciones y comprobar consistencia
- 13: Identificar sesgos en las evaluaciones de los LLMs
- 14: Comprobar el nivel de consenso
- 15: **Proceso de consenso:**
- 16: **repeat**
- 17: **Fase de deliberación:** Realizar una reunión virtual con expertos humanos y LLMs
- 18: Los expertos discuten sus opiniones y los LLMs aportan conocimientos adicionales
- 19: **Negociación y reevaluación:** Se pide a los expertos y LLMs que reevalúen las alternativas en base a nuevos criterios o perspectivas
- 20: Reevaluar alternativas según los nuevos datos
- 21: Combinar las nuevas puntuaciones y comentarios
- 22: Normalizar las puntuaciones, comprobar consistencia y abordar nuevos posibles sesgos
- 23: **until** se alcance el nivel de consenso requerido o se complete un número predeterminado de iteraciones
- 24: **Proceso de Selección final:**
- 25: Sintetizar las recomendaciones combinadas de expertos humanos y LLMs
- 26: Clasificar las alternativas según las puntuaciones finales
- 27: Revisión final para verificar coherencia y alineación con los criterios de decisión
- 28: **Salida:** Lista de alternativas seleccionadas $\{Alt_{selected}\}$

REFERENCES

- [1] F.J. Cabrerizo, J. Bernabé-Moreno, E. Herrera-Viedma, I.J. Pérez. “An Innovative Hybrid Soft Consensus Framework Leveraging Generative Large Language Models and Human Expertise,” Proceedings of the 58th Hawaii International Conference on System Sciences, HICSS 2025 pp 1793-1802.

Salvando ratones: ChenseNet121, una arquitectura de aprendizaje profundo para la estimación de toxicidad aguda (LD50)

Enol Junquera Álvarez*, Beatriz Remeseiro*, Ferdinando Febbraio†, Irene Díaz*

*Departamento de Informática, Universidad de Oviedo, España

Email: enoljunquera@uniovi.es

†Institute of Biochemistry and Cell Biology, CNR, Italia

Resumen—Presentamos un modelo de inteligencia artificial orientado a reducir la experimentación animal en toxicología. ChenseNet121 es una arquitectura profunda multimodal capaz de predecir la toxicidad aguda oral (LD50) de compuestos químicos, integrando información estructural y funcional procedente de diversas representaciones moleculares. La arquitectura alcanza una precisión competitiva en una comparativa completa, mediante una clasificación *post hoc*, identificar sustancias de bajo riesgo con alta confianza, contribuyendo así a reemplazar ensayos *in vivo* innecesarios. El modelo está listo para validación experimental en laboratorio, a falta únicamente de una interfaz de usuario accesible.

Index Terms—Toxicología computacional, LD50, aprendizaje profundo, multimodalidad, reducción de experimentación animal.

I. INTRODUCCIÓN

El test de toxicidad aguda oral (LD50) es una de las pruebas más comunes y controvertidas de la toxicología reglamentaria, tradicionalmente realizada con animales de laboratorio [1]. Cada nuevo pesticida, por ejemplo, implica decenas de ratas en protocolos de dosificación letal escalonada. Frente a este escenario, la combinación de ciencia de datos, biocomputación y *deep learning* permite hoy imaginar un camino alternativo [2]. En este trabajo presentamos ChenseNet121, una red neuronal profunda capaz de estimar el valor de LD50 a partir de información estructural de la molécula. Reduciendo la necesidad de experimentación animal, con sus correspondientes ventajas éticas y económicas.

II. METODOLOGÍA

II-A. Dataset

El conjunto de datos empleado en este estudio fue específicamente construido para facilitar la predicción de toxicidad aguda (LD50) de compuestos químicos, integrando información multimodal topológica, fisicoquímica y espacial. A diferencia de los conjuntos tradicionales, este dataset combina múltiples representaciones moleculares, lo que permite a redes neuronales profundas modelar patrones tóxicos desde distintas modalidades. El desarrollo del dataset, junto con su descripción técnica, ha sido publicado en el *EFSA Journal* como parte de un trabajo previo [3].

Los valores de LD50 fueron extraídos de documentos técnicos y guías regulatorias centradas en pesticidas [4]–[7]. Estos

datos se derivan principalmente de ensayos estandarizados en animales, como ratas. Las estructuras moleculares fueron representadas en formato SMILES, obtenidas desde la base de datos pública PubChem. A partir de esta notación se generaron descriptores moleculares clásicos usando la biblioteca RDKit, incluyendo: peso molecular, logP, área polar superficial topológica (TPSA), número de enlaces rotables y energía de unión estimada tras *docking*. Estos valores constituyen una de las tres modalidades de entrada. La segunda modalidad corresponde a imágenes 2D de las moléculas, descargadas directamente de PubChem. La tercera modalidad emplea información tridimensional simulada mediante *docking* flexible sobre la proteína diana humana acetilcolinesterasa (PDB ID: 7E3H). A partir del complejo proteína-ligando se generaron volúmenes voxelizados de tamaño $25 \times 25 \times 25$ con 8 canales por voxel, incluyendo energías atractivas (Gauss1, Gauss2), repulsión, hidrofobicidad y presencia de sitios activos (donadores y aceptores de enlaces de hidrógeno, etc.). Su nivel de afinidad se guardó como descriptor numérico. Este proceso se realizó con AutoDock Vina, Open Babel y RDKit, siguiendo una estrategia inspirada en RosENet [8].

Cada compuesto queda representado así mediante tres vistas complementarias:

- Una imagen estructural 2D.
- Un vector numérico con descriptores moleculares.
- Un volumen 3D voxelizado que captura interacciones con la proteína diana.

II-B. Arquitectura

El diseño de la arquitectura parte de una premisa clave: no existía un modelo previo específicamente adaptado a la predicción de toxicidad aguda a partir de datos químicos calculados *in silico*. Por ello, se propuso una familia de arquitecturas derivadas de redes convolucionales conocidas por su eficacia en visión por computador, adaptadas a la estructura del nuevo dataset construido.

Se consideraron inicialmente variantes de arquitecturas clásicas como ResNet50 [9], InceptionV3 [10], DenseNet121 [11] y EfficientNetV2-M [12], valoradas por su desempeño en ImageNet y por tanto en clasificación basada en patrones no temporales. De todas ellas, DenseNet121 ofreció los resultados más prometedores. Esto condujo al desarrollo de

una arquitectura derivada denominada **ChenseNet121** (Chemical DenseNet121).

ChenseNet121 consta de tres ramas paralelas especializadas:

- Una rama convolucional 2D basada en DenseNet121 para procesar imágenes moleculares.
- Una rama convolucional 3D para mapas de energía voxelizados tras *docking*.
- Una rama densa para descriptores fisicoquímicos y puntuaciones de afinidad.

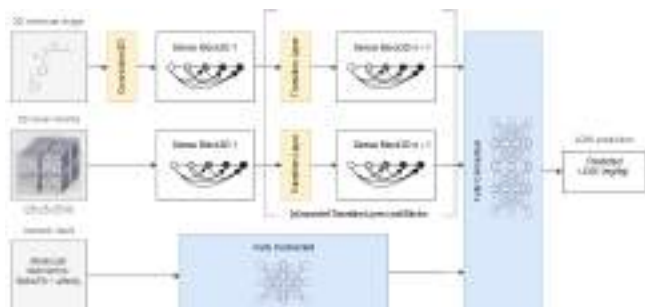


Figura 1. Esquema de la arquitectura ChenseNet121. Cada rama procesa una modalidad distinta antes de fusionarse en una capa densa multimodal.

III. RESULTADOS

ChenseNet121 supera sistemáticamente a otras arquitecturas (ResNet, Inception, EfficientNet) adaptadas al contexto químico [13], [14]. En configuración multimodal completa. El MAE de 1568.68 mg/kg, equivalente aproximadamente al 9.2 % del rango total (0–17,000 mg/kg).

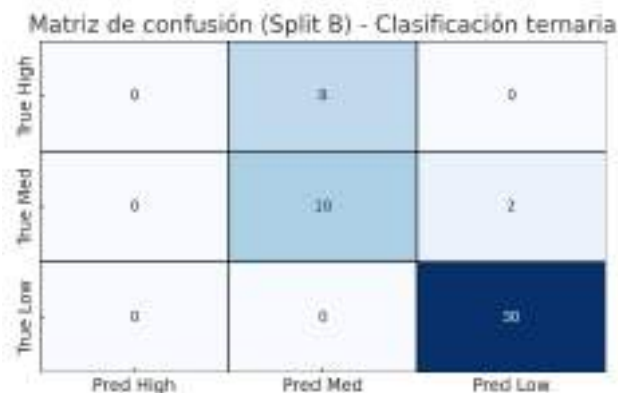


Figura 2. Matriz de confusión en clasificación ternaria (Split B). El modelo tiende a sobrepredecir toxicidad moderada, pero acierta plenamente en compuestos de baja toxicidad.

En clasificación binaria, suponiendo únicamente compuestos tóxicos frente a no-tóxicos (baja toxicidad), se obtuvo un F1-score de 0.96 y precisión del 96 %. Siendo ningún compuesto altamente tóxico considerado de baja toxicidad.

IV. DISCUSIÓN Y CONCLUSIONES

Los resultados muestran que la integración de datos estructurales 2D, numéricos y espaciales 3D permite capturar patrones de toxicidad complejos, más allá de lo visible en una sola

representación. El modelo resulta especialmente conservador, priorizando la seguridad sobre la sensibilidad.

Salvando ratones. Si ChenseNet121 evita el ensayo in vivo de solo 15 compuestos, se estima un ahorro de más de 300 ratones de laboratorio y 37.000€ en costes. La toxicología computacional, por tanto, ofrece no solo eficiencia económica, sino un avance ético ineludible.

Este trabajo se enmarca en una línea de investigación activa centrada en IA explicable, visualización de contribuciones moleculares, y reducción sistemática de la experimentación animal en fases tempranas del desarrollo químico.

El modelo se ha diseñado para integrarse como filtro previo en flujos regulatorios. Se requiere un aumento mecánico en los datos del dataset en combinación molecular con más proteínas que se espera mejoren la capacidad de detección de alta toxicidad. Actualmente como diferenciador de baja toxicidad, está listo para validación en laboratorio, y su siguiente paso es la creación de una interfaz gráfica accesible que permita su uso por toxicólogos y personal no técnico.

AGRADECIMIENTOS

Este trabajo ha sido parcialmente financiado en parte por el proyecto EUFORA. Agradecemos al CNR y a EFSA por los recursos compartidos.

REFERENCIAS

- [1] X. Li, M. Zhang, J. Zhao, and J. Ma, "Toxicity prediction of chemicals using ensemble learning," *BMC Bioinformatics*, vol. 21, no. 1, pp. 1–10, 2020.
- [2] E. N. Muratov, J. Bajorath, R. P. Sheridan, I. V. Tetko, D. Filimonov, V. Poroikov, T. I. Oprea, I. I. Baskin, A. Varnek, A. Roitberg *et al.*, "Qsar without borders," *Chemical Society Reviews*, vol. 49, no. 11, pp. 3525–3564, 2020.
- [3] E. Junquera *et al.*, "A curated dataset for ld50 estimation using multi-modal deep learning: chemical, spatial, and biological features," *EFSA Journal*, vol. 22, no. 1, p. e08923, 2024.
- [4] E. F. S. Authority, "Efsa guidance document on risk assessment for birds and mammals," 2009. [Online]. Available: <https://www.efsa.europa.eu>
- [5] U. E. N. C. for Computational Toxicology, "Toxrefdb: Toxicity reference database," 2007. [Online]. Available: <https://www.epa.gov/chemical-research/toxicity-forecasting>
- [6] U. E. P. Agency, "Epa: Guidelines for carcinogen risk assessment," 2009. [Online]. Available: <https://www.epa.gov/risk/guidelines-carcinogen-risk-assessment>
- [7] —, "Health effects test guidelines: Oppts 870.1100 acute oral toxicity," 1998. [Online]. Available: <https://www.epa.gov>
- [8] L. Zhao and B. Huang, "Rosenet: Improving binding affinity prediction by leveraging molecular mechanics energies with an ensemble of 3d-cnns," *Journal of Cheminformatics*, vol. 13, no. 1, pp. 1–14, 2021.
- [9] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," pp. 770–778, 2016.
- [10] C. Szegedy, V. Vanhoucke, S. Ioffe, J. Shlens, and Z. Wojna, "Rethinking the inception architecture for computer vision," pp. 2818–2826, 2016.
- [11] G. Huang, Z. Liu, L. Van Der Maaten, and K. Q. Weinberger, "Densely connected convolutional networks," pp. 4700–4708, 2017.
- [12] M. Tan and Q. V. Le, "Efficientnetv2: Smaller models and faster training," *arXiv preprint arXiv:2104.00298*, 2021.
- [13] Y. Yu, Y. Wang, J. Sun, Y. Cai, Y. Ma, S. Gao, W. Zhang, T. Huang, and X. Kong, "Deep learning for toxicity prediction: from chemical structure to transcriptome," *Scientific Reports*, vol. 12, no. 1, pp. 1–11, 2022.
- [14] A. Mayr, G. Klambauer, T. Unterthiner, and S. Hochreiter, "Deeptox: Toxicity prediction using deep learning," *Frontiers in Environmental Science*, vol. 3, p. 80, 2016.

Pooling over-time no-simétrico mediante funciones de pseudo-grouping en redes neuronales convolucionales

M. Ferrero-Jaurrieta, L. De Miguel, C. Lopez-Molina, H. Bustince
Dpto. de Estadística, Informática y Matemáticas
Universidad Pública de Navarra
Pamplona, España

A. Cruz, B. Bedregal
Universidade Federal do Rio Grande do Norte
Natal, Brasil

R. Paiva
Departamento de Física e Matemática
Instituto Federal de Educação, Ciência e Tecnologia do Ceará
Fortaleza, Brasil

Z. Takáč
Faculty of Materials Science and Technology in Trnava
Slovak University of Technology
Trnava, Slovakia

Abstract—Este artículo es un resumen del trabajo publicado en la revista *Engineering Applications of Artificial Intelligence* [1]. Las redes neuronales convolucionales utilizadas para el procesamiento de texto (TextCNN) fusionan características locales para generar características globales en sus capas de pooling over-time. Estas capas se han construido tradicionalmente utilizando funciones simétricas. Es importante señalar que el orden de las características locales de entrada es significativo (es decir, la simetría no es una característica inherente al modelo). Esta característica es contraproducente para el procesamiento de texto, donde el orden de las palabras es relevante. En este trabajo proponemos utilizar operadores no-simétricos para sustituir a las funciones de máximo o media aritmética. En concreto, proponemos realizar el proceso de pooling over-time utilizando funciones de pseudo-grouping, una familia de operadores de agregación no simétricos que generalizan la función máximo. Presentamos un método de construcción de funciones de pseudo-grouping y se aplican en capas de pooling over-time en TextCNNs en distintas arquitecturas y datasets de clasificación de texto.

Palabras clave—Funciones de agregación no simétricas, clasificación de texto, pooling over-time, funciones de pseudo-grouping.

I. INTRODUCCIÓN

En primer lugar, realizamos algunas definiciones necesarias. Una función $h: [0, 1] \rightarrow [0, 1]$ es un pseudo-automorfismo [2] si es no-creciente, continua, $h(x) = 0$ ssi $x = 0$ y $h(x) = 1$ ssi $x = 1$. Por otro lado, una función $N: [0, 1] \rightarrow [0, 1]$ es una negación difusa si es decreciente, $N(0) = 1$ y $N(1) = 0$; y estricta si es una función continua y estrictamente decreciente. Sea $\mathbf{x} = (x_1, \dots, x_n)$, para todo $\mathbf{x} \in [0, 1]^n$.

Definición 1.1 ([3]): Una función de pseudo-overlap (FPO) es una aplicación $PO: [0, 1]^n \rightarrow [0, 1]$ creciente y continua tal que para todo $\mathbf{x} \in [0, 1]^n$:

- 1) $PO(\mathbf{x}) = 0$ ssi $\prod_{i=1}^n x_i = 0$,
- 2) $PO(\mathbf{x}) = 1$ ssi $\prod_{i=1}^n x_i = 1$.

Este trabajo ha sido financiado por el proyecto PID2022-136627NB-I00, MCIN/AEI/10.13039/501100011033/FEDER, UE.

Definición 1.2 ([3]): Una función de pseudo-grouping (FPG) es una aplicación $PG: [0, 1]^n \rightarrow [0, 1]$ creciente y continua tal que para todo $\mathbf{x} \in [0, 1]^n$:

- 1) $PG(\mathbf{x}) = 0$ ssi $x_i = 0$, para todo $i \in \{1, \dots, n\}$;
- 2) $PG(\mathbf{x}) = 1$ ssi existe $i \in \{1, \dots, n\}$ tal que $x_i = 1$.

Ejemplo 1.3: La función $PG_{\min}^r: [0, 1]^n \rightarrow [0, 1]$ dada por $PG_{\min}^r(\mathbf{x}) = 1 - \min_{i=1}^n \{(1 - x_i)^{r_i}\}$, tal que $r_i \geq 0$, para todo $i \in \{1, \dots, n\}$ y $\mathbf{x} \in [0, 1]^n$ es una FPG.

II. FPGS OBTENIDAS MEDIANTE DISTORSIÓN

Definición 2.1: Sean $g, h_1, \dots, h_n: [0, 1] \rightarrow [0, 1]$ pseudo-automorfismos y $A: [0, 1]^n \rightarrow [0, 1]$ una función de agregación. La función $B: [0, 1]^n \rightarrow [0, 1]$ dada por $B(\mathbf{x}) = g(A(h_1(x_1), \dots, h_n(x_n)))$ es una función de agregación obtenida mediante $(g, A, h_1, \dots, h_n)$ -distorsión.

Proposición 2.2: La aplicación $B: [0, 1]^n \rightarrow [0, 1]$ dada en la Definición 2.1 es una FPG si y solo si A es una FPG.

Ejemplo 2.3: Consideramos los pseudo-automorfismos $g(x) = x$ y $h_i(x) = x^{r_i}$ tal que $r_i \geq 0$ para todo $i \in \{1, \dots, n\}$ y $PG: [0, 1]^n \rightarrow [0, 1]$ dado por $PG(\mathbf{x}) = 1 - \prod_{i=1}^n (1 - x_i)$ una FPG. La función $PG_{\text{prod 1}}^r(\mathbf{x}) = 1 - \prod_{i=1}^n (1 - x_i^{r_i})$ es una FPG, donde $r_i \geq 0$ para todo $i \in \{1, \dots, n\}$.

La función $PG_{\text{prod 1}}$ es no-simétrica si se cumple $r_i \neq r_j$, para algún $i, j \in \{1, \dots, n\}$, tal que $i \neq j$.

Corolario 2.4: Si $\eta, N_1, \dots, N_n: [0, 1] \rightarrow [0, 1]$ son negaciones estrictas y $PO: [0, 1]^n \rightarrow [0, 1]$ es una FPO, la función $PG(\mathbf{x}) = \eta(PO(N_1(x_1), \dots, N_n(x_n)))$ es una FPG obtenida mediante $(\eta, PO, N_1, \dots, N_n)$ -distorsión.

Ejemplo 2.5: Consideramos la función $PO(\mathbf{x}) = \prod_{i=1}^n x_i^{r_i}$, donde $r_i \geq 0$, para todo $i \in \{1, \dots, n\}$ y $\eta(x) = N_i(x) = 1 - x$, para todo $i \in \{1, \dots, n\}$. La función $PG_{\text{prod 2}}^r(\mathbf{x}) = 1 - \prod_{i=1}^n (1 - x_i)^{r_i}$ es una FPG, donde $r_i \geq 0$, para todo $i \in \{1, \dots, n\}$.

III. FUSIÓN DE CARACTERÍSTICAS DE POOLING OVER-TIME MEDIANTE FPG

Con respecto a los modelos utilizados, la arquitectura tipo utilizada en la experimentación sigue la siguiente estructura:

- *Embedding*. Transforma cada palabra en un vector d -dimensional. Se utilizan embeddings clásicos [4] (GloVe-6B, GloVe-Twitter27B y FastText) así como basados en transformers derivados de BERT [5] (BERT, ALBERT, RoBERTa y DeBERTa).
- *Convolución*. Sea $F \in \mathbb{N}$ el número de filtros de convolución y $\mathcal{Z} = \{z_1, \dots, z_S\}$ el conjunto de ventanas. Los vectores de características generados son $\mathbf{c}^{j,z} \in [0, \infty)^{m-z+1}$, para todo $z \in \mathcal{Z}$ y $j \in \{1, \dots, F\}$.
- *Pooling over-time*. Sea PG una FPG para $n = m - z + 1$, el valor agregado se calcula de la forma. $p^{j,z} = PG(\mathbf{c}^{j,z})$ para todo $j \in \{1, \dots, m\}$ y $z \in \mathcal{Z}$.
- *Capa lineal*. Los valores agregados de desnormalizan y se concatenan para ser la entrada de una capa densa.

En la Tabla I se muestran los datasets utilizados, donde #C es el número de clases; #inst es el número de instancias; y #test, el tamaño de la partición de test de cada dataset.

TABLA I
DATASETS UTILIZADOS

#	Name	#C	#inst	#test	#	Name	#C	#inst	#test
1	IMDB	2	50000	25000	5	TREC-CL	6	5950	500
2	SMS	2	5584	2240	6	TREC-FL	50	5950	500
3	TweetEval	2	14010	860	7	BigPatent	9	45000	5000
4	NewsData	6	5518	828					

Con respecto a los resultados numéricos, en la Tabla II mostramos los resultados obtenidos, utilizando el embedding GloVe-6B para los 7 datasets. En la Tabla III mostramos los resultados obtenidos utilizando distintos embeddings basados en transformers, para el Dataset 1. En ambos casos se muestra la media aritmética y desviación estándar para 10 ejecuciones independientes. A nivel numérico se obtiene la mejor métrica de F_1 -score y accuracy cuando utilizamos FPGs en lugar de funciones simétricas. Todos los mejores resultados numéricos basados en FPGs de las Tablas II y III son estadísticamente significativos (evaluados mediante el test estadístico de Mann Whitney U), excepto en los que se utiliza BERT (Tabla III).

A nivel de análisis, en el artículo se utilizan mapas de calor para comparar la fusión de información tras aplicar una FPG y máximo en la capa de pooling over-time por lotes del conjunto de validación. Se observa que la distribución es más dispersa cuando se aplica FPG en comparación con max, obteniéndose así una serie de características globales más definidas, debido a que para toda FPG, $PG(\mathbf{x}) \geq \max(\mathbf{x})$, para todo $\mathbf{x} \in [0, 1]^n$. Cuando se utiliza una FPG, la diferencia obtenida al utilizar diferentes tamaños de ventana en los filtros de convolución tiene un mayor impacto en comparación con la función máximo. Se observa que hay una mayor discriminación de rasgos a medida que la ventana utilizada en la convolución es mayor.

IV. OBSERVACIONES FINALES

A nivel teórico, a generalización de los métodos de construcción encontrados en la literatura a partir de distorsiones

TABLA II
ACCURACY Y F_1 -SCORE OBTENIDOS CON GLOVE-6B

		Dataset 1		Dataset 2		Dataset 3	
		Acc	F_1	Acc	F_1	Acc	F_1
max		.855 $\pm$.002	.819 $\pm$.003	.893 $\pm$.002	.860 $\pm$.003	.889 $\pm$.000	.856 $\pm$.001
Med. Ar.		.816 $\pm$.016	.778 $\pm$.020	.865 $\pm$.001	.832 $\pm$.000	.858 $\pm$.005	.824 $\pm$.009
$PG_{prod 1}^r$		.892 $\pm$.001	.859 $\pm$.001	.899 $\pm$.001	.867 $\pm$.001	.894 $\pm$.002	.861 $\pm$.002
PG_{min}^r		.877 $\pm$.003	.844 $\pm$.003	.894 $\pm$.001	.862 $\pm$.001	.888 $\pm$.001	.855 $\pm$.003
$PG_{prod 2}^r$		.883 $\pm$.001	.849 $\pm$.002	.898 $\pm$.002	.865 $\pm$.001	.886 $\pm$.007	.853 $\pm$.011
		Dataset 4		Dataset 5			
		Acc	F_1	Acc	F_1		
max		.865 $\pm$.002	.855 $\pm$.003	.909 $\pm$.005	.907 $\pm$.006		
Med. Ar.		.866 $\pm$.002	.860 $\pm$.002	.816 $\pm$.004	.805 $\pm$.004		
$PG_{prod 1}^r$		.870 $\pm$.001	.866 $\pm$.002	.902 $\pm$.004	.900 $\pm$.004		
PG_{min}^r		.865 $\pm$.003	.856 $\pm$.003	.915 $\pm$.004	.914 $\pm$.004		
$PG_{prod 2}^r$		.874 $\pm$.002	.869 $\pm$.002	.908 $\pm$.002	.907 $\pm$.002		
		Dataset 6		Dataset 7			
		Acc	F_1	Acc	F_1		
max		.830 $\pm$.006	.811 $\pm$.006	.625 $\pm$.003	.607 $\pm$.004		
Med. Ar.		.772 $\pm$.006	.755 $\pm$.007	.605 $\pm$.003	.589 $\pm$.004		
$PG_{prod 1}^r$		.835 $\pm$.009	.819 $\pm$.011	.635 $\pm$.004	.617 $\pm$.006		
PG_{min}^r		.841 $\pm$.007	.824 $\pm$.008	.632 $\pm$.004	.613 $\pm$.007		
$PG_{prod 2}^r$		.830 $\pm$.009	.813 $\pm$.012	.630 $\pm$.003	.612 $\pm$.003		

TABLA III
ACCURACY Y F_1 -SCORE OBTENIDOS PARA EL DATASET 1.

		BERT		ALBERT	
		Acc	F_1	Acc	F_1
max		.928 $\pm$.001	.865 $\pm$.000	.892 $\pm$.002	.830 $\pm$.001
Med. Ar.		.906 $\pm$.001	.846 $\pm$.000	.860 $\pm$.003	.799 $\pm$.005
$PG_{prod 1}^r$		.928 $\pm$.002	.864 $\pm$.002	.901 $\pm$.002	.837 $\pm$.002
PG_{min}^r		.929 $\pm$.001	.866 $\pm$.001	.893 $\pm$.002	.832 $\pm$.002
$PG_{prod 2}^r$		.925 $\pm$.002	.862 $\pm$.002	.893 $\pm$.001	.831 $\pm$.003
		RoBERTa		DeBERTa	
		Acc	F_1	Acc	F_1
max		.930 $\pm$.001	.875 $\pm$.001	.939 $\pm$.001	.849 $\pm$.001
Med. Ar.		.887 $\pm$.002	.835 $\pm$.002	.918 $\pm$.001	.831 $\pm$.001
$PG_{prod 1}^r$		.932 $\pm$.001	.877 $\pm$.002	.941 $\pm$.001	.852 $\pm$.001
PG_{min}^r		.929 $\pm$.001	.875 $\pm$.001	.941 $\pm$.001	.852 $\pm$.001
$PG_{prod 2}^r$		.934 $\pm$.001	.879 $\pm$.001	.937 $\pm$.003	.848 $\pm$.004

(geométricas), permite diseñar agregaciones en función de las necesidades de cada ejemplo particular. A nivel aplicado, los resultados justifican la eliminación de la conmutatividad en la fusión de información con dependencia secuencial/temporal. Como líneas futuras, a nivel teórico, se propone el desarrollo de este concepto en retículos acotados; y otro lado, se profundizará en el estudio de la simetría en la fusión de características en otros modelos de aprendizaje automático.

REFERENCIAS

- [1] M. Ferrero-Jaurrieta, R. Paiva, A. Cruz, B. Bedregal, L. De Miguel, Z. Takáč, C. Lopez-Molina, and H. Bustince, "Non-symmetric over-time pooling using pseudo-grouping functions for convolutional neural networks," *Eng. Appl. Artif. Intell.*, vol. 133, p. 108470, 2024.
- [2] G. P. Dimuro, B. Bedregal, H. Bustince, M. J. Asiain, and R. Mesiar, "On additive generators of overlap functions," *Fuzzy Sets Syst.*, vol. 287, pp. 76–96, 2016.
- [3] X. Zhang, R. Liang, H. Bustince, B. Bedregal, J. Fernandez, M. Li, and Q. Ou, "Pseudo overlap functions, fuzzy implications and pseudo grouping functions with applications," *Axioms*, vol. 11, no. 11, p. 593, 2022.
- [4] J. Pennington, R. Socher, and C. D. Manning, "GloVe: Global vectors for word representation," in *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing*, 2014, pp. 1532–1543.
- [5] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," in *Proceedings of the 2019 conference of the North American chapter of the Association for Computational Linguistics*, 2019, pp. 4171–4186.

Generando explicaciones fiables por medio de un modelo de lenguaje enriquecido por sistemas basados en reglas borrosas

Pablo Miguel Perez-Ferreiro, Alejandro Catala, Alberto Bugarín-Diz, Jose M. Alonso-Moral

Centro Singular de Investigación en Tecnologías Inteligentes (CITIUS)

Universidade de Santiago de Compostela

Santiago de Compostela, Spain

{pablo.perez.ferreiro, alejandro.catala, alberto.bugarin.diz, josemaria.alonso.moral}@usc.es

Resumen—En este *keyword* presentamos un resumen del trabajo titulado *Generating trustworthy explanations with a language model enriched by fuzzy rule-based systems*, aceptado para publicación en las actas del congreso *2025 IEEE International Conference on Fuzzy Systems* en julio de 2025 [1]. La versión inicial del trabajo fue enviada el 20 de enero de 2025, y la versión final fue aceptada el 3 de abril de 2025.

Index Terms—Grandes Modelos de Lenguaje, Lógica Borrosa, Inteligencia Artificial Explicable, Inteligencia Artificial Fiable, Árboles de Decisión

I. INTRODUCCIÓN

El campo de la Inteligencia Artificial se encuentra en una etapa de gran crecimiento, habiendo llegado ya a manos del grueso de la población por medio de servicios como ChatGPT. Sin embargo, esta creciente fama ha supuesto también una drástica disminución de la confianza que los usuarios tienen en algunas aplicaciones basadas en IA, debido a los riesgos que perciben como aparejados a la tecnología en sí. Este es un problema complejo y que se agrava si tenemos en cuenta que los riesgos percibidos y su severidad varían dependiendo de los ámbitos y áreas geográficas que se discutan. Por ejemplo, el mercado norteamericano se encuentra principalmente preocupado por aspectos relacionados con la fiabilidad y seguridad de las aplicaciones basadas en IA, mientras que el mercado europeo presta mayor atención a problemáticas como la transparencia, la equidad y la privacidad. Por si esto fuese poco, las vertientes de IA de mayor crecimiento son aquellas relacionadas con la IA generativa, y en concreto las aproximaciones basadas en modelos de lenguaje preentrenados a gran escala (*Large Language Models*, LLMs); una alternativa tecnológica notablemente difícil de evaluar en cuanto a sus riesgos y limitaciones, por falta de estandarización en los protocolos y por su diseño inherentemente opaco. Incluso aquellos

LLMs que obtienen valoraciones positivas en términos de seguridad y fiabilidad rara vez son modelos abiertos, lo que en cierto modo dificulta ejercer un control fino sobre su actuación y propiedades reales, al accederse a ellos principalmente como servicio.

En el trabajo “*Generating trustworthy explanations with a language model enriched by fuzzy rule-based systems*” publicado originalmente en [1], analizamos la potencial pérdida de confianza de los usuarios en aplicaciones basadas en IA como el principal elemento motivador de nuestra investigación. El objetivo era construir un sistema inteligente capaz de proporcionar explicaciones fiables, siguiendo las definiciones sobre fiabilidad aportadas por la Comisión Europea en sus “Directrices éticas para una IA fiable”. En ese sentido, el trabajo busca sobrepasar las características usuales de los sistemas explicables, en tanto a que consideramos la explicabilidad como un requisito necesario pero no suficiente por sí mismo para obtener la fiabilidad de la que hablamos, siendo esta una propiedad más completa y compleja.

Concretamente, en este trabajo buscamos potenciar explicadores previamente desarrollados con la capacidad expresiva de los LLMs. En trabajos previos, se implementaron estos explicadores haciendo uso de técnicas de IA explicable (XAI, por sus siglas en inglés) tradicionales, como son los árboles de decisión y los sistemas basados en reglas borrosas. Estos últimos son capaces de modelar adecuadamente la imprecisión inherente al lenguaje humano, lo que los hace especialmente adecuados para tareas explicativas, al ser el lenguaje natural el medio más directamente comprensible para un usuario inexperto. Sin embargo, la verbalización implementada para estos explicadores estaba basada en plantillas y reglas, lo que daba a veces un resultado artificial que podría dañar la percepción final de las explicaciones. Por otra parte, construir una explicación totalmente basada en LLMs es arriesgado, debido a la baja fiabilidad de estos por culpa del conocido fenómeno de las alucinaciones (esto es, generaciones de texto incoherentes o sin base real).

Nuestra hipótesis de partida fue que, mediante una combinación apropiada de los LLMs con nuestra base explicativa basada en lógica borrosa, las debilidades de ambas propuestas

Se reconoce el apoyo de la Consellería de Educación, Universidades y Formación Profesional de la Xunta de Galicia y del Fondo Europeo de Desarrollo Regional, FEDER Una manera de hacer Europa (ayudas “Centro de investigación de Galicia con acreditación 2024-2027 ED431G-2023/04” “Grupo de Referencia Competitiva con acreditación 2022-2025 ED431C 2022/19”). Este trabajo también cuenta con el apoyo del Ministerio de Ciencia, Innovación y Universidades de España (MCIN/AEI/10.13039/501100011033/) con las ayudas PID2021-123152OB-C21 y CNS2024-154915.

podrían cubrirse entre sí, obteniendo un mejor resultado final: el rigor y control aportados por la explicación basada en reglas borrosas paliarían la inconsistencia del LLM, mientras que la potencia expresiva del LLM verbalizaría esa explicación de manera más natural para el receptor final de la misma. En particular, nuestra validación se centró en el control de las alucinaciones, considerando que el poder expresivo de los LLMs está suficientemente demostrado en la literatura. Teniendo esto en cuenta, la hipótesis de investigación resultó ser la siguiente: “una fusión adecuada de técnicas de XAI con LLMs apropiadamente dirigidos puede reducir las alucinaciones de estos últimos, consiguiendo un sistema generalmente más fiable.”.

II. PROPUESTA

Nuestro modelo para la fusión de LLMs con técnicas de XAI tradicional se basa en la combinación de una herramienta explicativa que utiliza sistemas de reglas borrosas, capaz de integrar las explicaciones de diversas bases de conocimiento en una narrativa única verbalizada mediante plantillas, con una aplicación basada en LLMs debidamente instruida mediante técnicas como el razonamiento por cadena de pensamiento. Adicionalmente, y como fuente secundaria de información fiable, también se integra con el LLM un pequeño banco de documentos manualmente seleccionados.

En esta arquitectura, el LLM procesa la entrada del usuario en varios pasos de razonamiento, decidiendo si su herramienta explicativa o los documentos de los que dispone podrían ser útiles para responder a la consulta recibida (siendo más prioritaria la herramienta explicativa). Únicamente en caso de que ninguno de los dos recursos sea útil se produciría una generación libre por parte del modelo; ya que si cualquiera de ambos pareciera de utilidad, se obtendría de ellos la información pertinente para generar una respuesta más fiable.

Esta arquitectura fue validada a nivel técnico mediante su comparación con dos alternativas de aplicación que empleaban LLMs con un nivel de control menor: una de ellas no daba al modelo ninguna fuente de información o guía, dejándolo generar sus respuestas de forma totalmente libre; y en la otra únicamente se le proporcionaba al modelo un pequeño conjunto de ejemplos explicativos para guiar su tarea.

Estas alternativas se evaluaron con tres métricas: longitud de respuesta media (dado que los LLMs tienen una fuerte tendencia a extenderse demasiado, algo que creemos puede dañar la sensación de naturalidad), precisión predictiva, y número medio de potenciales alucinaciones. Para el cómputo de esta última, diseñamos un método de búsqueda de expresiones impropias para los valores de los atributos del caso a explicar (que son conocidos) en la narrativa final, véase, detectar elementos como “hombre” o “masculino” en ocasiones donde se está hablando de un sujeto de género femenino.

III. CASO DE USO Y VALIDACIÓN

En este trabajo, planteamos como prueba de concepto el uso del sistema propuesto en el contexto médico, considerándolo pertinente debido a que en este dominio existe una gran

demanda de fiabilidad. Para ello, nos centramos en dos dolencias, cardiopatía y diabetes, para las cuales se recopilamos y preparamos los respectivos conjuntos de datos. Tras terminar estas tareas, se reservaron un total de 7070 ejemplos de prueba para diabetes, y 5475 para cardiopatía.

Los ejemplos restantes (no reservados para la validación) se utilizaron para construir los árboles de decisión que posteriormente serían convertidos en sistemas basados en reglas borrosas mediante aproximación lingüística; integrando estos sistemas de reglas con la interfaz explicativa y la aplicación basada en LLMs. Con la arquitectura ya preparada, se efectuaron las pruebas de validación descritas en la sección II, utilizando todos los ejemplos de prueba de ambos conjuntos de datos. Las tablas I y II muestran los resultados de estas pruebas, reportando para alucinaciones potenciales y longitud de respuesta tanto su media como su desviación estándar (entre paréntesis).

Tabla I
RESULTADOS DE LAS PRUEBAS PARA LOS DATOS SOBRE DIABETES.

	Alucinaciones potenciales	Longitud de respuesta	Precisión
Estrategia A	1.305 (0.87)	423.955 (98.52)	0.675
Estrategia B	1.693 (1.00)	1649.319 (533.49)	0.587
Estrategia C	1.773 (1.28)	736.820 (345.82)	0.641

Tabla II
RESULTADOS DE LAS PRUEBAS PARA LOS DATOS SOBRE CARDIOPATÍA.

	Alucinaciones potenciales	Longitud de respuesta	Precisión
Estrategia A	1.215 (0.87)	451.514 (101.01)	0.713
Estrategia B	1.905 (1.06)	1550.181 (535.64)	0.582
Estrategia C	1.738 (1.31)	660.219 (268.09)	0.626

Los resultados de las pruebas resultaron prometedores, puesto que la Estrategia A, correspondiente con nuestra propuesta, muestra claramente las mejores métricas, superando a las alternativas menos controladas en todos los apartados: las alucinaciones potenciales disminuyen (siendo esto nuestro objetivo principal), la longitud de las respuestas se reduce significativamente, especialmente comparada con la Estrategia B (LLM sin ningún tipo de información o control adicional), lo que puede contribuir a una mejor percepción de naturalidad; y la precisión predictiva es prácticamente igual a la que mostraban los sistemas de reglas borrosas subyacentes, indicando que el modelo no se está desviando de las conclusiones obtenidas por estos.

Una vez terminadas las pruebas de validación, se llevó a cabo una evaluación post-hoc adicional, con 500 casos seleccionados al azar entre los disponibles en los conjuntos de datos de prueba y usando solo la estrategia correspondiente a nuestra propuesta. Con esto se confirmó que la arquitectura implementada seguía adecuadamente las instrucciones sobre formato, al no registrar ningún error relacionado con el mismo.

REFERENCIAS

- [1] P. Perez-Ferreiro, A. Catala, A. Bugarin-Diz, and J. M. Alonso-Moral, “Generating trustworthy explanations with a language model enriched by fuzzy rule-based systems”, in proceedings of 2025 IEEE International Conference on Fuzzy Systems

Pooling global de características en Redes Neuronales Convolucionales mediante funciones de agregación basadas en desviaciones moderados

1° Iosu Rodriguez-Martinez

Dept. de Estadística, Informática y Matemáticas
Universidad Pública de Navarra
Pamplona, España
iosu.rodriguez@unavarra.es

3° Mikel Ferrero-Jaurieta

Dept. de Estadística, Informática y Matemáticas
Universidad Pública de Navarra
Pamplona, España
mikel.ferrero@unavarra.es

5° Francisco Javier Fernández

Dept. de Estadística, Informática y Matemáticas
Universidad Pública de Navarra
Pamplona, España
fcojavier.fernandez@unavarra.es

2° Jana Špírková

Dept. of Quantitative Methods and Information Systems
Matej Bel University
Banská Bystrica, Slovakia
jana.spirkova@umb.sk

4° Humberto Bustince

Dept. de Estadística, Informática y Matemáticas
Universidad Pública de Navarra
Pamplona, España
bustince@unavarra.es

Obtener el mejor representante para un conjunto de valores de entrada es uno de los problemas más recurrentes en el campo de la fusión de información. Las funciones basadas en *penalty* [1], [2] o las medias de Daróczy [3] intentan resolver este problema, aunque ambas presentan importantes limitaciones.

En particular, las medias de Daróczy no siempre conservan la monotonía y, por tanto, no son funciones de agregación en general. Un enfoque similar a la construcción de medias de Daróczy se puede llevar a cabo mediante funciones de desviación moderada [4], garantizando que la función resultante cumpla las condiciones de una función de agregación. Durante los últimos años se han propuesto diversas construcciones de agregaciones basadas en estas funciones [5], [6]. En particular, en [7] se abordó el problema de agregar valores en el intervalo $[-1, 1]$ mediante funciones de agregación basadas en desviaciones moderadas. Esta escala es particularmente interesante dado que permite representar el tipo de fenómenos simétricos que encontramos en el comportamiento humano, lo que la hace útil en problemas como la toma de decisión, donde el uso de escalas simétricas con respecto al 0 permite representar la simetría del comportamiento de la persona que

toma la decisión.

En este trabajo, presentamos una extensión de dicho enfoque que permite agregar valores en un rango general $[a, b] \subset \mathbb{R} \cup \{-\infty, \infty\}$, mediante un isomorfismo del intervalo $[a, b]$ a $[-1, 1]$.

Creemos que estas expresiones pueden resultar útiles como operadores de fusión en tareas de aprendizaje automático. Con esta idea en mente, y como ejemplo ilustrativo, comprobamos la efectividad de las agregaciones basadas en desviaciones moderadas para fusionar las características extraídas por una Red Neuronal Convolucional (*Convolutional Neural Network* o CNN).

Las CNN son una familia de redes neuronales que se han convertido en el estado del arte en el campo de la visión artificial para resolver problemas de clasificación [8], [9] o segmentación [10], [11] entre otros. Estos algoritmos extraen de manera automática las características más discriminantes de una imagen de entrada, que pueden ser posteriormente empleadas como la entrada para un modelo de clasificación o segmentación. A la hora de reducir la alta dimensionalidad de las representaciones generadas, es común emplear capas de Pooling Promedio Global (*Global Average Pooling*), que simplemente las reemplazan por la media aritmética de sus valores [12]. Dado que estos valores pueden acotarse en la práctica a un intervalo $[a, b] \subset \mathbb{R}$, exploramos empíricamente si nuestra propuesta puede actuar como un reemplazo valioso para este proceso, obteniendo resultados favorables.

El trabajo de Jana Špírková está financiado por la beca de la Slovak Scientific Grant Agency VEGA no. 1/0124/24. El trabajo de Iosu Rodriguez-Martinez, Mikel Ferrero-Jaurieta, Humberto Bustince y Javier Fernandez está financiado por el proyecto de investigación PID2022-136627NB-I00 financiado por MCIN/AEI/10.13039/501100011033/FEDER, UE.

REFERENCES

- [1] T. Calvo and G. Beliakov, “Aggregation functions based on penalties,” *Fuzzy sets and Systems*, vol. 161, no. 10, pp. 1420–1436, 2010.
- [2] H. Bustince, G. Beliakov, G. P. Dimuro, B. Bedregal, and R. Mesiar, “On the definition of penalty functions in data aggregation,” *Fuzzy Sets and Systems*, vol. 323, pp. 1–18, 2017.
- [3] Z. Daróczy, “Über eine klasse von mittelwerten,” *Publ. Math. Debrecen*, vol. 19, pp. 211–217, 1972.
- [4] M. Decký, R. Mesiar, and A. Stupňanová, “Deviation-based aggregation functions,” *Fuzzy Sets and systems*, vol. 332, pp. 29–36, 2018.
- [5] A. Stupňanová and P. Smrek, “Generalized deviation functions and construction of aggregation functions,” in *11th Conference of the European Society for Fuzzy Logic and Technology (EUSFLAT 2019)*, pp. 96–100, Atlantis Press, 2019.
- [6] J. Fumanal-Idocin, Z. Takáč, J. Fernández, J. A. Sanz, H. Goyena, C.-T. Lin, Y.-K. Wang, and H. Bustince, “Interval-valued aggregation functions based on moderate deviations applied to motor-imagery-based brain–computer interface,” *IEEE Transactions on Fuzzy Systems*, vol. 30, no. 7, pp. 2706–2720, 2021.
- [7] J. Špírková, H. Bustince, J. Fernandez, and M. Sesma-Sara, “New classes of the moderate deviation functions,” in *19th World Congress of the International Fuzzy Systems Association (IFSA), 12th Conference of the European Society for Fuzzy Logic and Technology (EUSFLAT), and 11th International Summer School on Aggregation Operators (AGOP)*, pp. 661–666, Atlantis Press, 2021.
- [8] K. He, X. Zhang, S. Ren, and J. Sun, “Deep residual learning for image recognition,” in *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- [9] G. Huang, Z. Liu, L. Van Der Maaten, and K. Q. Weinberger, “Densely connected convolutional networks,” in *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 4700–4708, 2017.
- [10] O. Ronneberger, P. Fischer, and T. Brox, “U-net: Convolutional networks for biomedical image segmentation,” in *International Conference on Medical image computing and computer-assisted intervention*, pp. 234–241, Springer, 2015.
- [11] V. Badrinarayanan, A. Kendall, and R. Cipolla, “Segnet: A deep convolutional encoder-decoder architecture for image segmentation,” *IEEE transactions on pattern analysis and machine intelligence*, vol. 39, no. 12, pp. 2481–2495, 2017.
- [12] M. Lin, Q. Chen, and S. Yan, “Network in network,” 2014.

Modelado de Perfiles Difusos para la Interpretación de Redes Neuronales en la Recomendación de Vinos

1st Arturo Fernández-Mora

Escuela de Doctorados

Universidad de Jaén

Jaén, España

afm00108@red.ujaen.es

2nd Juan Moreno-García

E. Ingeniería Industrial y Aeroespacial Toledo

Universidad de Castilla-La Mancha

City, Country

Juan.Moreno@uclm.es

3rd David Muñoz-Valero

E. Ingeniería Industrial y Aeroespacial Toledo

Universidad de Castilla-La Mancha

City, Country

David.Munoz@uclm.es

4th Luis Martínez

Escuela Politécnica Superior de Jaén

Universidad de Jaén

Jaén, España

martin@ujaen.es

Abstract—Este trabajo se centra en el modelado de distintos perfiles de usuario mediante técnicas difusas, con el objetivo de entrenar una red neuronal orientada a la recomendación de vinos. La finalidad principal es interpretar el funcionamiento del modelo a partir de los valores de activación de sus neuronas, analizando cómo se distribuyen los conceptos aprendidos a lo largo de las distintas unidades.

Para ello, se emplea un conjunto de datos que recoge características organolépticas del vino [1] —como el dulzor, el cuerpo o la temperatura óptima de consumo—, que se utilizan para definir perfiles difusos asociados a diferentes maridajes. Estos perfiles se modelan mediante reglas difusas generadas a partir de combinaciones de atributos y valores representativos [2]. Además, se asignan vectores de pesos a cada relación entre atributos, lo que permite modular la importancia relativa de cada uno o de sus combinaciones.

Utilizando el enfoque TSK FIS [3], se calcula un porcentaje de afinidad entre cada vino y cada perfil. Este valor se obtiene aplicando los pesos a las reglas difusas, de modo que la suma de los resultados aportados por cada regla equivale al 100%. Estos porcentajes de afinidad se emplean como salidas objetivo durante el entrenamiento de la red neuronal.

Una vez entrenado el modelo, se extraen los vectores de activación de las capas ocultas y se aplican algoritmos de agrupamiento con el fin de identificar patrones en los conceptos representados. Los resultados preliminares obtenidos son prometedores y respaldan la continuidad del estudio.

Index Terms—Fuzzy rules, TSK FIS, Explainability, Deep Learning

I. INTRODUCCIÓN

Este trabajo presenta un enfoque híbrido para la construcción de un sistema de recomendación de vinos, combinando técnicas de lógica difusa y redes neuronales. El objetivo principal es modelar perfiles de usuario mediante reglas difusas que reflejen preferencias en función de características organolépticas del vino, y utilizar dichos perfiles para guiar el entrenamiento de una red neuronal. Además, se persigue la

interpretabilidad del modelo a través del análisis de las activaciones neuronales, lo que permite estudiar la representación interna de los conceptos aprendidos. Este enfoque contribuye al desarrollo de sistemas de recomendación más transparentes y comprensibles, alineados con los principios de la inteligencia artificial explicable.

Este artículo se va a centrar en el proceso de modelización de los diferentes perfiles de maridaje. Para definir un perfil es necesario determinar varios factores cruciales.

- En primer lugar, se identifican los valores de los atributos que mejor caracterizan cada perfil. Por ejemplo, platos como la comida asiática, las ensaladas, el marisco o las frituras ligeras maridan adecuadamente con vinos muy secos de acidez elevada, bajo contenido tánico, cuerpo ligero y temperatura de servicio baja. Además, dado que estos alimentos suelen tener un bajo contenido graso —lo que implica una menor absorción del alcohol—, se recomienda acompañarlos con vinos de baja graduación alcohólica, como los vinos blancos o espumosos. El resumen del perfil se puede ver en la Tabla I, donde se muestran varios atributos extraídos del conjunto de datos que han sido usados para definir los perfiles. En este caso los atributos de cuerpo (CU), acidez (AC), taninos (TA) y dulzor (D) se encuentran en formato categórico con 5 niveles: [Muy Bajo, Bajo, Medio, Alto y Muy Alto] para AC y TA, [Muy Ligero, Ligero, Medio, Con Cuerpo y Con Mucho Cuerpo] para CU y [Muy seco, Seco, Semi-Seco, Dulce y Muy Dulce] para D; y C contiene cada uno de los tipos de vino como categorías (Tinto, Blanco, Rosado, Espumoso, etc). En cambio los valores de temperatura de servicio del vino (T) y del porcentaje de alcohol presente en el vino (AL) se encuentran en formato numérico, a partir del cual se han definido varios conjuntos difusos para cada atributo. Para T se han definido los conjuntos mostrados en la Figura 1 y para AL se han definido los mostrados en la Figura 2.

This work was supported by grant PID2020-112967GBC32 funded by MCIN/AEI/10.13039/501100011033, by ERDF A Way of Making Europe and by grant ProyExcel_00257.

- Posteriormente, es necesario definir un vector de preferencias que represente la importancia relativa de las distintas coocurrencias entre atributos. Este vector permite asignar pesos diferenciados a combinaciones de diferente número de atributos, permitiendo representar diferentes complejidades en la representación de los perfiles. Por otro lado, también se definen vectores de pesos para cada uno de los ordenes de las combinaciones específicas de características organolépticas del vino —como la relación entre acidez, temperatura de servicio y porcentaje de alcohol, o entre dulzor y taninos—, en función de su relevancia dentro de un perfil determinado. De este modo, se introduce un mecanismo de ponderación que permite reflejar con mayor precisión las preferencias implícitas en cada perfil de usuario, facilitando una modelización más flexible y expresiva mediante reglas difusas.

TABLE I
ATRIBUTOS PARA PLATOS DE COMIDA ASIÁTICA, ENSALADAS, FRITURA Y MARISCO

Cuerpo (CU)	Muy ligero Ligero
Temperatura (T)	Baja
Acidez (AC)	Muy alta Alta
Taninos (TA)	Muy Bajos
Alcohol (AL)	Bajo
Dulzor (D)	Muy seco
Color (C)	Blanco Espumoso

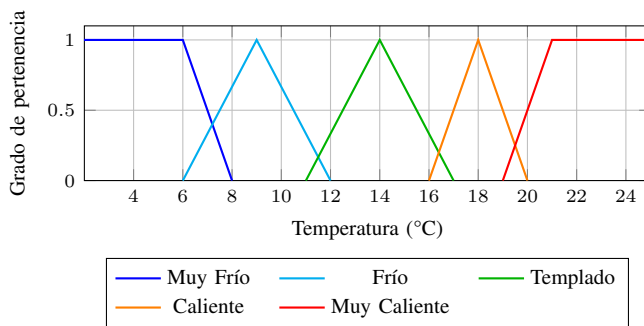


Fig. 1. Conjuntos difusos para la temperatura de servicio del vino

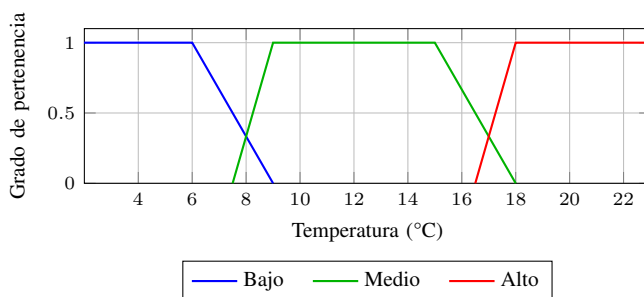


Fig. 2. Conjuntos difusos para el porcentaje de alcohol presente en el vino

A partir de estos datos, se generarán diferentes reglas difusas que capturan el conocimiento experto sobre maridajes, permitiendo evaluar el grado de afinidad entre cada vino y los distintos perfiles definidos. Estas reglas, construidas a partir de combinaciones relevantes de atributos facilitan la recomendación de maridajes adecuados para cada vino en función de sus características organolépticas.

II. CONCLUSIONES

Los perfiles generados mediante este método permiten representar de forma flexible y expresiva las preferencias enológicas asociadas a distintos contextos de consumo o tipos de comida. La utilización de lógica difusa en la construcción de estos perfiles ofrece una ventaja significativa al capturar la naturaleza gradual e incierta de muchas decisiones de maridaje, evitando clasificaciones rígidas. Además, al integrarse con el entrenamiento de una red neuronal, este enfoque no solo mejora la calidad de las recomendaciones, sino que también habilita un análisis interpretativo posterior mediante el estudio de las activaciones neuronales. En conjunto, los resultados obtenidos muestran el potencial de combinar técnicas de inteligencia artificial simbólica y conexionista para desarrollar sistemas de recomendación más transparentes, explicables y alineados con las preferencias reales de los usuarios.

REFERENCES

- [1] Kim, G.-r. *Wine information Version 7*. Kaggle. Retrieved May 28, 2025, from <https://www.kaggle.com/datasets/dev7halo/wine-information>
- [2] Muñoz-Valero, D., Moreno-García, J., López-Gómez, J. A., Villarrubia-Martín, E. A., & Jiménez-Linares, L. (2025). A knowledge-driven fuzzy logic framework for supporting decision-making entities. *Applied Soft Computing*, 139, 113415. <https://doi.org/10.1016/j.asoc.2025.113415>
- [3] Takagi, T., & Sugeno, M. (1985). Fuzzy identification of systems and its applications to modeling and control. *IEEE transactions on systems, man, and cybernetics*, (1), 116-132.

On the use of Aggregation Operators to improve Human Identification using Dental Records

Antonio D. Villegas-Yeguas
Dept. Computer Science and
Artificial Intelligence (DECSAI),
University of Granada, and
Panacea Coop. Research
Granada, Spain
Email: advy99@correo.ugr.es

Guillermo R-García
Panacea Coop. Research
Granada, Spain
Email: guiramgar96@gmail.com

Tzipi Kahana
Faculty of Criminology,
The Hebrew University of Jerusalem
Jerusalem, Israel
Email: kahana.tzipi@gmail.com

Oscar Ibáñez
Ciencias da Computación e Tecnoloxías da Información,
University of A Coruña, and Panacea Coop. Research
A Coruña, Spain
Email: oscar.ibanez@udc.es

Oscar Cordon
DECSAI and Andalusian Research Institute in
Data Science and Computational Intelligence
University of Granada. Granada, Spain
Email: ocordon@decsai.ugr.es

Abstract—The comparison of dental records is a standardized technique in forensic dentistry used to speed up the identification of individuals, which is based on the aggregation of different criteria. Existing automatic methods either make use of simple techniques, without utilizing the full potential of the information obtained from a comparison, or are proprietary and have not been peer-reviewed. This work aims to define high performance aggregation mechanisms that can be understood and validated by experts. We will use fuzzy aggregation operators and a modification of the state-of-the-art proposal by introducing expert knowledge. The results obtained in a database of 215 real-world cases show how the new approaches can improve the results of the state of the art without compromising the explainability and interpretability of the method.

Index Terms—Human identification, forensic anthropology, odontogram comparison, fuzzy aggregation operators.

I. INTRODUCTION AND BACKGROUND

There has been a recent increase in the number of geopolitical events that claim the lives of numerous victims. Although DNA and fingerprints help forensic experts for human identification in Disaster Victim Identification (DVI) scenarios with a high reliability, these techniques cannot be applied in situations with lack of ante-mortem (AM) or post-mortem (PM) data (e.g. lack of a second DNA sample for comparison or remain decomposition). For this reason, forensic odontology have been used widely in DVI scenarios [1].

A first step in the human identification process within forensic odontology is the dental record comparison [2]. Forensic experts record the condition of each tooth of an individual on a sheet (*odontogram*), describing different characteristics and pathologies. Given different odontogram pairs (one PM and several AM ones) to be compared, the number of coincidences (matches), of explainable differences (possible matches), and of unexplainable differences (mismatches) are obtained. These

characteristics are used to rank the odontograms pairs from most to least plausible using only the dentition description, allowing us to shortlist the candidates.

Different proposals have been made both to represent the dental information in these sheets [3], [4] and to automatically shortlist candidates [5], [6]. One of the main problems is the coding used to represent the dental information. A codification with a wide set of codes, as Interpol's proposal [4], can represent much more detailed dental information but is more complex and error-prone. Simple codification codes avoid this problem but the final record may not be sufficiently individualizing. Another common problem is that these automatic systems do not explain how they make a decision, which is problematic for experts to use. Another challenging problem is the lack of standardization, validation and access to the inner workings of the algorithms as most implementations are proprietary software or have not been peer reviewed.

The proposal of Adams and Aschheim [6] solves these problems by using a *simple coding system*, with only seven codes to represent the dental information and an automatic method considering a lexicographic order to rank pairs of odontograms in descending order of similarity. In this work we will use this codification to represent the dental records in our dataset and will extend the method in two different ways proposing a simple shortlisting method introducing expert knowledge and a fuzzy aggregation operator-based approach to combine the matching criteria.

II. PROPOSAL

A. Using expert knowledge to improve the method

With the first approach we aim to improve Adam's method by designing a different lexicographical order, ranking first by the lower percentage of mismatches followed by the higher

number of matches. The system will be able to avoid scenarios where a comparison with a high number of matches but a few mismatches would be better positioned than a comparison with no mismatches, many matches and a few possible matches. We will also consider a data-driven approach to explore all possible combinations of criteria used in the original proposal.

B. Exploring fuzzy aggregation operators

Two different types of aggregation operators, ordered-based aggregation operators [7], [8] and fuzzy integrals defined with respect a fuzzy measure aggregation operators [9], [10] will be considered. For the OWA and OWHM operator, we tried different sets of weights to test different scenarios:

- Maximum: For each comparison, the maximum value is chosen from all the criteria. In this case, the set of weights W will have $w_1 = 1$ and $w_i = 0$ for $i \neq 1$.
- Minimum: The minimum value is chosen from all criteria. This will choose the odontogram pair with the best minimum support of all its variables. The set of weights W will have $w_n = 1$ and $w_i = 0$ for $i \neq n$.
- Average: The final score is the average of all the criteria. The set of weights W will have $w_i = \frac{1}{n}$ for all $i \in [1, n]$.
- Linear: Since the proposed operators order the criteria by support, the characteristics with less support will have a lower weight in the final decision, and those with a strong support will have a higher decision power. The set of weights W will have $w_i = i$ for all $i \in [1, n]$.

For the Choquet and Sugeno integrals, we consider two different measures, the cardinality and a Sugeno λ -measure [11]. In the Sugeno λ -measure the values $\mu_\lambda^1, \dots, \mu_\lambda^n$ are called the fuzzy densities, which in our case can be interpreted as the importance of the a decision method considering the single criterion. To obtain a value for λ , we will use the normalized average ranking obtained by considering only the variable i to generate the rank as the value of μ_λ^i . In this way, we assign the relevance of each criteria based on its discrimination power.

III. EXPERIMENTS AND ANALYSIS OF RESULTS

After reviewing and processing more than 3400 orthopantomographies and 256 dental records and filtering the cases for which both AM and PM information were available, 215 cases were selected, providing a total of 430 dental charts from two different sources.

The best lexicographical order found from the expert approach is to order by i) % Mismatch, ii) % Match, iii) # Mismatch, iv) # Match, v) Exact group, vi) % Possible matches, vii) # Possible matches.

Table I shows how our lexicographical order-based aggregation obtains the best results on average and on maximum ranking, significantly outperforming Adam's method. We can also see how selecting the maximum or the minimum have very poor results, while averaging the criteria to obtain a score slightly improves Adams' proposal.

	Average	Q1	Q2	Q3	P95	P99	Max
Adams' method	10.51	1	1	2	77	172	206
Our lex. order	3.88	1	1	1	14	61	173
OWA-Min	14.82	1	1	3	102	178	204
OWA-Max	25.75	7	16	33	86	136	191
OWA-Average	7.62	1	1	1	41	120	186
OWA-Linear	5.57	1	1	2	27	84	184
OWHM-Average	18.44	3	5	14	88	175	202
OWHM-Linear	16.69	2	4	12	86	169	198
Choquet integral: μ : Cardinality	7.73	1	1	1	44	119	187
Choquet integral: μ : λ	7.52	1	1	1	45	123	186
Sugeno integral: μ : Cardinality	19.01	1	1	7	116	184	210
Sugeno integral: μ : λ	21.98	5	11	25	82	136	191

TABLE I: Performance statistics of every proposal.

The OWA operator with the linear weights obtains the second best result, both in average and in maximum ranking. In contrast, OWHM aggregators obtains worst results than the state-of-the-art. We attribute this result to the fact that the harmonic mean is a conservative aggregator, and is not the most appropriate approach for this problem.

Regarding the fuzzy integrals, Choquet improves the state-of-the-art results in both experiments, performing better with the Sugeno's λ -measure. This shows that including information on the performance of each criterion can improve the performance. However, Sugeno's integral obtains worse results than the state of the art with both measures.

For a more detailed comparison, Figure 1 show the CMC curve of the OWA and OWHM proposals, and Figure 2 the CMC curve of the Sugeno and the Choquet integral proposals. In both figures we added the results of the state-of-the-art method and the best method, our lexicographical order-based aggregation according to Table I to have a reference of the fuzzy aggregation operators' behavior.

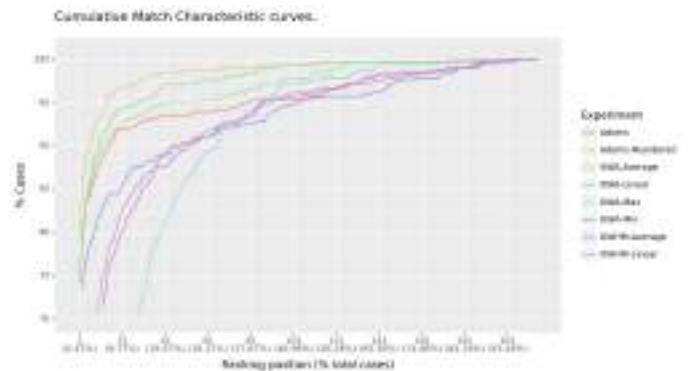


Fig. 1: Results of the fuzzy aggregation operators, our lexicographical order-based approach, and Adams proposal.

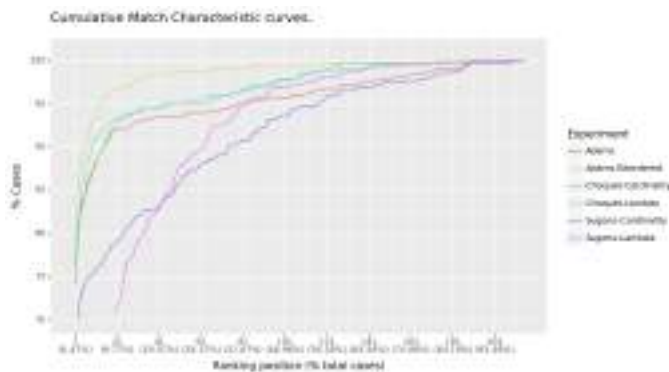


Fig. 2: Results of the fuzzy aggregation operators, our lexicographical order-based approach, and Adams proposal.

According to the results in Figures 1 and 2 and in Table I, it is clear that the best proposal is our expert-defined lexicographical order, followed by the OWA with linear weights. We think that this is because, although the aggregation of all the information improves Adam's method as in the case of using the average, the weights of the aggregation should not be uniform.

These experiments show how the application of different aggregation operators is useful to improve human identification methods in the field of forensic odontology. Further research is required to improve the performance of fuzzy aggregation operators.

ACKNOWLEDGMENT

This work was supported by grant CONFIA (PID2021-122916NB-I00) funded by MCIN/AEI/10.13039/501100011033, funded by 'ERDF A way of making Europe'.

This research project was made possible through the access granted by the Galician Supercomputing Center (CESGA) to its supercomputing infrastructure. The supercomputer Finis-Terrae III and its permanent data storage system have been funded by the NextGeneration EU 2021 Recovery, Transformation and Resilience Plan, ICT2021-006904, and also from the Pluri-regional Operational Programme of Spain 2014-2020 of the European Regional Development Fund (ERDF), ICTS-2019-02-CESGA-3, and from the State Programme for the Promotion of Scientific and Technical Research of Excellence of the State Plan for Scientific and Technical Research and Innovation 2013-2016 State subprogramme for scientific and technical infrastructures and equipment of ERDF, CEG15-DE-3114

REFERENCES

[1] H. P. Kirsch, K. Röttscher, C. Grundmann, and R. Lessig, "The tsunami disaster in the kingdom of thailand 2004," 2007.

[2] P. Pittayapat, R. Jacobs, E. De Valck, D. Vandermeulen, and G. Willems, "Forensic odontology in the disaster victim identification process," *Journal of Forensic Odonto-stomatology*, vol. 30, no. 1, pp. 1–12, Jul. 1, 2012.

[3] T. Chomdej, W. Pankaow, and S. Choychumroon, "Intelligent dental identification system (IDIS) in forensic medicine," *Forensic Science International*, vol. 158, no. 1, pp. 27–38, Apr. 20, 2006. DOI: 10.1016/j.forsciint.2005.05.001.

[4] A. Franco, S. G. F. Orestes, E. d. F. Coimbra, P. Thevissen, and Â. Fernandes, "Comparing dental identifier charting in cone beam computed tomography scans and panoramic radiographs using INTERPOL coding for human identification," *Forensic Science International*, vol. 302, p. 109860, Sep. 1, 2019. DOI: 10.1016/j.forsciint.2019.06.018.

[5] S. H. Al-Amad, J. G. Clement, M. McCullough, et al., "Evaluation of two dental identification computer systems: DAVID and WinID3.," *The Journal of forensic odonto-stomatology*, vol. 25, no. 1, pp. 23–29, Jun. 1, 2007, MAG ID: 2405115824 S2ID: 7889a79c73918ef6c8ae3f3862a5064db35b7632.

[6] B. J. Adams and K. W. Aschheim, "Computerized dental comparison: A critical review of dental coding and ranking algorithms used in victim identification," *Journal of Forensic Sciences*, vol. 61, no. 1, pp. 76–86, Jan. 1, 2016, MAG ID: 1943029493. DOI: 10.1111/1556-4029.12909.

[7] R. R. Yager, "On ordered weighted averaging aggregation operators in multicriteria decisionmaking," in *Readings in Fuzzy Sets for Intelligent Systems*, D. Dubois, H. Prade, and R. R. Yager, Eds., Morgan Kaufmann, Jan. 1, 1993, pp. 80–87. DOI: 10.1016/B978-1-4832-1450-4.50011-0.

[8] Z. Xu, "Fuzzy harmonic mean operators," *International Journal of Intelligent Systems*, vol. 24, no. 2, pp. 152–172, 2009, eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1002/int.20330>. DOI: 10.1002/int.20330.

[9] Sugeno and Michio, "Theory of fuzzy integrals and its applications," *doctoral thesis tokyo institute of technology*, 1974.

[10] G. Choquet, "Theory of capacities," *Annales de l'Institut Fourier*, vol. 5, pp. 131–295, 1954. DOI: 10.5802/aif.53.

[11] M. Sugeno, "Fuzzy measures and fuzzy integrals - a survey," in *Readings in Fuzzy Sets for Intelligent Systems*, D. Dubois, H. Prade, and R. R. Yager, Eds., Morgan Kaufmann, Jan. 1, 1993, pp. 251–257. DOI: 10.1016/B978-1-4832-1450-4.50027-4.

Caracterización del rendimiento de jugadores de balonmano mediante lógica difusa y algoritmos bioinspirados

Eusebio Angulo Sánchez-Herrera
Departamento de Matemáticas
Universidad de Castilla-La Mancha
Ciudad Real, España.
Eusebio.Angulo@uclm.es

Francisco P. Romero
Dep. Tecnolog. y Sist. de Información
Universidad de Castilla-La Mancha
Ciudad Real, España.
FranciscoP.Romero@uclm.es

Julio Alberto López-Gómez
Dep. Tecnolog. y Sist. de Información
Universidad de Castilla-La Mancha
Ciudad Real, España.
Julioalberto.lopez@uclm.es

Abstract—Este trabajo desarrolla una metodología basada en algoritmos de optimización bioinspirados y lógica difusa para poder caracterizar el rendimiento de jugadores profesionales de balonmano. En una primer etapa, los algoritmos bioinspirados se aplican para optimizar los pesos de las acciones más relevantes que definen el rendimiento de un jugador en partidos de alto rendimiento. Para mejorar la interpretabilidad de los pesos obtenidos y poder caracterizar a los jugadores, en una segunda etapa se aplicará lógica difusa para determinar los factores que caracterizan el rendimiento excelente de los jugadores de balonmano, dependiendo de la posición específica en la que jueguen. De esta forma, la metodología propuesta permitirá identificar las características más importantes que deben tener un jugador, en función de la posición específica que ocupe en la pista. Para llevar a cabo este estudio, se utilizarán datos reales provenientes de la federación internacional (IHF) y europea de balonmano (EHF) de campeonatos europeos y mundiales recientes en las categorías masculina y femenina. La finalidad de este trabajo es profundizar en el entendimiento de las variables que juegan un factor clave en el balonmano, un deporte que ha sido escasamente estudiado debido a su gran dinamismo y complejidad.

Index Terms—Sports Analytics, Handball Performance, Bio-inspired Algorithms, Fuzzy Sets

I. INTRODUCCIÓN

De forma general, el deporte de alto rendimiento tiene un notable impacto social y económico, atrayendo inversiones y empleo a través de grandes eventos como Juegos Olímpicos, Mundiales o Europeos. Desde el punto de vista particular, el balonmano no tiene la misma popularidad en España que el fútbol y el baloncesto, pese a ello es un deporte relevante especialmente dentro del continente europeo.

El principal objetivo y motivación de este trabajo es establecer una metodología que evalúa el rendimiento de jugadores profesionales de balonmano femenino y masculino tras su actuación en los partidos de competición. En primer lugar, se calculan los pesos de las distintas acciones del juego para los jugadores, utilizando algoritmos bioinspirados. Después, para interpretar los pesos de las acciones, se utilizan conjuntos difusos para definir variables lingüísticas que permiten caracterizar a los mejores jugadores en cada posición de juego en

balonmano (porteros, extremos, laterales, centrales y pivote); todo ello centrado en las siguientes aportaciones:

- Recopilación de datos reales de equipos, jugadores y eventos de balonmano de alto rendimiento.
- Aplicación de algoritmos bioinspirados para la optimización de los pesos de las distintas variables determinantes en los partidos de balonmano.
- Utilización de lógica difusa para la caracterización de los jugadores más determinantes en las distintas posiciones de balonmano.

El resto del trabajo se estructura de la siguiente manera: En la Sección II se estudian los trabajos existentes que abordan esta temática. Posteriormente, la Sección III describe los algoritmos y metodología de lógica difusa utilizada. Finalmente, la Sección IV señala las conclusiones de este trabajo y las líneas de investigación futuras.

II. ANTECEDENTES

Existen pocos trabajos que evalúen el rendimiento de los jugadores de balonmano en base estadísticas controlando todas las acciones de los mismos en los partidos de alto rendimiento. Por ejemplo en [1], los autores evalúan diferentes modelos de aprendizaje automático para predecir determinados tipos de rendimiento exclusivamente físico en jugadoras de balonmano, así como para determinar los factores significativos que influyen en el rendimiento esperado. Por otro lado, en [2] se presentó un enfoque novedoso basado en un modelo difuso de agregación de juicio de expertos con el fin seleccionar los criterios que mejor representan el rendimiento del jugador de balonmano en un partido y para establecer su relevancia en el juego. Posteriormente, en [3] se proporcionan los criterios necesarios para evaluar a porteros de balonmano aplicando metodologías de *soft-computing* para estimar sus pesos. Una evolución del trabajo anterior se presenta en [4] ampliando su alcance y aplicando algoritmos metaheurísticos para ajustar los pesos de las distintas acciones durante un partido. Por último, en [5] se aborda mediante algoritmos bioinspirados el ajuste de los pesos para jugadoras de todas las posiciones específicas de balonmano.

III. METODOLOGÍA

La metodología de partida de este trabajo se basa en la presentada en [5] con el fin de cuantificar los factores que caracterizan a los mejores jugadores de alto rendimiento y entender cómo los expertos evalúan y deciden la alineación ideal de la competición, es decir, el mejor jugador en cada posición de pista durante el torneo (Figura 1). A partir de este punto se incorpora, por un lado el abordaje de la categoría masculina y por otro, un paso adicional tras la obtención de los pesos de las distintas acciones del juego con el fin de mejorar la interpretabilidad de los resultados obtenidos y caracterizar el rendimiento de los jugadores mediante la aplicación de conjuntos difusos y variables lingüísticas.

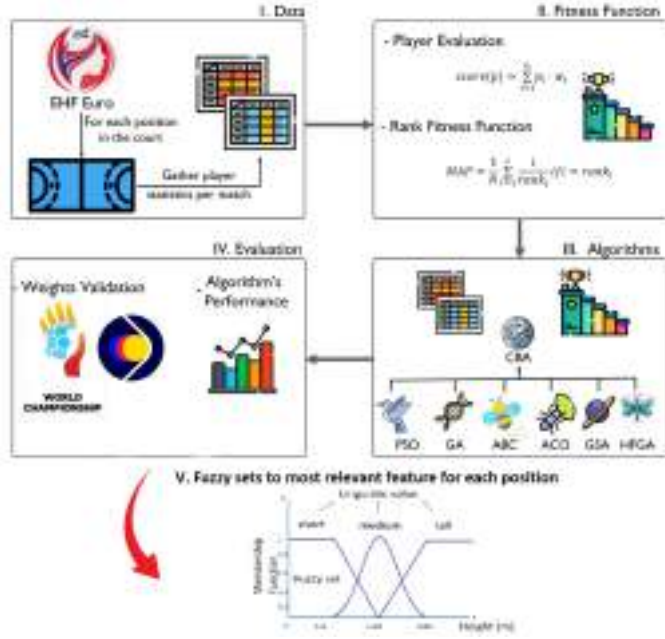


Fig. 1. Metodología basada en datos para identificar a los mejores jugadores y comprender el juicio de los expertos en las competiciones de balonmano.

A. Datos

Los datos corresponden a las estadísticas de acciones relevantes (definidas en [5] - Tabla 1) de los jugadores en partidos correspondientes a varias competiciones de alto nivel proporcionados por la *European Handball Federation* (EHF).

B. Formulación del problema

En primer lugar se calcula el rendimiento del jugador p en un torneo es la suma ponderada de sus acciones durante los partidos (véase la ecuación (1) propuesta por los autores [5]).

$$score(p) = \sum_{i=1}^n p_i \cdot w_i \quad (1)$$

Donde n es el número de acciones que se han recogido para evaluar a los jugadores, p_i es el valor recogido para i acción (agregado global de partidos jugados), y w_i es el peso de la acción. Ordenados de mayor a menor valor se define

un ranking de jugadores para cada posición. A partir de este resultado se define una función de *fitness* basada en *Mean Average Precision* (MAP) tal y como está definido en [5].

C. Algoritmos

Una vez determinados los datos y la función *fitness*, se aplicarán seis algoritmos bioinspirados: *Particle Swarm Optimization* (PSO), *Genetic Algorithm* (GA), *Artificial Bee Colony* (ABC), *Ant Colony Optimization* (ACO), *Gravitational Search Algorithm* (GSA) and *Hybrid Firefly - Genetic Algorithm* (HFCA).

D. Evaluación e interpretación

Con algoritmos de optimización se ajustan las ponderaciones óptimas, se resuelve 14 veces el problema de optimización con cada algoritmo, para cada posición específica (Hay 7 posiciones en balonmano) y hay que repetirlo para femenino y masculino. Posteriormente, se validarán utilizando datos reales de mundiales de balonmano femenino y masculino, comparando los resultados obtenidos con el *ranking* propuesto por los expertos. En este trabajo se incorpora el uso de conjuntos difusos y etiquetas lingüísticas con el fin de definir las variables más determinantes en cada posición específica del juego. Basándonos en el trabajo realizado por [3], cada criterio se evalúa mediante etiquetas lingüísticas pertenecientes a uno de los siguientes conjuntos de etiquetas predefinidos.

$$\begin{aligned} S^3 &= \{s_1^3 = H - High, s_2^3 = M - Medium, s_3^3 = L - Low\} \\ S^5 &= \{s_1^5 = VH - Very_High, s_2^5 = H - High, \\ &\quad s_3^5 = M - Medium, s_4^5 = L - Low, \\ &\quad s_5^5 = VL - Very_Low\} \end{aligned} \quad (2)$$

IV. CONCLUSIONES Y TRABAJOS FUTUROS

Este trabajo presenta una metodología que completa los trabajos existentes para evaluar jugadores de alto rendimiento, primero porque abordará el problema para todas las posiciones específicas en categoría femenina y masculina. A su vez, se introduce una capa de interpretabilidad con conjuntos difusos y variables lingüísticas para caracterizar el rendimiento de los jugadores.

REFERENCES

- [1] M. Oytun, C. Tinazci, B. Sekeroglu, C. Acikada and H.U. Yavuz, "Performance prediction and evaluation in female handball players using machine learning models," *IEEE Access* 8, pp. 116321–116335, 2020.
- [2] F.P. Romero, E. Angulo, J. Serrano-Guerrero and J.A. Olivas, "A Fuzzy Framework to Evaluate Players' Performance in Handball," *Int. Journal of Computational Intelligence Systems* 13 (1), pp. 549–558, 2020.
- [3] E. Angulo, F.P. Romero and J.A. López-Gómez, "A comparison of different soft-computing techniques for the evaluation of handball goalkeepers," *Soft Computing* 26 (6), pp. 3045–3058, 2022.
- [4] J.A. López-Gómez, F.P. Romero and E. Angulo, "A feature-weighting approach using metaheuristic algorithms to evaluate the performance of handball goalkeepers," *IEEE Access* 10, pp. 30556–30572, 2022.
- [5] J.A. López-Gómez, F.P. Romero and E. Angulo, "Bio-inspired algorithms for the characterization of excellent performance in handball players: a data-driven methodology," *Expert Systems with Applications* 274, pp. 126821, 2025.

Fuzzy processing applied to improve multimodal sensor data fusion to discover frequent behavioral patterns for smart healthcare

1st Carlos Fernández-Basso
Department of Computer Science
University of Granada
Granada, Spain
ORCID: 0000-0002-8809-8676

2nd David Díaz-Jiménez
Department of Computer Science
University of Jaén
Jaén, Spain
ORCID: 0000-0003-1791-4258

3rd José L. López Ruiz
Department of Computer Science
University of Jaén
Jaén, Spain
ORCID: 0000-0003-2583-8638

4th Macarena Espinilla
Department of Computer Science
University of Jaén
Jaén, Spain
ORCID: 0000-0003-1118-7782

Abstract—This work summarizes the proposal entitled “Fuzzy processing applied to improve multimodal sensor data fusion to discover frequent behavioral patterns for smart healthcare”, which introduces an automatic methodology for the fusion and fuzzification of multisource sensor data. This fuzzy logic-based method enhances the interpretability and interoperability of complex datasets and has been validated in sensorized homes of Type II diabetic patients located in Cabra (Spain). The proposed system enables the identification of behavioral patterns and the optimization of clinical monitoring, demonstrating high potential to improve healthcare delivery through advanced data fusion techniques.

Index Terms—Data fusion, Sensor data, Sensor fuzzification, Smart healthcare

I. INTRODUCTION

The use of sensor technologies in healthcare environments, such as hospitals, nursing homes, and domestic settings, generates large volumes of data with high potential to provide valuable insights into patients’ habits and behaviors. The increasing interest in exploiting such information has driven the development of advanced processing and analysis methods, particularly through data mining techniques aimed at detecting behavioral patterns that may support clinical decision-making.

In this context, systems based on the Internet of Medical Things (IoMT) enable continuous data acquisition through multimodal sensors. However, the heterogeneity and volume

of the data, together with their continuous and fine-grained nature, make direct analysis challenging. Therefore, it is necessary to implement data fusion and transformation processes that produce more compact and interpretable representations. In this regard, fuzzy logic constitutes an effective tool for modeling the uncertainty inherent in sensory data, facilitating its interpretation through linguistic labels.

The work entitled “Fuzzy processing applied to improve multimodal sensor data fusion to discover frequent behavioral patterns for smart healthcare” [1] presents an automatic methodology that combines multisensor data fusion with a fuzzification process that does not require expert knowledge. Its aim is to extract association rules that describe behavioral routines relevant to healthcare professionals. The proposal has been validated in a real-world scenario: sensorized homes of patients with type II diabetes, demonstrating its practical applicability in smart healthcare systems oriented toward the monitoring and improvement of daily habits.

The document is structured as follows. Section II presents the fuzzy logic-based rule extraction process. Section I introduces the use case in which the proposed methodology has been applied. Finally, Section IV outlines the conclusions and future work.

II. A FUZZY-BASED SYSTEM FOR PATTERN MINING IN BIG DATA

The proposal presented in this work is based on an intelligent system designed to detect behavioral patterns in older people from data collected by multimodal sensors deployed in real home environments. The system is structured into four main phases: data acquisition, pre-processing and fusion, fuzzification, and pattern mining through fuzzy association rules.

This result has been partially supported by grant PID2021-123960OB-I00, PDC2023-145863-I00 and grant PID2021-127275OB-I00 funded by MICIU/AEI/10.13039/501100011033 and by “ERDF A way of making Europe”, grant PDC2023-145863-I00 funded by MICIU/AEI/10.13039/501100011033 and by “European Union NextGenerationEU/PRTR”, and grant M.2 PDC_000756 funded by Consejería de Universidad, Investigación e Innovación and by ERDF Andalusia Program 2021-2027. Finally, the research reported in this paper is also funded by the European Union (BAG-INTEL project, grant agreement no. 101121309).

First, data acquisition is carried out using sensors that capture physical activity, indoor location, and environmental conditions, integrated within an edge-fog-cloud architecture and protected through secure communication protocols. In the pre-processing phase, data are temporally normalized into one-minute intervals, and outliers or irrelevant transactions are discarded. Subsequently, during the fusion phase, contextual information is enriched by associating sensors with the user's position in the home using spatial anchors and proximity-based membership functions. This information is transformed through a fuzzification process that generates automatic linguistic labels, enabling the intuitive interpretation of continuous data and facilitating further analysis through data mining algorithms. Finally, fuzzy association rules are applied within a Big Data environment using Apache Spark, allowing the extraction of frequent behavioral patterns.

Experimental results demonstrate the system's scalability, computational efficiency, and ability to produce interpretable outcomes for healthcare professionals. The combination of data fusion techniques with fuzzy logic not only enhances the quality of the dataset, but also enables the discovery of relevant routines, such as hygiene habits or treatment adherence, key elements for monitoring in smart healthcare environments.

III. USE CASE: SMART HOME

To validate the effectiveness of the system, it was integrated into a real home (see Figure 1) inhabited by an older person, where a sensor network was deployed to monitor the user's daily routines.

For this purpose, the following devices were installed across all key areas of the home, including the user through wearable technologies:

- Motion sensors
- Environmental sensors
- Open/close sensors
- Fixed location devices (Raspberry Pi)
- Activity wristband
- Central node (Raspberry Pi)

This infrastructure enables the detection of daily activities related to hygiene, meals, rest, medication intake, and user localization, even in multi-occupant settings, thanks to the integration of indoor positioning technologies based on Bluetooth Low Energy.

In this case, the application of the fuzzy association rule algorithm enabled the identification of meaningful relationships between behaviors and time slots. For example, routines such as tooth brushing followed by medication intake in the early morning were discovered, as well as shower usage prior to that same activity.

$$\{teeth = using, Medication = open\} \rightarrow \{time = earlyMorning\}$$

The proposed approach, based on contextual data fusion and the use of fuzzy linguistic labels, significantly reduced the generation of irrelevant rules compared to traditional methods,

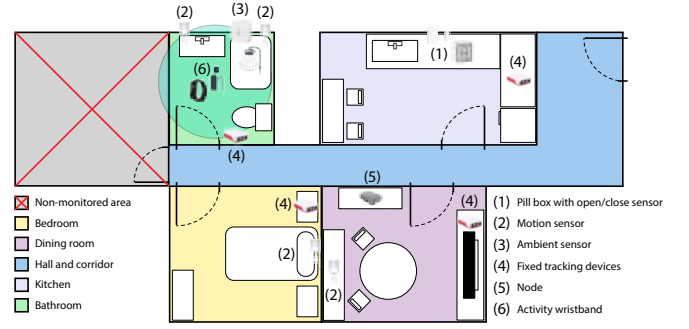


Fig. 1. Distribution of the dwelling used in the case study, location of the devices used and example of the activation range of the monitored elderly person obtained based on the proximity of these devices

thereby enhancing the interpretability and practical usefulness of the results for healthcare professionals.

Moreover, a substantial improvement in performance and scalability was observed when applied to scenarios involving large volumes of data. The reduction in dataset size, achieved by filtering out non-relevant information, combined with distributed processing in Big Data environments, enabled faster pattern extraction and optimized the use of computational resources.

IV. CONCLUSIONS AND FUTURE WORK

In this work, a data mining system was developed and applied to a real-world setting, specifically in a home located in Cabra, Spain, with the objective of discovering behavioral patterns through the integration of data fusion techniques and fuzzy logic. The proposed system demonstrated significant improvements in computational efficiency and interpretability, due to the reduction in data volume and the use of linguistic labels that facilitate the understanding of results by end users, such as medical professionals.

Furthermore, the system has proven to be scalable and adaptable to other contexts, including hospitals and nursing homes, as well as to different population profiles, such as patients with chronic diseases. The extraction of patterns, such as hygiene routines or treatment adherence, shows great potential for anticipating behavioral deviations and enabling early interventions. Future work will focus on increasing the number of sensors, improving result visualization, and strengthening security mechanisms through techniques such as federated learning, in addition to conducting a qualitative and quantitative evaluation of the system's impact in clinical practice.

REFERENCES

- [1] C. Fernandez-Basso, D. Díaz-Jimenez, J. L. López, and M. Espinilla, "Fuzzy processing applied to improve multimodal sensor data fusion to discover frequent behavioral patterns for smart healthcare," *Inf. Fusion*, vol. 123, p. 103307, Nov. 2025.

A discrete Median filter for RGB images

Pablo Flores-Vidal^{1,2}, Daniel Gómez^{1,2}, Luis Magdalena³

^{1,2}Dpto. de Estadística y Ciencia de Datos, Instituto Universitario de Estadística y Ciencia de Datos,
Facultad de Estudios Estadísticos, Universidad Complutense de Madrid, Spain

³ETSI Informáticos, Universidad Politécnica de Madrid, Madrid, Spain

¹pflores@ucm.es, ²dagomez@estad.ucm.es, ³luis.magdalena@upm.es

Abstract—This work proposes a new approach to smoothing RGB images. A new median filter is introduced considering smoothing filters as penalty functions that must respect two fundamental properties: discrete domain and monotonicity. For applying this to the RGB images, decomposable functions are defined. The effectiveness of this approach is illustrated by adapting two traditional methods based on the mean and median. Finally, the advantages of the proposed decomposable median method compared to traditional average approach are shown.

Index Terms—Image Pre-processing, Smoothing, RGB, Median filter, Decomposable functions

I. INTRODUCTION

In the digital image analysis field, pre-processing is a critical step to ensure the effectiveness of any downstream algorithms. One of the most important techniques in this step is smoothing [1]. The proper application of smoothing techniques is key to the success of subsequent tasks such as edge detection, region segmentation, etc.

Several well-known approaches and filters, such as mean, median, and Gaussian, are applied, each with different properties and applications [2]. However, most of these techniques were originally developed for single-channel images (grayscale), and their application to multichannel typically involves processing each channel independently, disregarding potential inter-channel dependencies. This limitation is particularly evident in the context of smoothing techniques, where the added complexity of preserving both spatial and chromatic consistency in color images remains an open and often overlooked challenge [3].

In this work, which is based in [4], we propose a reinterpretation of these approaches within the framework of separable or ‘decomposable’ monotonic functions (as we defined them). This allows studying their properties from a mathematical perspective. We explore how traditional smoothing methods can be reformulated under these decomposable functions and analyzing their behavior in terms of monotonicity, ensuring that if the aggregation is monotonic for each channel independently, the smoothing result will also preserve this property.

The presented approach provides a unified formulation for different smoothing methods and opens the possibility of designing new image processing techniques with well-defined properties in terms of aggregation and monotonicity.

Present research has been supported by Government of Spain (grants PID2021-122905NB-C21, PID2021-122905NB-C22 and PID2023-146540OB-C42).

II. SMOOTHING IN GRAY-SCALE IMAGES

Smoothing functions aim to reduce the noise and attenuate abrupt intensity variations, facilitating subsequent tasks such as segmentation, edge detection, among others [5]. Smoothing functions are usually defined on images with a single dimension (i.e. gray-scale images where $I_{i,j} \in \mathcal{T} = \{0, \dots, 255\}$) and can be seen as functions (convolutions) that, for each pixel $I_{i,j}$ (the ‘central pixel’), return a value $I_{i,j}^S$ resulting from the aggregation of neighbors’ image values. It is important to notice that pixel (i,j) belongs to its neighborhood: $(i,j) \in N(i,j)$.

Two well-known smoothing filters are the **mean filter** (which calculates the average of the values in $N(i,j)$) and the **Gaussian filter** (which calculates the weighted average of the values in $N(i,j)$ with the weights derived from a Gaussian distribution).

For a given pixel $I_{i,j}$, most of the smoothing functions are monotonic and can be defined as functions from $\mathcal{T}^{|N(i,j)|} \rightarrow \mathcal{T}$ as:

$$f_S(x_1, \dots, x_{|N(i,j)|}) = g \left(\sum_{i=1}^{|N(i,j)|} \lambda_i x_i \right),$$

where g is a function from $[0, 255] \rightarrow \{0, \dots, 255\}$ discretizing the continuous result of the average function into the discrete space $\mathcal{T} = \{0, \dots, 255\}$.

The discretization or truncation function to transform a real number into a $\{0, \dots, 255\}$ value can result in a loss of subtle intensity variations and may introduce artifacts in the smoothed image. Moreover, the truncation in the discretization process can blur edges or fail to preserve critical structural information, leading to loss of image detail.

Another issue arises when dealing with the median filter that is usually defined only for windows with odd number of pixels (The median of 9 values $x_1, \dots, x_9 \in \mathcal{T}$ can be always expressed as $x_{(5)}$, which belongs to $\mathcal{T}$.) But this fact could be changed when the cardinality of $N(i,j)$ is even.

Definition 1: Let us define the neighborhood of radius k for pixel (i,j) , as $N_k(i,j) = \{I_{r,s}, \text{ with } |r-i| \leq k \text{ and } |s-j| \leq k\}$, the set of neighbors at distance less or equal than k .

The median of a set of points in $\mathbb{R}$ can be viewed as the element that minimizes the sum of distances between that point and the set, and the mean minimizes the sum of the squared distances. These properties can be used to redefine smoothing functions.

Definition 2: Given a neighborhood of radius k for pixel (i, j) , $N_k(i, j) = \{x_1, \dots, x_{|N_k(i, j)|}\}$, the Median function $\phi^M : \mathcal{T}^{|N_k(i, j)|} \rightarrow \mathcal{T}$ can be defined as:

$$\phi^M(x_1, \dots, x_{|N_k(i, j)|}) = \arg \min_{a \in \mathcal{T}} \sum_{i=1}^{|N_k(i, j)|} |x_i - a|.$$

Definition 3: Given a neighborhood of radius k for pixel (i, j) , $N_k(i, j) = \{x_1, \dots, x_{|N_k(i, j)|}\}$, the Mean function $\mathcal{T}^{|N_k(i, j)|} \rightarrow [0, 255]$ can be defined as:

$$\text{Mean}(x_1, \dots, x_{|N_k(i, j)|}) = \arg \min_{a \in [0, 255]} \sum_{i=1}^{|N_k(i, j)|} (x_i - a)^2.$$

Let us observe that Definition 3 implies that Mean must be truncated to ensure that the output belongs to $\{0, \dots, 255\}$.

III. SMOOTHING IN RGB IMAGES

The most natural way to approach the problem of smoothing an RGB image is to decompose the partially ordered space $\mathcal{T} = \{0, \dots, 255\}^3$ into each of its components, so that all the measures defined in the gray-scale case can be extrapolated to the RGB or other multichannel contexts.

We will address the concept of decomposable as a partially ordered set $\mathcal{T}$ that is expressed by a Cartesian product, i.e. $\mathcal{T} = \mathcal{T}_1 \times \dots \times \mathcal{T}_k$ of linearly ordered spaces $\mathcal{T}_i$.

Definition 4: Given a partial order $\mathcal{T}$, we will say that it is decomposable when $\mathcal{T} = \mathcal{T}_1 \times \dots \times \mathcal{T}_k$, $\mathcal{T}_i$ being a linear order set.

From this definition, we can see that the RGB space of digital images $\mathcal{T} = \{0, \dots, 255\}^3$ is decomposable into three gray-scale spaces $\mathcal{T} = R \times G \times B$, where $R, G, B = \{0, \dots, 255\}$ are linear orders.

Let us note that here exists a natural way to build smoothing functions in decomposable sets from monotonic functions of their projections.

Given a decomposable set $\mathcal{T} = \mathcal{T}_1 \times \dots \times \mathcal{T}_k$ with the natural order obtained from the respective linear orders of $\mathcal{T}_i$, and given $f^i : \mathcal{T}_i^n \rightarrow \mathcal{T}_i$ a monotonic function for all $i = 1, \dots, k$, the following function f is a monotonic function from $\mathcal{T}^n$ to $\mathcal{T}$: $f = (f^1, \dots, f^k)$

Given a RGB image defined in $\mathcal{T} = \{0, \dots, 255\}^3$, the function $f^M = (\phi^M, \phi^M, \phi^M) : \mathcal{T}^n \rightarrow \mathcal{T}$ (**Decomposable Median**) is a monotonic function where,

$$\phi^M(x_1, \dots, x_{|N_k(i, j)|}) = \arg \min_{a \in \{0, \dots, 255\}} \sum_{i=1}^{|N_k(i, j)|} |x_i - a|.$$

IV. RESULTS AND DISCUSSION

The previously defined ϕ^{Median} smoothing filter was applied over the 50 first grayscale images of BSDS500 [5].

In order to evaluate the aggregation produced by the smoothing methods, they were tested with Sobel operator employing different strategies as it was made in [6].

The results show a superiority when employing our ‘decomposable’ method. To reflect this it is shown in Figure 1 a visual example of outputs of both smoothing approaches.

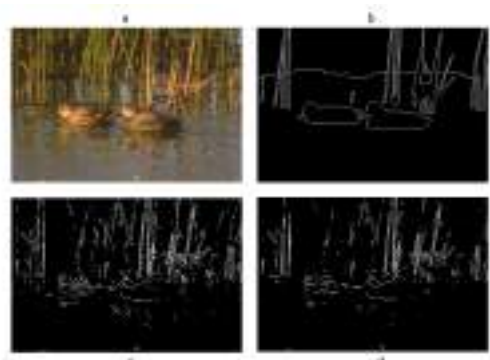


Fig. 1. An example of outputs of both smoothing approaches (average and median). a=original image, b=human ground-truth, c=Median method thr=180 and d=Average method thr=180.

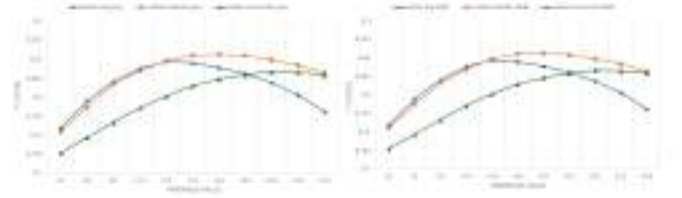


Fig. 2. Sobel performance of different methods with gray-scale and RGB images

The presented approach opens the possibility of designing new image processing techniques with well-defined properties in terms of aggregation and monotonicity. As demonstrated, our way of smoothing defined in non-linearly orderable spaces (because we are working in RGB), is more efficient than other classical methods, and furthermore, they are methods with great mathematical potential by respecting monotony in all directions. Median method’s was proved to be superior to Average method.

REFERENCES

- [1] R. C. González and R. E. Woods. *Digital Image Processing*, 3rd edition. 2008.
- [2] Anil K Jain. *Fundamentals of digital image processing*. Prentice-Hall, Inc., 1989.
- [3] Alasdair McAndrew. An introduction to digital image processing with matlab notes for scm2511 image processing. *School of Computer Science and Mathematics, Victoria Univ. of Technology*, 264(1):1–264, 2004.
- [4] P. Flores-Vidal, D. Gómez, and L. Magdalena. Smoothing in RGB images: A discrete 3-d median filter (submitted). In *Proceedings IEEE SMC 2025*. IEEE, 2025.
- [5] D. Martin, C. Fowlkes, D. Tal, and J. Malik. A database of human segmented natural images and its application to evaluating segmentation algorithms and measuring ecological statistics. In *Proceedings of IEEE International Conf. on Computer Vision*, pages 416–423, 2001.
- [6] Pablo Flores-Vidal, Daniel Gómez, Javier Castro, and Javier Montero. New aggregation approaches with hsv to color edge detection. *International Journal of Computational Intelligence Systems*, 15(1):78, 2022.

Modelado Interpretable de Datos de Glucosa

1° Juan F. Gaitán-Guerrero
Departamento de Informática
Universidad de Jaén
Jaén, España
ORCID: 0009-0007-6872-1401

2° José L. López
Departamento de Informática
Universidad de Jaén
Jaén, España
ORCID: 0000-0003-2583-8638

3° Macarena Espinilla Estévez
Departamento de Informática
Universidad de Jaén
Jaén, España
ORCID: 0000-0003-1118-7782

4° David Díaz-Jiménez
Departamento de Informática
Universidad de Jaén
Jaén, España
ORCID: 0000-0003-1791-4258

5° Carmen Martínez-Cruz
Departamento de Lenguajes y Sistemas Informáticos
Universidad de Granada
Granada, España
ORCID: 0000-0002-8117-0647

Abstract—La diabetes plantea un reto sanitario global que requiere interpretar adecuadamente los datos generados por sensores. Este trabajo propone un sistema basado en lógica difusa (FL) para generar resúmenes lingüísticos de datos de monitorización continua de glucosa (CGM). Integrado en una arquitectura del Internet de las Cosas (IoT) multipaciente, permite describir series temporales (TS) en lenguaje natural, facilitando su comprensión por profesionales, pacientes y familiares. La propuesta se compara con GPT-4, destacando ventajas en interpretabilidad, eficiencia y sostenibilidad.

Index Terms—Lógica difusa, Monitorización continua de glucosa, IoT, Resúmenes lingüísticos de series temporales, Generación de lenguaje natural, diabetes, GPT.

I. INTRODUCCIÓN

La diabetes *mellitus* es una enfermedad metabólica causada por una insuficiencia parcial o total de insulina, afectando a más de 537 millones de personas en el mundo [1]. Pese a que los sistemas sanitarios actuales dotan a sus pacientes con dispositivos inteligentes para la CGM a través de sensores, aun se enfrentan al desafío de ofrecer una atención personalizada, en gran parte por la complejidad y volumen de datos generados.

En este contexto, la FL surge como un enfoque clave para modelar la incertidumbre inherente a los datos de glucosa. Pese a que existen numerosos trabajos en el manejo de datos de diabetes, estos se enfocan en la predicción de valores futuros o en determinar si el paciente padece o no la enfermedad [2], no abordando así técnicas interpretables que permitan llevar a cabo un control adecuado de los pacientes.

Como consecuencia, se propone una arquitectura IoT multipaciente para la recogida de datos en tiempo real, combinada con un sistema basado en FL para generar descripciones lingüísticas interpretables de TS de glucosa. La validación del

sistema se realiza a través de una evaluación del mismo por parte de los usuarios finales, así como una comparativa con modelos de generación de lenguaje natural como GPT-4, en términos de interpretabilidad, eficiencia, contenido semántico y sostenibilidad.

II. TRABAJOS RELACIONADOS

En el ámbito de la CGM en pacientes con diabetes, numerosos trabajos proponen el uso de dispositivos IoT para recopilar datos biométricos. Sin embargo, muchas de estas propuestas no describen una arquitectura concreta multipaciente ni ofrecen información completa sobre los sensores empleados.

Por otra parte, para afrontar la incertidumbre inherente en los datos monitorizados, se recurre a la FL [3], la cual ha demostrado ser fundamental para generar descripciones lingüísticas de TS (GLiDTS), tal y como se refleja en [4]. Aunque los modelos de lenguaje como GPT-4 han comenzado a utilizarse para resumir datos de CGM [5], estos presentan errores, alucinaciones o subjetividad en descripciones lingüísticas que abarcan varios días, revelando así la necesidad de análisis diarios detallados.

III. ARQUITECTURA DEL SISTEMA

La arquitectura del sistema combina el sensor comercial Freestyle Libre 3 [6] con una aplicación móvil que recoge los datos en tiempo real; a diferencia del glucómetro tradicional, el dispositivo IoT permite una recopilación completa de los niveles de glucemia sin intervención del usuario, durante un período de 14 días. Para ello, se emplea la aplicación xDrip+ [7], que facilita la lectura continua del sensor y el envío automático de los datos al servidor. Los valores de glucosa, el timestamp y el utcOffset se almacenan y preprocesan en un servidor propio, que además aloja una plataforma web con el motor de generación de resúmenes lingüísticos. La comunicación se realiza mediante una API REST que replica la funcionalidad de Nightscout [8], y adaptada para gestionar múltiples pacientes de forma simultánea.

Este resultado ha sido parcialmente respaldado por el proyecto PID2021-127275OB-I00 y el proyecto PID2021-126363NB-I00 financiados por el MICIU/AEI/10.13039/501100011033 y por “ERDF A way of making Europe”, y por el proyecto PDC2023-145863-I00 financiado por el MICIU/AEI/ 10.13039/501100011033 y por la “European Union NextGenerationEU/PRTR”.

Este resultado es un resumen del trabajo publicado con referencia doi.org/10.1109/JIOT.2024.3419260

IV. METODOLOGÍA DE GENERACIÓN DE RESÚMENES LINGÜÍSTICOS

Las TS registradas en tiempo real son convertidas a lenguaje natural de manera interpretable, abordando así las limitaciones encontradas en la literatura.

En primer lugar, los datos con preprocesados, destacando el uso del algoritmo *Moving Average Window* para la imputación de valores perdidos en intervalos de tiempo no superiores a media hora. Posteriormente, se aplica el algoritmo *Ramer Douglas Peucker* [9] para la segmentación de la TS, obteniendo así dos TS (la original y la segmentada). Cada segmento se caracteriza usando funciones de pertenencia difusas asociadas a los niveles de glucosa, la tendencia de los segmentos y, los momentos del día al que pertenece cada uno de los puntos. Así, la TS segmentada permite la caracterización de eventos (i.e.: hipoglucemia, hiperglucemia, picos y mesetas). Por su parte, a partir de la TS original con el etiquetado de niveles de glucosa, los momentos del día y el uso de cuantificadores difusos, se permite la generación de Protoformas de tipo 2 (p.ej.: Por la mañana, la mayoría de los niveles de glucosa fueron altos).

Una vez se genera la descripción de la TS, cada una de las sentencias se evalúa en función del grado de verdad (DoT), el cual debe ser superior o igual a 0.7 para ser admitida en el resumen final. Adicionalmente, un marco de reglas de calidad es utilizado, evitando repeticiones y redundancias, dando así prioridad a eventos significativos de manera progresiva en el tiempo, proporcionando una descripción desde lo más general a lo más específico.

V. RESULTADOS Y DISCUSIÓN

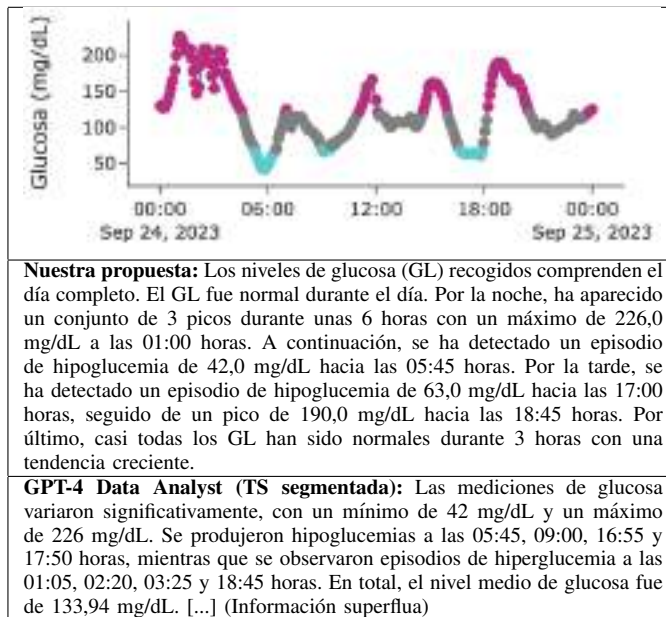
El sistema ha sido evaluado sobre un conjunto de datos reales de un paciente con diabetes tipo 1, con más de 200 días monitorizados. Para validar el conjunto de datos recogidos, este fue computado de forma tal que se hizo el cálculo de todas las posibles protoformas que se podrían generar. El resultado fue que el 89% de estas fueron activadas, lo que evidencia una amplia cobertura de los eventos esperados en los datos.

Una encuesta realizada a pacientes, familiares y personal clínico arrojó una puntuación media de usabilidad del 73%, destacando la claridad, brevedad y relevancia de los resúmenes. Comparado con GPT-4 y su *plug-in* Data Analyst (variando el *prompt* del modelo entre la TS original y la TS segmentada), nuestra propuesta se ensalzó en aspectos clave como la precisión, la dinámica temporal, la relevancia del contenido, la claridad narrativa, la escalabilidad del sistema y su sostenibilidad a largo plazo, esto último considerando los Objetivos de Desarrollo Sostenible #3 (Salud y Bienestar) #10 (Reducción de las desigualdades) y #12 (Producción y Consumo Responsables). Igualmente, cabe destacar como el *plug-in* Data Analyst obtuvo la mayor valoración en el escenario en que el *prompt* contenía la TS segmentada.

VI. CONCLUSIONES Y TRABAJOS FUTUROS

Este estudio ha presentado una arquitectura completa para la generación de resúmenes lingüísticos de alta calidad a partir

TABLE I
COMPARATIVA DE RESUMEN LINGÜÍSTICO CON GPT-4 DATA ANALYST.



de TS de glucosa, mediante un modelo de representación del conocimiento y una arquitectura de monitorización en tiempo real para múltiples pacientes, basada en sensores continuos de glucosa intersticial. De esta manera, se persigue fomentar una mejora en el bienestar de las personas a la vez que se alivia la carga de los sistemas sanitarios actuales.

El sistema permite visualizar e interpretar los datos de manera comprensible, implicando a expertos y usuarios en su diseño, aplicando FL para transformar patrones relevantes en descripciones en lenguaje natural. Como trabajo futuro, se plantea ampliar el análisis a periodos más largos, con el objetivo de evaluar la evolución del paciente y posibles cambios en sus hábitos alimentarios.

REFERENCES

- [1] International Diabetes Federation. (n.d.) What Is Diabetes. [Online]. Available: <https://idf.org/about-diabetes/what-is-diabetes/>
- [2] K. Liu, L. Li, Y. Ma, J. Jiang, Z. Liu, Z. Ye, S. Liu, C. Pu, C. Chen, Y. Wan et al., "Machine Learning Models for Blood Glucose Level Prediction in Patients With Diabetes Mellitus: Systematic Review and Network Meta-Analysis," *JMIR Medical Informatics*, vol. 11, no. 1, p. e47833, 2023.
- [3] L. A. Zadeh, "Fuzzy sets," *Information and control*, vol. 8, no. 3, pp. 338–353, 1965.
- [4] N. Marín and D. Sánchez, "On generating linguistic descriptions of time series," *Fuzzy Sets and Systems*, vol. 285, pp. 6–30, 2016.
- [5] E. Healey, A. Tan, K. Flint, J. Ruiz, and I. Kohane, "Leveraging large language models to analyze continuous glucose monitoring data: A case study," *medRxiv*, pp. 2024–04, 2024.
- [6] Abbot. (n.d.) Freestyle Libre 3 sensor. [Online]. Available: <https://www.freestyle.abbott/es-es/productos/freestylelibre-3.html>
- [7] Nightscout contributors, "Nightscout xDrip+." [Online]. Available: <https://github.com/NightscoutFoundation/xDrip>
- [8] Nightscout contributors, "Nightscout RESTful service framework." [Online]. Available: <https://github.com/cosmonaut/nightscout/blob/develop/swagger.yaml>
- [9] A. Saalfeld, "Topologically consistent line simplification with the douglas-peucker algorithm," *Cartography and Geographic Information Science*, vol. 26, no. 1, pp. 7–18, 1999.

Sistemas Difusos Lingüísticos que Evolucionan para Regresión

Javier Martín Moreno
Centro de Estudios Avanzados en
Física, Matemática y Computación
(CEAFMC)
Universidad de Huelva
Huelva, España
0000-0002-7002-1064

Francisco Alfredo Márquez Hernández
Centro de Estudios Avanzados en
Física, Matemática y Computación
(CEAFMC)
Universidad de Huelva
Huelva, España
0000-0003-4916-4703

Antonio Peregrín Rubio
Instituto Andaluz Interuniversitario
en Ciencia de Datos e
Inteligencia Computacional (DaSCI)
Universidad de Huelva
Huelva, España
0000-0001-7105-2615

Abstract—Los Sistemas Difusos mantienen su interés debido a su alto nivel de explicabilidad frente a otros modelos en el ámbito de la IA actual. En entornos dinámicos, los Sistemas Difusos que Evolucionan permiten modelar el conocimiento en tiempo real, adaptándose a los cambios en el entorno. No obstante, en tareas de regresión y control, estos modelos han sido abordados principalmente desde un punto de vista centrado en la precisión, usando Sistemas Difusos de tipo Takagi-Sugeno y Neuro-Difusos. En este trabajo se aborda el diseño de Sistemas Difusos Lingüísticos Adaptativos en Línea para Regresión y Control, primando por tanto su interpretabilidad frente a la precisión de las predicciones.

Keywords—Sistemas Difusos Lingüísticos, Sistemas Difusos Adaptativos en Línea, Sistemas Difusos Evolutivos, Aprendizaje Online, Inteligencia Artificial Explicable

I. INTRODUCCIÓN

Desde su concepción, los Sistemas Basados en Reglas Difusas (SBRD) han favorecido el manejo y modelado del conocimiento de manera similar a como lo hacen los humanos, lo que permite interpretar de manera sencilla sus criterios y razonamiento. Además, son ampliamente preferidos por su manejo de la incertidumbre y la imprecisión

Circunscribiéndonos a los empleados para regresión y control, según su arquitectura podemos distinguir entre diferentes paradigmas:

- *Mamdani* o Lingüísticos: Usan reglas con etiquetas lingüísticas tanto en los antecedentes como en los consecuentes de las reglas. A menudo, utilizan la misma representación de los conceptos lingüísticos para todas las reglas: de esta manera, se favorece en gran medida la interpretabilidad sacrificando precisión, al tener menor capacidad de adaptación a los datos [1].
- *Takagi-Sugeno-Kang* (TSK o TS): Usan etiquetas lingüísticas solo en los antecedentes de las reglas, representando habitualmente cada etiqueta un concepto distinto en cada regla. Los consecuentes están formados por polinomios cuyos coeficientes se ajustan en base a los datos de entrada [1].
- *Neuro*: Estos sistemas hibridan el manejo de la incertidumbre y la imprecisión de la Lógica Difusa con la arquitectura de conexión de las Redes Neuronales Artificiales [2].

Según la técnica utilizada para ajustar los parámetros del modelo, podemos distinguir entre:

- *Metaheurísticas Bio-inspiradas*: Se basan en diferentes metáforas biológicas para optimizar la función de proyección entre los datos de entrada y la salida del modelo. Cabe destacar dentro de este grupo a los Sistemas Difusos Evolutivos, *Evolutionary Fuzzy Systems*, llamados así por emplear algoritmos genéticos o evolutivos en general, ampliamente utilizados por sus buenos resultados [3].
- *Descenso por gradiente*: Estas técnicas guiadas por la derivada del error permiten una búsqueda de parámetros más exacta, aunque habitualmente con escasa exploración, viéndose perjudicados frecuentemente por los mínimos locales [4].

Estas técnicas de ajuste a menudo se ejecutan de manera *offline*, sobre un conjunto de datos establecido. De esta manera, el modelo obtenido permanece estático, independientemente de los posibles cambios en el entorno. En entornos dinámicos, donde el comportamiento de los datos puede variar con el paso del tiempo, es útil utilizar técnicas de ajuste en línea u *online* para evitar que el modelo quede obsoleto.

Los Sistemas Difusos que Evolucionan, *Evolving Fuzzy Systems*, se basan en técnicas de ajuste *online* para mantener los parámetros del modelo actualizados, manteniendo la precisión de las predicciones en tiempo real [5]. Desde que fueron propuestos por P. Angelov a principio de los 2000, han destacado por su ajuste incremental sobre datos generados en tiempo real en una sola pasada.

Sin embargo, la gran mayoría de aportaciones en esta área son de tipo TS o Neuro-Difusos. Consideramos que el uso de modelos difusos lingüísticos en entornos dinámicos es de interés en la actualidad, propiciando alternativas más interpretables.

La Inteligencia Artificial Explicable no sólo se basa en la interpretabilidad de los modelos, pues comprende varias características que permiten a los humanos confiar en ellos: Interpretabilidad, Explicabilidad, Comprensibilidad, Intelligibilidad y Comprensibilidad. En este sentido, no sólo la arquitectura del modelo resultante debe ser explicable: el método de ajuste debe ser también transparente, y por ello consideramos oportuno basar el diseño de este modelo en técnicas evolutivas.

En este trabajo se aborda preliminarmente el desarrollo de SBRDs Lingüísticos para regresión y control que Evolucionan, favoreciendo la interpretabilidad de los modelos mientras se mantienen ajustados en tiempo real.

II. SISTEMAS DIFUSOS LINGÜÍSTICOS QUE EVOLUCIONAN

Se propone un modelo basado en el uso de una arquitectura Lingüística, ajustada mediante algoritmos evolutivos, sobre un flujo de datos online. Para ello, se propone el flujo de trabajo representado en el Algoritmo 1.

Algoritmo 1 Pseudocódigo del modelo propuesto	
1:	Para cada instancia $\vec{x}$ do
2:	Preprocesar instancia $\vec{x}$
3:	Si BR es vacía Entonces
4:	Inicializar primera regla
5:	Si no
6:	Si instancia es no cubierta Entonces
7:	Generar Regla cubrimiento
8:	A. Evolutivo de los parámetros de la BC
9:	Si hay reglas similares Entonces
10:	Combinar Reglas
11:	Si hay reglas irrelevantes Entonces
12:	Eliminar Reglas

El modelo propuesto se compone principalmente de dos módulos:

- Aprendizaje de la Base de Conocimiento (BC): La Base de Datos (BD) se obtiene particionando el universo de discurso con etiquetas lingüísticas triangulares. La Base de Reglas (BR) se obtiene generando reglas por cubrimiento: se genera cada regla usando las etiquetas de los antecedentes y consecuentes con mayor grado de pertenencia para la instancia. Adicionalmente, se combinarán las reglas que superen el umbral de similitud y se podarán las que no superen el umbral de relevancia.
- Ajuste evolutivo de la BC: Los parámetros de la BC se ajustan mediante un algoritmo genético CHC online. Este ajuste se realiza de forma incremental, manteniendo el mejor cromosoma entre ejecuciones sucesivas. Para cada nueva instancia, se ejecuta el algoritmo genético.

Inicialmente, tanto la BD como la BR están vacías. La primera instancia inicializa la BD asignando a cada variable, tanto de entrada como de salida, una etiqueta lingüística cuyo valor coincida con el centro de la etiqueta. A partir de estas etiquetas se genera la primera regla.

A medida que el entorno evoluciona, cada vez que se recibe una nueva instancia, se actualiza primero la BD. Para ello, se modifican los valores máximos y mínimos del universo de discurso de cada variable, y se redistribuyen las etiquetas en consecuencia. Si la nueva instancia no está cubierta por ninguna de las reglas existentes en la BR, se genera una nueva regla que la cubra.

Posteriormente, se ejecuta el algoritmo genético para ajustar los parámetros de la BC teniendo en cuenta únicamente la instancia actual. Concretamente, el ajuste se realiza sobre los parámetros correspondientes a la posición y la anchura de

cada una de las etiquetas en la BD, y se ejecuta hasta cumplir un criterio de parada predefinido.

A continuación, se realiza una comprobación de similitud entre reglas, empleando criterios de distancia entre etiquetas. Las reglas que superan un umbral de similitud se combinan. Finalmente, se evalúa la relevancia de cada regla, calculada en función de su antigüedad y del número de instancias que cubre. Las reglas que superen un umbral de irrelevancia son eliminadas.

Este proceso de ajuste de parámetros mediante el algoritmo genético y de reestructuración de la BR, a través de la generación, combinación y eliminación de reglas, se realiza de forma continua, manteniendo la BC actualizada y garantizando la precisión en las predicciones.

III. ESTUDIO EXPERIMENTAL

Actualmente, el modelo propuesto se encuentra aún en fase de experimentación y se prevé, en breve, en la versión final de este documento incluir dichos resultados que consisten en mostrar y evaluar su rendimiento empleando conjuntos de datos representativos comparado con otros SBRDs, con el fin de analizar las características de los sistemas obtenidos y su nivel de desempeño. En cualquier caso, su fortaleza se encuentra en la comprensibilidad del modelo frente a otras opciones clásicas en Sistemas Difusos que Evolucionan, tipo TS, o a modelos de caja negra.

IV. CONCLUSIONES

En este trabajo se ha propuesto abordar el vano detectado entre las propuestas de la literatura comprendido por el desarrollo de Sistemas Difusos Lingüísticos para regresión en el ámbito de los Sistemas Difusos que Evolucionan. Se ha elaborado una propuesta sobre cómo diseñar estos modelos desde un enfoque que prima la explicabilidad del modelo y el mecanismo de ajuste frente a la precisión de las predicciones.

Si bien se trata de una propuesta preliminar, consideramos que es una vía de trabajo con proyección, dado que se trata de un ámbito de aplicación real, por ejemplo, integrado en dispositivos para control, que pudiesen requerir confianza y certificación.

REFERENCIAS

- [1] Alonso Moral, J. M., Castiello, C., Magdalena, L. & Mencar, C. (2021). An overview of fuzzy systems. *Explainable Fuzzy Systems: Paving the Way from Interpretable Fuzzy Systems to Explainable AI Systems*, 25-47. https://doi.org/10.1007/978-3-030-71098-9_2
- [2] de Campos Souza, P. V. (2020). Fuzzy neural networks and neuro-fuzzy networks: A review the main techniques and applications used in the literature. *Applied soft computing*, 92, 106275. <https://doi.org/10.1016/j.asoc.2020.106275>
- [3] Fernandez, A., Herrera, F., Cordon, O., del Jesus, M. J., & Marcelloni, F. (2019). Evolutionary fuzzy systems for explainable artificial intelligence: Why, when, what for, and where to?. *IEEE Computational intelligence magazine*, 14(1), 69-81. <https://doi.org/10.1109/MCI.2018.2881645>
- [4] Wu, D., Yuan, Y., Huang, J., & Tan, Y. (2019). Optimize TSK fuzzy systems for regression problems: Minibatch gradient descent with regularization, DropRule, and AdaBound (MBGD-RDA). *IEEE Transactions on Fuzzy Systems*, 28(5), 1003-1015. <https://doi.org/10.1109/TFUZZ.2019.2958559>
- [5] Angelov, P. (2023). Evolving fuzzy systems. In *Granular, fuzzy, and soft computing* (pp. 725-741). New York, NY: Springer US. https://doi.org/10.1007/978-1-0716-2628-3_192

Estudio comparativo de funciones de agregación inspiradas en Choquet para la reducción de dimensionalidad en imágenes con ruido

Joaquin Arellano, Humberto Bustince, Francisco Javier Fernández, Carlos Guerra, Cedric Marco-Detchart

Departamento de Estadística, Informática y Matemáticas

Universidad Pública de Navarra, Campus Arrosadía, 31006 Pamplona, España

arellanojoaquin@me.com, {bustince, fcojavier.fernandez, carlos.guerra, cedric.marco}@unavarra.es

Abstract—La reducción de dimensionalidad de imágenes con ruido es un reto clave en sistemas de visión con recursos limitados. Este trabajo propone una metodología basada en funciones de agregación inspiradas en la integral de Choquet, combinando operadores clásicos (media, mediana), la integral de Choquet y variantes de su generalización. En cada región de la imagen, el enfoque adapta de manera dinámica el tipo de operador y la medida difusa que fusiona las características, capturando relaciones entre píxeles y aumentando la robustez frente al ruido. Este esquema reduce la dimensionalidad preservando la información relevante. La novedad del estudio radica en integrar la flexibilidad de la integral de Choquet con una selección local inteligente de funciones de agregación. Se espera mayor calidad en la representación compacta y una mejor robustez ante ruido en comparación con métodos tradicionales.

Index Terms—Reducción de dimensionalidad, procesamiento de imágenes, funciones de agregación, integral de Choquet, operadores inspirados en d-Choquet

I. INTRODUCCIÓN

La reducción de dimensionalidad en procesamiento de imágenes juega un papel fundamental cuando el almacenamiento y la eficiencia son críticos, como en sistemas de visión en tiempo real o en dispositivos empujados. Técnicas clásicas como el análisis de componentes principales (PCA) permiten proyectar los datos en un espacio de menor dimensión preservando la mayor parte de la varianza [1], mientras que en el ámbito de redes neuronales convolucionales (CNN) se emplean capas de pooling para resumir regiones locales de la imagen [2]. Sin embargo, estos métodos suelen ser sensibles al ruido y pueden perder detalles estructurales importantes.

Para mejorar la robustez frente a diversos tipos de ruido, en la literatura se han explorado operadores de agregación basados en la integral de Choquet con medidas difusas [3] y, más recientemente, familias de operadores d-Choquet inspirados que combinan funciones de disimilitud y estrategias de fusión de manera flexible [4]. Adicionalmente, principios de desviación moderada ofrecen un marco para cuantificar similitud entre datos ruidosos [5]. Estos enfoques permiten capturar relaciones no lineales entre píxeles, pero habitualmente requieren diseñar una configuración fija de la medida difusa, lo que limita su adaptabilidad a regiones con características distintas.

II. METODOLOGÍA PROPUESTA

La estrategia desarrollada se basa en el procesamiento local de la imagen mediante una ventana deslizante de dimensiones 3×3 . Para cada bloque de nueve píxeles se generan múltiples estimaciones del valor representativo utilizando distintos operadores de agregación: funciones clásicas (media, mediana, mínimo y máximo), la integral de Choquet con medida difusa, variantes de operadores d-Choquet y un esquema basado en principios de desviación moderada. Cada operador aporta diferentes propiedades de conservación de información y eliminación de ruido.

Una vez procesadas todas las regiones, la imagen reducida se reconstruye mediante interpolación lineal a su resolución original. Esta etapa facilita la evaluación cuantitativa de la calidad de la representación utilizando métricas estándar, al comparar directamente con la referencia sin ruido.

El método propuesto presenta un diseño modular y adaptable: la incorporación de nuevos operadores o ajustes en el criterio de selección puede realizarse de forma transparente. Su capacidad de adaptación local lo convierten en una alternativa prometedora.

AGRADECIMIENTOS

Este trabajo ha sido financiado por el Ministerio de Ciencia e Innovación a través del proyecto PID2022-136627NB-I00, con el apoyo de MCIN/AEI/10.13039/501100011033/FEDER, Unión Europea.

REFERENCES

- [1] I. T. Jolliffe, *Principal component analysis for special types of data*. Springer, 2002.
- [2] D. Paternain, J. Fernández, H. Bustince, R. Mesiar, and G. Beliakov, "Construction of image reduction operators using averaging aggregation functions," *Fuzzy Sets and Systems*, vol. 261, pp. 87–111, 2015.
- [3] G. Beliakov, H. Bustince Sola, and T. Calvo, *A Practical Guide to Averaging Functions*, 1st ed., ser. Studies in Fuzziness and Soft Computing. Cham, Switzerland: Springer International Publishing, 2016.
- [4] H. Bustince, R. Mesiar, G. Dimuro, J. Fernández, J. Lafuente, and B. De Baets, "Choquet-inspired aggregation functions," SSRN, Tech. Rep. 5000794, 2025.
- [5] A. H. Altalhi, J. I. Forcén, M. Pagola, E. Barrenechea, H. Bustince, and Z. Takáč, "Moderate deviation and restricted equivalence functions for measuring similarity between data," *Information Sciences*, vol. 501, pp. 19–29, 2019.

Integración de datos clínicos y modelos de lenguaje para resúmenes lingüísticos de historias médicas

Andrea Morales-Garzón*, Raúl Martínez-Alonso*, Hossam El Amraoui-Leghzali*, Maria J. Martin-Bautista*, Carlos Fernandez-Basso†

*Departamento de Ciencias de la Computación e Inteligencia Artificial, Universidad de Granada, Granada, España
amoralesg@decsai.ugr.es, {raulmartinez03, hossam}@correo.ugr.es, mbautis@decsai.ugr.es

†Departamento de Lenguajes y Sistemas Informáticos, Universidad de Granada, Granada, España
cjferba@ugr.es

Abstract—El auge de datasets a gran escala en el ámbito médico ha impulsado la generación automática de texto clínico. Sin embargo, la descentralización y fragmentación de los historiales clínicos dificultan la creación de resúmenes contextualizados y completos sobre los pacientes. En este trabajo, proponemos una metodología para la unificación de historias clínicas a nivel de paciente centrada en el dataset MIMIC-IV y recursos complementarios como PhysioNet. Nuestra estrategia permite consolidar múltiples fuentes de datos hospitalarios en una única representación estructurada, que luego es empleada para generar prompts personalizados según la especialidad médica. Estos prompts son utilizados por modelos de lenguaje de gran escala (LLMs) para producir resúmenes clínicos claros y eficientes. Los resultados respaldan la viabilidad de nuestra propuesta para mejorar la comprensión integral del paciente y favorecer aplicaciones clínicas de apoyo a la decisión.

Index Terms—large language models, healthcare, text summarisation, data fusion

I. INTRODUCCIÓN

La publicación de datasets a gran escala en el ámbito médico ha marcado un antes y un después en la investigación médica computacional [1]. Han permitido avances significativos en la comprensión, modelado y análisis de datos clínicos en tareas muy diversas, como puede ser la generación de texto clínico [2]. Entre estos recursos, destaca el *Medical Information Mart for Intensive Care* (MIMIC) [1], un dataset de acceso abierto que proporciona información detallada y variada sobre pacientes atendidos en unidades de cuidados intensivos. La publicación de este dataset es especialmente relevante por su riqueza multimodal (incluye tanto datos estructurados como notas clínicas, señales fisiológicas o imágenes) y por haber sido adoptado como estándar de referencia en numerosos trabajos [3], [4]. Sin embargo, su potencial se puede ver limitado por la descentralización en la organización y publicación de los datos disponibles. Un caso de ello es su aplicación a entornos lingüísticos. A pesar de que existen conjuntos complementarios (p.ej., PhysioNet [5]) no hay un único punto

centralizado que integre todos estos recursos, lo que dificulta la construcción de soluciones generalizables a partir de múltiples fuentes de datos. Como consecuencia, los historiales clínicos están desagregados, dificultando el desarrollo de aplicaciones contextuales que pongan en contexto todo el conocimiento existente sobre el paciente, y como consecuencia, la calidad y confianza en los modelos de IA desarrollados para ello.

Estudios recientes en este campo han explorado herramientas basadas en IA generativa para la generación automática de informes clínicos [6]. Destaca el uso de large language models (LLMs) para resumir registros electrónicos de salud (EHR), con el fin de reducir la carga de trabajo clínico, generando borradores de respuestas para asistir al profesional médico [3]. Asimismo, MIMIC se ha usado en tareas relacionadas, por ejemplo, para detectar el consumo de sustancias (p.ej., alcohol) por parte del paciente a partir de la información almacenada en los informes clínicos [7].

Este trabajo busca una estrategia global para unificar historias clínicas con el objetivo de obtener resúmenes clínicos con LLMs, para una comprensión rápida y eficiente del paciente.

II. MÉTODO PROPUESTO

La Figura 1 ilustra la metodología propuesta para resumir el historial del paciente, basada en los siguientes pasos:



Fig. 1. Metodología propuesta para la generación de resúmenes clínicos.

a) *Recopilación de datos*: Para la generación de informes clínicos, hemos empleado el conjunto de datos MIMIC-IV, el cual recopila información detallada sobre pacientes ingresados en unidades de cuidados intensivos, incluyendo historiales clínicos, tratamientos administrados y otros datos hospitalarios relevantes. Lo hemos combinado con PhysioNet (dataset con notas clínicas de los pacientes en MIMIC-IV).

Research funded by the project BAG-INTEL (Ref. 101121309) funded by the European Commission, Grant PID2021-123960OB-I00 funded by MCIN/AEI/10.13039/501100011033 and by ERDF/EU and Grant TED2021-129402B-C21 funded by MCIN/AEI/10.13039/501100011033 and by the European Union NextGenerationEU/PRTR. Also funded by “Consejería de Transformación Económica, Industria, Conocimiento y Universidades de la Junta de Andalucía” through a pre-doctoral fellowship program (Grant Ref. PREDOC_00298).

b) *Unificación de datos a nivel de paciente*: Selección, extracción y unificación de múltiples tablas contenidas en el conjunto de datos. Estas tablas incluyen información general del paciente, servicios hospitalarios utilizados y diagnósticos registrados, entre otros. La integración de los conjuntos de datos se realizó a través del identificador único del paciente, lo que permite reconstruir su historial. Las principales etapas fueron: (1) Análisis detallado de los datos y sus relaciones; (2) Depuración de datos inconsistentes o incompletos; (3) Integración en una única estructura.

c) *Creación de prompts*: Generación automática de la entrada al LLM (prompt) a partir del historial clínico consolidado, adaptadas a cada especialidad médica. El considerar la especialidad médica en la gestión del paciente nos permite personalizar el prompt para adecuarlo a distintas patologías.

d) *Instrucción al LLM para generación de resúmenes*: Resúmenes clínicos con los datos unificados, utilizando un LLM. El proceso se ha personalizado a nivel de dominio utilizando roles. Evaluamos el resumen con métricas que nos permiten cuantificar la polaridad, subjetividad, y claridad del prompt, así como medidas semánticas de dominio médico.

e) *Generación del resumen clínico*: Obtenemos un resumen breve del paciente combinando distintas fuentes, para facilitar una comprensión rápida de su historial médico.

III. EXPERIMENTOS Y RESULTADOS

Se unificaron los datos de pacientes de los conjuntos en MIMIC-IV (hospitalización y UCI). La fusión se realizó en distintos batches debido a limitaciones de memoria. Una vez fusionados a nivel de paciente, se realizó una segunda fusión con las notas clínicas de PhysioNet, ocupando un total de 98 GB. El conjunto final fusionado contiene 736,152,332 admisiones de pacientes en distintos servicios hospitalarios. Dado que las notas clínicas están en inglés, el procesamiento textual y la elección de modelos de lenguaje se realizó teniendo en cuenta el rendimiento de los modelos para este idioma.

Para la generación de resúmenes, usamos 150 historiales de pacientes, obtenidos aleatoriamente de la categoría ‘cardiothoracic’. Esta categoría tiene diagnósticos de distinta gravedad y casuísticas. Como LLMs, utilizamos **Gemma-2-2b-it** [8], y diseñado para ofrecer un equilibrio entre rendimiento y eficiencia computacional) y **Llama 3.2-1B-Instruct** [9], afinado para seguir instrucciones y generar informes claros y concisos.

Para evaluar la calidad del resumen, usamos las métricas de comprensión lingüística *Flesch Reading Ease*, *SMOG index* *Dale-Chall index* y la métrica de similitud *Bert-score*. Esta última la ejecutamos con BERT (propósito general) y Bio_ClinicalBERT, entrenado en textos clínicos. La Figura 2 muestra una comparativa de la legibilidad de los resúmenes obtenidos con Gemma y LLaMA, bajo dos condiciones: generación con rol (contextualizada en un escenario médico) y sin rol. Se evalúan tres métricas clásicas de legibilidad (Flesch Reading Ease, SMOG, y Dale-Chall) a distintas temperaturas de muestreo. Los resultados indican que la inferencia con LLaMA usando rol genera textos más fáciles de leer, con puntuaciones más altas en Flesch y más bajas en SMOG y

Dale-Chall, especialmente a valores de temperaturas intermedios y bajos (es decir, en escenarios más controlados). En comparación, Gemma muestra una legibilidad más estable pero menor, y el impacto del rol en este modelo es menos pronunciado. Los resultados sugieren que en lugar de dificultar la comprensión del resumen, el uso de roles mejora la claridad y accesibilidad de los resúmenes, siendo más efectivo en el caso de LLaMA. BertScore nos permitió comprobar que LLaMA genera textos con mayor diversidad de contenido cuando se usan roles, en comparación con cuando no.

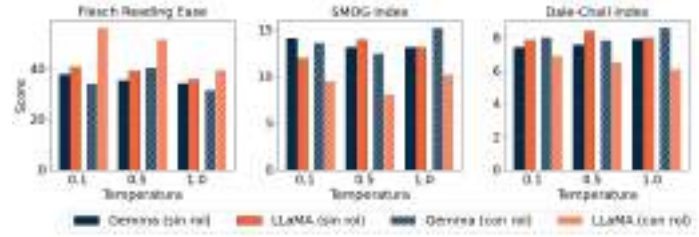


Fig. 2. Comparativa de las medidas de legibilidad de los resúmenes.

IV. CONCLUSIONES Y LÍNEAS FUTURAS

En este estudio, hemos podido comprobar que el uso de roles de especialidad clínica permiten mejorar la calidad de los resúmenes lingüísticos. En trabajos futuros, evaluaremos si a partir del resumen generado, un modelo automático puede predecir con éxito eventos clínicos relevantes, con el objetivo de valorar la utilidad clínica del resumen. Asimismo, exploraremos métodos para la generación de resúmenes controlables, donde podamos definir el foco del contenido, así como mecanismos de *style transfer* en LLMs para adaptar el estilo y el nivel técnico del resumen según el perfil del lector.

REFERENCES

- [1] C. Meng, L. Trinh, N. Xu, J. Enouen, and Y. Liu, “Interpretability and fairness evaluation of deep learning models on mimic-iv dataset,” *Scientific Reports*, vol. 12, no. 1, p. 7166, 2022.
- [2] A. Latif and J. Kim, “Evaluation and analysis of large language models for clinical text augmentation and generation,” *IEEE Access*, 2024.
- [3] S. L. Baxter, C. A. Longhurst, M. Millen, A. M. Sitapati, and M. Tai-Seale, “Generative artificial intelligence responses to patient messages in the electronic health record: early lessons learned,” *JAMIA open*, vol. 7, no. 2, p. ooae028, 2024.
- [4] I. Aden, C. H. Child, and C. C. Reyes-Aldasoro, “International classification of diseases prediction from mimic-iii clinical text using pre-trained clinicalbert and nlp deep learning models achieving state of the art,” *Big Data and Cognitive Computing*, vol. 8, no. 5, p. 47, 2024.
- [5] A. Johnson, L. Bulgarelli, T. Pollard, L. A. Celi, R. Mark, and S. Horng, “MIMIC-IV-ED (version 2.2),” 2023. [Online]. Available: <https://physionet.org/content/mimic-iv-ed/2.2/>
- [6] R. Luo, L. Sun, Y. Xia, T. Qin, S. Zhang, H. Poon, and T.-Y. Liu, “Biogpt: generative pre-trained transformer for biomedical text generation and mining,” *Briefings in bioinformatics*, vol. 23, no. 6, p. bbac409, 2022.
- [7] F. Shah-Mohammadi and J. Finkelstein, “Extraction of substance use information from clinical notes: Generative pretrained transformer-based investigation,” *JMIR Medical Informatics*, vol. 12, p. e56243, 2024.
- [8] G. Team, M. Riviere, S. Pathak, P. G. Sessa, C. Hardin, S. Bhupatiraju, L. Husseint, T. Mesnard, B. Shahriari, A. Ramé et al., “Gemma 2: Improving open language models at a practical size,” *arXiv preprint arXiv:2408.00118*, 2024.
- [9] A. Grattafiori, A. Dubey, A. Jauhri, A. Pandey, A. Kadian, A. Al-Dahle, A. Letman, A. Mathur, A. Schelten, A. Vaughan et al., “The llama 3 herd of models,” *arXiv preprint arXiv:2407.21783*, 2024.

APIA: Sistema Borroso para la estimación del riesgo de incumplimiento en la ejecución de contratos públicos.

1st Andrés Montoro-Montarroso

Tecnologías y Sistemas de Información
Universidad de Castilla-La Mancha
Ciudad Real, España
andres.montoro@uclm.es

2nd Guadalupe Plaza

Tecnologías y Sistemas de Información
Universidad de Castilla-La Mancha
Ciudad Real, España
guadalupe.plaza@uclm.es

3rd Jose A. Olivas

Tecnologías y Sistemas de Información
Universidad de Castilla-La Mancha
Ciudad Real, España
joseangel.olivas@uclm.es

4th Jesus Serrano-Guerrero

Tecnologías y Sistemas de Información
Universidad de Castilla-La Mancha
Ciudad Real, España
jesus.serrano@uclm.es

5th Jared D.T. Guerrero-Sosa

Tecnologías y Sistemas de Información
Universidad de Castilla-La Mancha
Ciudad Real, España
jareddavid.guerrero@uclm.es

6th Francisco P. Romero

Tecnologías y Sistemas de Información
Universidad de Castilla-La Mancha
Ciudad Real, España
franciscop.romero@uclm.es

Palabras Clave—Administración pública, contratos, recuperación de información, lógica borrosa

I. INTRODUCCIÓN

En el ámbito de la Administración Pública, la adopción de sistemas de Inteligencia Artificial (IA) promete mejorar la eficiencia de la gestión, optimizar la toma de decisiones y potenciar la calidad de los servicios [1], [2]. Sin embargo, su adopción exige un enfoque responsable que priorice la gobernanza del dato, la ética y la transparencia, evitando que los algoritmos deterioren la confianza ciudadana [3].

En este contexto bajo el marco del proyecto “Administración Pública e Inteligencia Artificial: regulación y uso de la IA en el ámbito de la contratación Pública” (APIA) se presenta un sistema de IA borroso con un enfoque centrado en las personas. Su objetivo es doble: detectar de forma temprana posibles irregularidades y riesgos de incumplimiento durante la ejecución de contratos públicos y, asistir en la redacción de Pliegos de Cláusulas Administrativas Particulares (PCAP) para que las condiciones y requisitos se formulen adecuadamente, minimizando así riesgos futuros [4].

II. PROPUESTA

La incorporación de IA en la contratación pública responde a la necesidad de modernizar la Administración y, al mismo tiempo, reforzar los mecanismos de control para prevenir riesgos en el uso de recursos públicos. El sistema borroso propuesto (ver Figura 1) pretende anticipar incumplimientos y apoyar la toma de decisiones, siempre dentro de un enfoque

centrado en las personas y orientado al interés público, en este caso el de la Universidad de Castilla-La Mancha (UCLM).



Fig. 1. Metodología del sistema para la prevención del riesgo de incumplimiento de contratos públicos.

A. Modelo conceptual de partida

Reunidos con las partes interesadas (servicio de gerencia de la UCLM) se establece la siguiente hipótesis: la mayoría de los incumplimientos de los contratos públicos licitados se producen en su fase de ejecución por la falta de control. Dado que en esta fase no existe material para alimentar al sistema de IA se acordó anticiparse al potencial incumplimiento en la fase de preparación del contrato por parte de los funcionarios. En concreto, el sistema propuesto toma como entrada los PCAP.

B. Modelo operativo

Durante el proceso de adquisición de conocimiento se ha contado con diversos actores y organizaciones expertas en derecho y, específicamente, en contratación pública. Algunas de estas entidades son: el Centro de Estudios Europeos “Luis

Este trabajo ha sido financiado por MCIN/AEI /10.13039/501100011033 y por la Unión Europea NextGenerationEU/PRTR, proyecto “Administración Pública e Inteligencia Artificial: regulación y uso de la IA en el ámbito de la contratación pública” (TED2021-130682B-I00).

Ortega Álvarez¹, el Instituto de Derecho Penal Europeo e Internacional², el equipo de Gerencia de la UCLM y expertos en contratación de la Junta de Comunidades de Castilla-La Mancha (JCCM). Otras fuentes de conocimiento fueron las siguientes:

- Encuesta anónima³ enviada a más de 20.000 entidades públicas para que aportaran expedientes incumplidos.
- Decisiones judiciales relativas a incumplimientos contractuales extraídas de la base de datos Aranzadi.
- Diversas fuentes bibliográficas e informes de juntas de contratación entre 1990 y 2024.

C. Organización del conocimiento

Una vez obtenido el conocimiento se ha organizado en forma de taxonomía:

- Fase de preparación (ej.: plazos reducidos que limitan la concurrencia).
- Fase de licitación (ej.: falta de verificación de la solvencia).
- Fase de ejecución (ej.: modificaciones injustificadas del contrato).
- Otros riesgos horizontales (ej.: insuficiente planificación de la actividad contractual, conflicto de interés, colusión).

La taxonomía recoge todos los riesgos que se pueden dar según la fase del procedimiento de contratación.

D. Selección y refinamiento del conocimiento

Como se ha mencionado en el apartado II-A se han seleccionado aquellos factores que se pueden identificar en los PCAP, por ejemplo: inexistencia de régimen específico y concreto de penalidades, falta de personalización del responsable del contrato, entre otros. Una vez identificados los factores y con ayuda de un comité de expertos en contratación de la UCLM y de la JCCM se ha definido la semántica de la taxonomía a partir de una base de reglas borrosas y funciones de pertenencia asociadas a cada concepto (ver Figura 2).

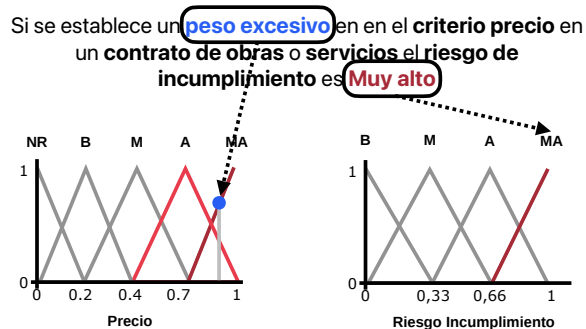


Fig. 2. Ejemplo de base de conocimiento del sistema propuesto.

Para la obtención de los PCAPs se ha implementado un sistema de extracción sobre la plataforma de contratación del

sector público⁴ (PLACSP). Dado que las relaciones entre los órganos de contratación, las licitaciones y los documentos están organizados de forma compleja, estos se han almacenado en una base de datos Neo4j (ver Figura 3).

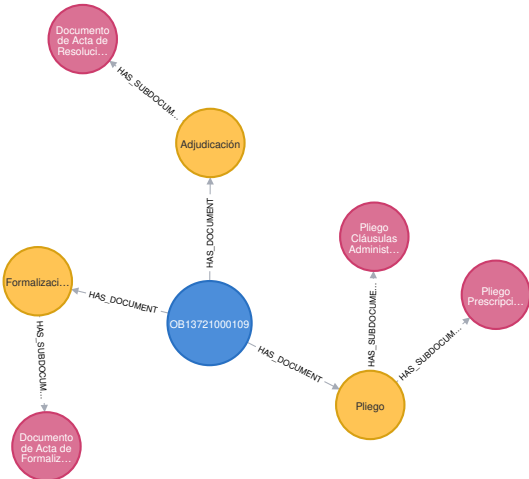


Fig. 3. Ejemplo de grafo sobre la PLACSP.

E. Razonamiento

Obtenido y representado el conocimiento el proceso de razonamiento se divide en dos partes bien diferenciadas:

- 1) Recuperación de información. Dado que la representación del conocimiento parte de unas sentencias imprecisas se ha empleado Llama-3.3-70B-Instruct⁵ como modelo de recuperación de información [5] que extrae de los PCAPs, los factores de riesgo identificados.
- 2) Inferencia borrosa. Una vez identificados los potenciales riesgos en un PCAP, estos se combinan mediante inferencia borrosa y agregación borrosa ponderada para obtener el nivel de riesgo de incumplimiento del contrato. Esto permite al funcionario identificar los riesgos potenciales y la posterior mejorar en la redacción de los PCAPs para minimizarlos.

BIBLIOGRAFÍA

- [1] E. M. Mota Sánchez, A. Montoro-Montarros, A. Nieto Martín, and J. A. Olivas, "Inteligencia Artificial y el control interno en el sector público local," Red Localis, Tech. Rep. 16/2021, 2021. [Online]. Available: <https://bit.ly/3FPaa4s>
- [2] B. W. Wirtz, J. C. Weyerer, and C. Geyer, "Artificial Intelligence and the Public Sector—Applications and Challenges," *International Journal of Public Administration*, vol. 42, no. 7, pp. 596–615, 2019.
- [3] Y.-F. Wang, Y.-C. Chen, S.-Y. Chien, and P.-J. Wang, "Citizens' trust in AI-enabled government systems," *Information Polity*, vol. 29, no. 3, pp. 293–312, 2024.
- [4] A. Montoro-Montarros, "Propuesta de metodología para la implantación de sistemas de Inteligencia Artificial para la prevención del fraude en la contratación pública," in *Inteligencia artificial y administraciones públicas: una triple visión en clave comparada*. Iustel Portal Derecho, 2024, pp. 401–410.
- [5] Y. Zhu, H. Yuan, S. Wang, J. Liu, W. Liu, C. Deng, H. Chen, Z. Liu, Z. Dou, and J.-R. Wen, "Large Language Models for Information Retrieval: A Survey," 2023.

¹<https://www.uclm.es/centros-investigacion/cee>

²<https://www.uclm.es/centros-investigacion/idp>

³<https://forms.office.com/e/GiayrFq085>

⁴<https://contrataciondelestado.es/>

⁵<https://huggingface.co/meta-llama/Llama-3.3-70B-Instruct>

Consistencia semántica y patrones de estilo en el tiempo para la detección de bots en redes sociales

Salvador Lopez-Joya*, Jose A. Diaz-Garcia*, M. Dolores Ruiz*, Maria J. Martin-Bautista*,

*Departamento de Ciencias de la Computación e Inteligencia Artificial, Universidad de Granada, Granada, España
slopezjoya@ugr.es, {jagarcia, mdruiz, mbautis}@decsai.ugr.es

Abstract—En los últimos años, las redes sociales se han integrado profundamente en la interacción humana, haciendo de su seguridad e integridad un objetivo clave. Un problema creciente es la proliferación de cuentas automatizadas (bots) con fines maliciosos. Este trabajo propone un nuevo enfoque para detectarlas mediante el análisis de consistencia semántica y patrones estilísticos en el tiempo. Nuestra hipótesis postula que, mientras los usuarios humanos muestran variaciones naturales tanto en su estilo de escritura como en los temas que discuten a lo largo del tiempo, los bots suelen exhibir patrones más uniformes y consistentes. Nuestra propuesta introduce un conjunto de métricas que capturan tanto características estilísticas y de legibilidad del texto, como tres medidas diseñadas específicamente para evaluar la coherencia semántica en las publicaciones de los usuarios. Los resultados experimentales indican que esta estrategia combinada es prometedora, ofreciendo un nuevo enfoque para la detección de bots basado en diferencias conductuales fundamentales.

Index Terms—bot detection, natural language processing, social networks, data mining

I. INTRODUCTION

La rápida expansión de las redes sociales ha revolucionado la comunicación, pero también ha facilitado la propagación de desinformación mediante bots sociales (cuentas automatizadas que simulan comportamiento humano), amenazando la integridad del ecosistema digital [1]. Estos bots son utilizados para interferir en procesos políticos, difundir noticias falsas y realizar actividades maliciosas como spam o ciberataques [2].

Aunque la Inteligencia Artificial ha surgido como una herramienta prometedora para su detección [3], los enfoques actuales aún presentan limitaciones. La literatura se centra principalmente en dos líneas: métricas basadas en grafos y análisis de contenido/perfil de usuario [4]. En este trabajo, nos enfocamos en este último enfoque, proponiendo un método que combina consistencia semántica y patrones estilísticos.

Los trabajos que tratan la consistencia de estilo están clásicamente relacionados con la detección de autoría, pero no es la primera vez que se usan en detección de bots. En [5] se emplean las medias y desviaciones típicas de una lista variada de características sobre el dataset Cresci-17 [6]. Este último estudio, a diferencia del aquí presentado, no contempla la consistencia semántica, pero demuestra la eficacia de las

características basadas en la consistencia estilística aplicada a la detección de bots.

II. MÉTODO PROPUESTO

A. Métricas para consistencia estilística y de legibilidad

Para cuantificar la consistencia en el estilo de escritura y patrones de legibilidad de cada usuario, hemos calculado estadísticos descriptivos básicos para 27 características lingüísticas clásicas en estilometría y legibilidad. Formalmente, dado un conjunto de n tweets publicados por un usuario y una característica específica f , calculamos:

$$\mu_f = \frac{1}{n} \sum_{i=1}^n f(t_i) \quad (1)$$

$$\sigma_f = \sqrt{\frac{1}{n} \sum_{i=1}^n (f(t_i) - \mu_f)^2} \quad (2)$$

donde:

- μ_f representa la media del valor de la característica f a través de todos los tweets $(t_1, t_2, \dots, t_n)$ del usuario
- σ_f corresponde a la desviación estándar de dicha característica
- $f(t_i)$ es el valor de la característica f calculada para el tweet t_i

El estilo promedio revela preferencias lingüísticas constantes, mientras que la variabilidad mide su inconsistencia. Para legibilidad, la media indica complejidad habitual y la desviación, fluctuaciones entre tweets simples y complejos.

En la Figura 1 se pueden ver el top 20 características derivadas de las originales, que abarcan desde características superficiales (como el número de hashtags) hasta medidas sofisticadas de legibilidad calculadas mediante algoritmos estandarizados en procesamiento de lenguaje natural.

B. Métricas para consistencia semántica

Para el cálculo de las métricas basadas en la consistencia semántica se obtienen los embeddings de los tweets por cada usuario. El modelo usado para este proceso es all-MiniLM-L6-v2 que está especializado en tareas de clustering y búsqueda semántica. Sea entonces $\mathcal{E} = \{e_1, e_2, \dots, e_n\}$ un conjunto de embeddings de dimensión d . Se proponen dos métricas diferentes:

The research reported in this paper was funded by the European Union (BAG-INTEL project, grant agreement no. 101121309), and FederaMed project: Grant PID2021-123960OB-I00 funded by MICIU/AEI/10.13039/501100011033 and by ERDF/EU. Finally, the research reported in this paper is also supported by the DesinfoScan project: Grant TED2021-129402B-C21 funded by MICIU/AEI/10.13039/501100011033 and by the European Union NextGenerationEU/PRTR.

a) *Coherencia semántica*: En primer lugar se calcula la similitud del coseno entre todos los pares posibles de embeddings. Entonces, la coherencia semántica media y su desviación estándar se definen como:

$$\mu_{\text{sim}} = \frac{2}{n(n-1)} \sum_{i=1}^{n-1} \sum_{j=i+1}^n \text{sim}(\mathbf{e}_i, \mathbf{e}_j) \quad (3)$$

$$\sigma_{\text{sim}} = \sqrt{\frac{2}{n(n-1)} \sum_{i=1}^{n-1} \sum_{j=i+1}^n (\text{sim}(\mathbf{e}_i, \mathbf{e}_j) - \mu_{\text{sim}})^2} \quad (4)$$

b) *Concentración semántica*: Primero, se calcula el centroide semántico de los embeddings:

$$\mathbf{c} = \frac{1}{n} \sum_{i=1}^n \mathbf{e}_i \quad (5)$$

Luego, la concentración semántica se define como la media de las distancias de coseno entre cada embedding y el centroide:

$$cs = \frac{1}{n} \sum_{i=1}^n \text{sim}(\mathbf{e}_i, \mathbf{c}) \quad (6)$$

III. EXPERIMENTOS Y RESULTADOS

Se han realizado una serie de experimentos sobre el dataset TwiBot-20 [7]. Este dataset, extraído de la plataforma X, está enfocado en la detección de bots y está compuesto por cuentas variadas etiquetadas de forma binaria. Se han seleccionado de forma aleatoria un subconjunto de 2500 cuentas cuya cantidad de posts superase los 100. Se ha entrenado un Random Forest usando una validación cruzada de 5 partes. Los resultados se muestran en la Tabla I.

TABLE I
RESULTADOS EN EL CONJUNTO TWIBOT-20

ACC	Prec.	Recall	F1-Macro	F1-Micro	AUC
0.7065	0.7090	0.7065	0.7057	0.7065	0.8088

Además, se han extraído los valores SHAP para identificar la contribución de cada variable a la predicción final. Tal y como se puede ver en la Figura 1, entre las 6 características que más han contribuido a la predicción de la clase encontramos dos medias, dos desviaciones típicas y dos de las tres características de consistencia semántica propuestas. De hecho, la variable más influyente ha sido la desviación típica de la longitud de los tweets que ha contribuido a que la predicción se decante por la clase positiva (bot) cuando los valores son bajos. Es decir, cuanto menos dispersión hay en la longitud de los tweets de un usuario más probabilidad hay de que el modelo se decante por considerarlo un bot.

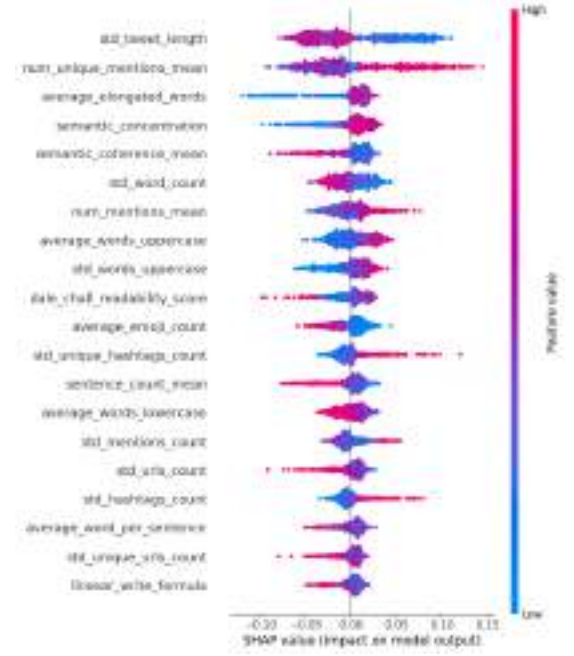


Fig. 1. Valores SHAP de las características usadas.

IV. CONCLUSIONES Y TRABAJOS FUTUROS

En este trabajo se ha mostrado como la combinación de características basadas en la consistencia del estilo de escritura y la consistencia semántica tienen una capacidad predictiva notable a la hora de identificar cuentas automáticas en redes sociales. En trabajos futuros se pretende estudiar más profundamente las distintas alternativas basadas en embeddings para la consistencia semántica con el objetivo de desarrollar métodos todavía más adecuados para modelar esta característica con la intención de que sirva como apoyo para futuros trabajos en detección de bots.

REFERENCES

- [1] A. M. Almars, E.-S. Atlam, T. H. Noor, G. Elmarhomy, R. Alagamy, I. Gad, Users opinion and emotion understanding in social media regarding covid-19 vaccine, *Computing* 104 (6) (2022) 1481–1496.
- [2] D. L. Linvill, P. L. Warren, Troll factories: Manufacturing specialized disinformation on twitter, *Political Communication* 37 (4) (2020) 447–467.
- [3] N. Hajli, U. Saeed, M. Tajvidi, F. Shirazi, Social bots and the spread of disinformation in social media: the challenges of artificial intelligence, *British Journal of Management* 33 (3) (2022) 1238–1253.
- [4] E. Ferrara, Social bot detection in the age of chatgpt: Challenges and opportunities, *First Monday* (2023).
- [5] M. Cardaioli, M. Conti, A. D. Sorbo, E. Fabrizio, S. Laudanna, C. A. Visaggio, It's a matter of style: Detecting social bots through writing style consistency, in: *2021 International Conference on Computer Communications and Networks (ICCCN)*, 2021, pp. 1–9. doi:10.1109/ICCCN52240.2021.9522339.
- [6] S. Cresci, R. Di Pietro, M. Petrocchi, A. Spognardi, M. Tesconi, Social fingerprinting: detection of spambot groups through dna-inspired behavioral modeling, *IEEE Transactions on Dependable and Secure Computing* 15 (4) (2017) 561–576.
- [7] S. Feng, H. Wan, N. Wang, J. Li, M. Luo, Twibot-20: A comprehensive twitter bot detection benchmark, in: *Proceedings of the 30th ACM International Conference on Information & Knowledge Management*, 2021, pp. 4485–4494.

Metodología híbrida para el aprendizaje y la evaluación de modelos de lenguaje bajo esquemas de incertidumbre difusa

1st Juan F. Gaitán-Guerrero 2nd Carmén Martínez-Cruz 3rd Macarena Espinilla 4th David Díaz-Jiménez
Departamento de Informática Departamento de Informática Departamento de Informática Departamento de Informática
Universidad de Jaén Universidad de Granada Universidad de Jaén Universidad de Jaén
Jaén, España Granada, España Jaén, España Jaén, España
jgaitan@ujaen.es cmcruz@ugr.es mestevez@ujaen.es ddjimene@ujaen.es

5th José L. López
Departamento de Informática
Universidad de Jaén
Jaén, España
llopez@ujaen.es

Abstract—Este artículo propone una metodología híbrida para el entrenamiento y evaluación de modelos de generación de lenguaje natural en dominios con presencia de datos temporales cuyo procesamiento y descripción están caracterizados por incertidumbre y ambigüedad. El enfoque central se basa en la lógica difusa como herramienta transversal: en la generación y validación del conjunto de datos, en la guía del modelo mediante ingeniería de prompts, y en la evaluación del rendimiento a través de una metodología de evaluación objetiva. La propuesta se valida en el dominio de la monitorización continua de glucosa en pacientes con diabetes, demostrando mejoras significativas en precisión, claridad y seguridad de las salidas generadas por modelos de la familia GPT.

Index Terms—Lógica difusa, GPT, Series temporales, Protoformas, Ingeniería de prompts, IoT, Fine-tuning, Resúmenes lingüísticos, Metodología de evaluación.

I. INTRODUCCIÓN

La interpretación automática de datos provenientes de sensores en tiempo real representa uno de los grandes retos en la inteligencia artificial aplicada. En dominios críticos como la salud, donde los datos suelen ser complejos, ruidosos e inciertos, resulta esencial contar con metodologías robustas que permitan a los modelos generar interpretaciones confiables y comprensibles para los usuarios finales.

En la actual era de la Inteligencia Artificial, muchas personas ya usan modelos de lenguaje. Esto coincide con las limitaciones del sistema sanitario, donde las citas son breves

y falta una visión clara de la evolución de pacientes crónicos monitorizados. Ante ello, tecnologías accesibles como las de OpenAI —con más de 400 millones de usuarios— abren la posibilidad de que tanto profesionales como pacientes analicen datos temporales sin conocimientos técnicos, gracias a interfaces amigables.

Tradicionalmente, la lógica difusa ha sido crucial para representar conocimiento experto bajo incertidumbre, ofreciendo una solución interpretable para traducir series temporales numéricas en lenguaje natural. Como consecuencia, en este trabajo proponemos una dualidad entre ambos paradigmas, pero donde la lógica difusa cumple un rol central en tres niveles:

- 1) En la generación de un conjunto de entrenamiento etiquetado con descripciones lingüísticas expertas.
- 2) En el diseño de prompts que guían al modelo en función de conocimiento clínico impreciso.
- 3) En la evaluación de los resultados generados por los modelos mediante criterios lingüísticos definidos en la metodología.

II. METODOLOGÍA GENERAL

La metodología se estructura en cuatro fases principales: (1) adquisición y preprocesamiento de los datos, (2) modelado experto y generación del dataset, (3) ingeniería de prompts y entrenamiento, y (4) evaluación.

A. Fase 1: Arquitectura de adquisición y preprocesamiento

Se emplea una arquitectura IoT centrada en sensores como el Freestyle Libre 3, que proporciona valores de glucosa cada 5 minutos. Los datos son recolectados mediante una app móvil, almacenados en una base de datos MongoDB y posteriormente procesados para su análisis.

Este resultado ha sido parcialmente respaldado por el proyecto PID2021-127275OB-I00 y el proyecto PID2021-126363NB-I00 financiados por el MICIU/AEI/10.13039/501100011033 y por “ERDF A way of making Europe”, y por el proyecto PDC2023-145863-I00 financiado por el MICIU/AEI/ 10.13039/501100011033 y por la “European Union NextGenerationEU/PRTR”.

Este resultado es un resumen del trabajo publicado con referencia doi.org/10.1016/j.cmpb.2025.108878

Cada serie diaria consta de hasta 288 muestras, que son segmentadas mediante el algoritmo de Ramer-Douglas-Peucker para simplificar su representación y facilitar la detección de patrones significativos como picos, valles o tendencias.

B. Fase 2: Modelado experto y generación del dataset

Esta es la fase clave, donde interviene el modelado experto mediante lógica difusa para transformar los datos numéricos en lenguaje humano interpretable. Se definen las siguientes categorías lingüísticas:

- Valores de glucosa: muy baja, baja, normal, alta, muy alta.
- Tendencias: descendente, estable, ascendente.
- Momentos del día: madrugada, mañana, tarde, noche.
- Cuantificadores: pocos, muchos, casi todos.

A partir de estas categorías se construyen protoformas lingüísticas, siguiendo la propuesta de Marín y Sánchez [1], de la forma:

Durante la mañana, muchos valores de glucosa fueron altos.

Cada protoforma se evalúa mediante su grado de verdad (DoT), calculado con funciones de pertenencia trapezoidales, y se filtran aquellas que no superan un umbral mínimo. También se aplican reglas de calidad textual para garantizar que las descripciones sean claras, no redundantes y respeten la coherencia temporal.

El dataset final se organiza en formato JSONL utilizando el sistema de roles estructurados de OpenAI [2]: el rol *system* para proporcionar las instrucciones iniciales y el contexto de la tarea, el rol *user* especificando los datos de la serie temporal, y el rol *assistant* con la salida esperada (el resumen modelado).

C. Fase 3: Ingeniería de prompts guiada por lógica difusa

La guía del modelo durante la generación textual es fundamental para reducir errores semánticos y alucinaciones. Para ello, se diseña un prompt estructurado en cuatro partes:

- 1) **Encabezado:** presentación del objetivo y estructura de los datos.
- 2) **Dominio:** contexto clínico relevante (por ejemplo, identificar hipoglucemias).
- 3) **Instrucción:** rol deseado (ej. endocrinólogo experto), nivel de detalle y periodo a analizar.
- 4) **Conocimiento:** explicitación de etiquetas difusas y cuantificadores que debe usar.

Esta estructura, inspirada en el enfoque prompt-as-prefix (PaP), facilita que el modelo interiorice las reglas del dominio y genere descripciones que reflejan el conocimiento experto. A través del rol *system* mencionado anteriormente, se introduce esta información, estableciendo un proceso de experimentación a través de la variación del contenido.

D. Fase 4: Evaluación

La evaluación de los resultados generados se basa en la metodología propuesta en [3], que propone una serie de criterios objetivos para valorar descripciones generadas automáticamente. Las dimensiones evaluadas son:

- **Calidad de la información:** precisión, relevancia, cobertura, corrección.
- **Capacidad de razonamiento:** comprensión y lógica de la salida.
- **Calidad comunicativa:** claridad del lenguaje, accesibilidad al usuario y uso de referencias temporales.
- **Seguridad del contenido:** ausencia de afirmaciones infundadas o no derivables de los datos.

Cada criterio es cuantificado a partir de la comparación con descripciones expertas, obviando la necesidad de evaluadores expertos y la subjetividad implícita a esta actuación, siendo una evaluación objetiva, y ampliando además la metodología [3] a escenarios con datos temporales.

III. RESULTADOS

Se aplicó esta metodología para entrenar y evaluar versiones de GPT-3.5, GPT-4o y GPT-4o-mini. Los mejores resultados se obtuvieron con GPT-4o, utilizando la serie original (no segmentada) y prompts enriquecidos con conocimiento difuso, alcanzando un 96% de conformidad con los criterios.

Los modelos sin ajuste fino o sin guía explícita tendieron a producir descripciones imprecisas, redundantes o incluso alucinaciones. El uso sistemático de lógica difusa en el dataset y en el prompt redujo significativamente estos errores.

IV. DISCUSIÓN

El estudio demuestra que la lógica difusa no solo es útil para representar conocimiento incierto, sino que constituye una herramienta poderosa para estructurar datasets, guiar el comportamiento generativo de modelos y evaluar su rendimiento. Este enfoque híbrido resulta especialmente valioso en contextos clínicos donde los errores pueden tener consecuencias serias.

Además, el diseño del sistema permite escalar la metodología a otros dominios como monitoreo cardíaco, análisis de actividad física o predicción meteorológica, siempre que exista una base de conocimiento experto susceptible de representación difusa.

V. CONCLUSIONES

Este trabajo propone una metodología robusta, replicable y segura para entrenar modelos de lenguaje bajo esquemas de incertidumbre. A través del uso transversal de la lógica difusa, desde la generación del dataset hasta su evaluación, se logra alinear modelos como GPT con los objetivos y limitaciones del razonamiento experto humano.

REFERENCES

- [1] N. Marín and D. Sánchez, "On generating linguistic descriptions of time series," *Fuzzy Sets and Systems*, vol. 285, pp. 6–30, 2016.
- [2] OpenAI, "Fine-tuning - openai api." <https://platform.openai.com/docs/guides/fine-tuning>, 8 2024. (Accessed on 11/10/2024).
- [3] T. Y. C. Tam, S. Sivarajkumar, S. Kapoor, A. V. Stolyar, K. Polanska, K. R. McCarthy, H. Osterhoudt, X. Wu, S. Visweswaran, S. Fu, *et al.*, "A framework for human evaluation of large language models in healthcare derived from literature review," *NPJ digital medicine*, vol. 7, no. 1, p. 258, 2024.

Knowledge base tuning by particle swarm optimization for centralized wireless sensor network clustering algorithm

Juan-Luis Herreros-Bodalo

*Department of Telecommunication Engineering
Universidad de Jaén
Linares, Spain
ORCID: 0009-0009-1257-2245*

Jose-Enrique Muñoz-Exposito

*Department of Telecommunication Engineering
Universidad de Jaén
Linares, Spain
ORCID: 0000-0002-7483-2964*

Antonio-Jesús Yuste-Delgado

*Department of Telecommunication Engineering
Universidad de Jaén
Linares, Spain
ORCID: 0000-0002-9321-1709*

Alicia Triviño-Cabrera

*Department of Electrical Engineering
Universidad de Málaga
Málaga, Spain
ORCID: 0000-0002-7516-2878*

Macarena Espinilla-Estevez

*Department of Informatics
Universidad de Jaén
Jaen, Spain
ORCID: 0000-0003-1118-7782*

David Diaz-Jimenez

*Department of Informatics
Universidad de Jaén
Jaen, Spain
ORCID: 0000-0003-1791-4258*

Juan Carlos Cuevas Martínez

*Department of Telecommunication Engineering
Universidad de Jaén
23700 Linares, Spain
ORCID: 0000-0003-3749-5986*

Abstract—Wireless Sensor Networks are integral to various Internet of Things applications [1], where energy efficiency and adaptive data routing significantly influence overall performance. Among different methods, clustering appears as one of the most suitable for efficient energy management [2]. This paper reviews CFC3PSO, a centralized fuzzy clustering method enhanced by particle swarm optimization to dynamically adjust rule weights within a fuzzy rule-based system. By considering multiple input parameters and adapting the knowledge base of the fuzzy rule-based system at the base station across different network life phases, the proposed method outperforms traditional and recent clustering strategies in simulations while dealing with the high amount of uncertainty of these kind of approaches. The results in multiple deployment scenarios demonstrate substantial gains in key performance metrics in wireless sensor networks.

Index Terms—wireless sensor network, particle swarm optimization, fuzzy logic, clustering

Clustering is a common strategy for energy conservation in WSN. Clustering consists of grouping nodes around other nodes selected as cluster heads (CHs). These CHs receive the information for other nodes, aggregate and send the aggregated data to a central node called base station (BS). Thus, this article discusses CFC3PSO [3], a novel centralized clustering approach that combines fuzzy rule-based systems (FRBS) and particle swarm optimization (PSO) to improve network performance and service life. In CFC3PSO, due to the intrinsic uncertainty of these kinds of system and, consequently, the difficulties to set up the semantic rules of a knowledge base (KB) base for a FRBS, an optimization process is carried out to obtain the weights to apply for each rule of KB to maximize the performance of the network.

I. INTRODUCTION

Wireless Sensor Networks (WSNs) are pivotal in Internet of Things (IoT) environments such as healthcare, smart cities, industrial automation, and precision agriculture. These networks comprise numerous battery-powered sensor nodes, where energy efficiency is crucial to operational longevity.

This research has been funded by the Ministry of Science, Innovation, and Universities, as part of the 2023 call for proposals for grants for proof-of-concept projects under the State Plan for Scientific, Technical, and Innovation Research for the period 2021-2023, within the framework of the Recovery, Transformation, and Resilience Plan, with grant number PDC2023-145863-I00

II. CFC3PSO CLUSTERING METHOD SUMMARY

CFC3PSO uses a Type-1 Mamdani fuzzy inference system to evaluate the suitability of each node to become a CH based on five input variables:

- Distance to the base station (dBS)
- Centrality (C) which is calculated based on the distance of the nodes in a CH to the center of the sensing area of their cluster.
- Node degree (ND) represents the number of neighbors that each node has within a given distance.
- Residual energy (RE) of the node.

- Cluster head rotation frequency (CHR) informs about the number of times a node has been selected as a CH [4].

Each input is fuzzified using three triangular membership functions. The system output, which indicates the likelihood that a node will become a CH (*chance*), is represented using 11 triangular fuzzy sets.

The knowledge base consists of 243 fuzzy rules (3 sets per 5 inputs). Initially, each rule is equally weighted. The novelty lies in the dynamic weighting of these rules, optimized for three different network life cycle stages:

- 1) Initial phase: from startup to first node death (FND).
- 2) Intermediate phase: from FND to when half of the nodes are dead (HND).
- 3) Final phase: from HND to last node death (LND).

In order to achieve rule weights that would be applied to a wider number of different wireless sensor network layouts, the optimization process uses three typical scenarios whose details can be found in Section III.

The tuning process employed a particle swarm optimization (PSO) strategy, referred to as PSO rule weight tuning (PSO-RWT). PSO was used due to its fast convergence, which makes it suitable and widely used in other clustering methods in WSN [5]. The PSO-RWT algorithm is employed to optimize the weight of the rules in the three phases enumerated above. In order to obtain the best FND value, the PSO-RWT is applied to 30 different network layouts, 10 for each of the scenarios described in Section III. Therefore, a set of rule weights is obtained and applied to the clustering algorithm from startup until the first node depletes its battery. Following this, the process is repeated twice: firstly from FND to HND, and then from HND to LND. This is done to obtain the intermediate weights and the weights that maximize the lifetime of the network.

Once the optimization process is complete, the rule weights are loaded into the BS that runs the clustering algorithm. The weights of the rules are then changed according to the stage of the network.

III. SIMULATION SCENARIOS

Three test scenarios were developed, each involving a square area measuring 100×100 m², with 100 nodes randomly placed within this area. The primary distinction between these scenarios lies in the BS location, as outlined below:

- Scenario 1: BS at corner of the sensing field
- Scenario 2: BS outside the sensing area at (150, 50) m.
- Scenario 3: BS in the center (50, 50) m.

Each scenario includes 30 different deployments of the 100 nodes, with the same initial energy. The simulations use the first-order radio model detailed in [6]. Table 2 in the original article details the values of the setup parameters for the energy model.

IV. COMPARATIVE EVALUATION

In order to assess the CFC3PSO, its performance was compared against that of four pre-existing clustering algorithms:

- LEACH: A distributed and probabilistic CH selection method [6]. LEACH is used just as a pattern because it is a well-known clustering method.
- CHEF: Uses residual energy and local density as fuzzy inputs [7]
- CRT2FLACO: Combines Type-2 fuzzy logic with ant colony optimization [8]
- ICUU: Applies the silhouette index and DBSCAN for intelligent clustering [9]

Performance metrics evaluated included:

- FND (First Node Dies).
- HND (Half of the Nodes Dead).
- LND (Last Node Dies).

Tables 4–6 in the original article summarize the results [3]. In all scenarios, CFC3PSO outperformed the other algorithms in all three metrics. For instance, in Scenario 3 (central BS), CFC3PSO achieved an FND of 1847.2 rounds compared to LEACH's 764.8, and an LND of 2344.7 versus LEACH's 1020 rounds.

V. CONCLUSION

The CFC3PSO approach presents a powerful method for energy-efficient CH selection in WSNs. Its use of a dynamically adapted fuzzy knowledge base, tuned through PSO for different network lifespans, significantly boosts performance. It offers a robust, computationally manageable solution that adapts well to diverse deployment scenarios.

REFERENCES

- [1] Awatef Benfradj Guiloufi, Salim El Khediri, Nejah Nasri, and Abdenaceur Kachouri. A comparative study of energy efficient algorithms for iot applications based on wsns. *Multimedia Tools and Applications*, 82(27):42239–42275, 2023.
- [2] Amin Shahraki, Amir Taherkordi, Øystein Haugen, and Frank Eliassen. Clustering objectives in wireless sensor networks: A survey and research direction analysis. *Computer Networks*, 180:107376, 2020.
- [3] Jose-Enrique Muñoz-Exposito, Antonio-Jesus Yuste-Delgado, Alicia Triviño-Cabrera, and Juan-Carlos Cuevas-Martinez. Optimizing rule weights to improve frbs clustering in wireless sensor networks. *Sensors*, 24(17), 2024. URL: <https://www.mdpi.com/1424-8220/24/17/5548>, doi:10.3390/s24175548.
- [4] A. J. Yuste-Delgado, J. C. Cuevas-Martínez, and A. Triviño-Cabrera. Statistical normalization for a guided clustering type-2 fuzzy system for wsn. *IEEE Sensors Journal*, 22(6):6187–6195, 2022. doi:10.1109/JSEN.2022.3150066.
- [5] Richa Sharma, Vasudha Vashisht, and Umang Singh. Metaheuristics-based energy efficient clustering in wsns: challenges and research contributions. *IET Wireless Sensor Systems*, 10(6):253–264, 2020.
- [6] Wendi B Heinzelman, Anantha P Chandrakasan, and Hari Balakrishnan. An application-specific protocol architecture for wireless microsensor networks. *IEEE Transactions on wireless communications*, 1(4):660–670, 2002.
- [7] Indranil Gupta, Denis Riordan, and Srinivas Sampalli. Cluster-head election using fuzzy logic for wireless sensor networks. In *3rd Annual communication networks and services research conference (CNSR'05)*, pages 255–260. IEEE, 2005.
- [8] Wei-Xian Xie, Qi-Ye Zhang, Ze-Ming Sun, and Feng Zhang. A clustering routing protocol for wsn based on type-2 fuzzy logic and ant colony optimization. *Wireless Personal Communications*, 84:1165–1196, 2015.
- [9] LAXMINARAYAN Sahoo, S Sen, KS Tiwary, SARBAST Moslem, and TAPAN Senapati. Improvement of wireless sensor network lifetime via intelligent clustering under uncertainty. *IEEE Access*, 2024.

Uso de las dCF -integrales en sistemas de clasificación basados en reglas difusas para abordar problemas no balanceados

José Antonio Sanz y Mikel Sesma-Sara
Departamento de Estadística, Informática y Matemáticas
Universidad Pública de Navarra
Pamplona, España
joseantonio.sanz, mikel.sesma@unavarra.es

Resumen—Los problemas de clasificación no balanceados son muy habituales en el mundo real y suponen un reto para los sistemas de aprendizaje puesto que tienden a ignorar la clasificación de los ejemplos de la clase minoritaria, que generalmente es la de mayor interés. Existen numerosas formas de abordarlos, pero en esta contribución utilizamos un sistema de clasificación basado en reglas difusas, llamado FARCI, que afronta estos problemas sin necesidad de aplicar técnicas de muestreo de datos. Nuestra propuesta es modificar la fase de agregación de la información dada por las reglas disparadas en el proceso de inferencia. En concreto, proponemos utilizar las dCF -integrales, que son una generalización de la integral Choquet, que además de reemplazar el producto por una función F , también cambia la diferencia por una función de disimilitud restringida, d .

Index Terms—Clasificación no balanceada, sistemas de clasificación basados en reglas difusas, funciones de agregación, generalizaciones de la integral Choquet.

I. INTRODUCCIÓN

Los problemas de clasificación no balanceados representan un gran desafío en aprendizaje automático, ya que la distribución de clases en los datos no es homogénea. Esta situación aparece con frecuencia en contextos críticos para la sociedad, como el diagnóstico médico, la detección de fraudes o la predicción de fallos en sistemas. En tales casos, la clase minoritaria suele ser la más relevante, por lo que es esencial diseñar métodos de clasificación que manejen adecuadamente este desequilibrio para no comprometer la toma de decisiones.

Los sistemas de clasificación basados en reglas difusas (SCBRD) han demostrado ser herramientas eficaces para resolver problemas de clasificación gracias a su capacidad para representar el conocimiento de forma interpretable y gestionar la incertidumbre presente en los datos. Estos sistemas permiten generar reglas lingüísticas que facilitan la comprensión del modelo por parte de expertos humanos, a la vez que ofrecen resultados competitivos en términos de precisión.

Una propuesta reciente dentro de los SCBRD para abordar problemas de clasificación no balanceados es el método llamado FARCI [1] (Fuzzy Association Rule-based Classification

for Imbalanced datasets), una extensión del clasificador FARC-HD [2] específicamente adaptada para manejar el desbalanceo de clases. En FARCI se modifican todas las fases del proceso de aprendizaje de FARC-HD para mejorar su rendimiento en contextos no balanceados, logrando resultados competitivos frente a otros métodos del estado del arte.

Un componente fundamental de FARCI es su método de inferencia difusa, llamado de combinación aditiva, en el que usa la suma normalizada como función de agregación para fusionar la información de todas las reglas disparadas durante la clasificación de un ejemplo. Es decir, se suman los grados de asociación de todas las reglas disparadas por cada clase, y se selecciona aquella con el valor agregado más alto. Esta estrategia resulta sencilla e intuitiva, pero puede generar sesgos cuando existe un número significativamente mayor de reglas disparadas de una clase respecto a la otra, lo cual es frecuente en escenarios de clasificación no balanceada.

Dado este posible sesgo, surge la hipótesis de que el uso de funciones de agregación alternativas podría mejorar el comportamiento del sistema. En este sentido, se han propuesto numerosas generalizaciones de la integral Choquet [3], que permiten una fusión más flexible y controlada de la información. Estas generalizaciones incluyen tanto funciones de agregación como de pre-agregación y han sido ampliamente estudiadas en la literatura reciente para abordar problemas de clasificación estándar. Recientemente, se ha desarrollado una nueva generalización de la integral Choquet, denominada dCF -integral [4], basada en funciones de disimilitud restringidas. En concreto, se sustituye la diferencia clásica en la integral de Choquet por una función de disimilitud restringida.

En este trabajo proponemos el uso de las dCF -integrales para reemplazar la suma normalizada utilizada actualmente en el método de inferencia de FARCI. Concretamente, utilizamos las dos mejores funciones de disimilitud y 19 de las 21 funciones de agregación presentadas en [4]. Evaluamos el rendimiento de la nueva propuesta en el mismo marco experimental que el utilizado en [1], pero centrado únicamente en los 44 conjuntos de datos con un elevado ratio de no balanceo ($IR \geq 9$).

II. PROPUESTA

En este trabajo utilizamos FARCI [1] como SCBRD, que está basado en FARC-HD [2] que es un clasificador diseñado para abordar problemas de clasificación estándar. Para abordar problemas de clasificación no balanceados, en FARCI se propuso el uso de funciones de agregación promedio para modelar la intersección de los antecedentes de las reglas y además se modificaron las 3 etapas de aprendizaje de FARC-HD:

- En el algoritmo A-priori reemplaza la confianza como medida de calidad de las reglas por el *lift*.
- Se modifica la filosofía del uso del esquema de ponderación de ejemplos, aplicada para seleccionar las reglas más interesantes, para potenciar la selección de reglas de la clase minoritaria.
- El algoritmo genético utiliza el F-score en la función de *fitness*, puesto que es una métrica apropiada para este tipo de problemas.

En la fase de inferencia, FARCI utiliza la combinación aditiva como método de razonamiento difuso. Es decir, para agregar la información local de las reglas disparadas y obtener la información global (por clases), se aplica como función de agregación la suma normalizada. Posteriormente, se predice la clase asociada al valor agregado más alto.

Sin embargo, en la práctica se disparan más reglas de la clase positiva; por ejemplo, en la mejor configuración de FARCI el número de reglas generadas de la clase positiva casi duplica al de la negativa. Este hecho puede provocar que el uso de la suma no sea equitativo para ambas clases, ya que el valor agregado de la clase positiva tiende a ser mayor incluso en situaciones en las que los valores, de manera individual, sean inferiores a los de la clase negativa por el hecho de tener más reglas disparadas de dicha clase.

Por ello, en este trabajo proponemos el uso de las *dCF*-integrales en vez de la suma normalizada para obtener la información global en el proceso de inferencia ya que su carácter no promedio puede mejorar el comportamiento del sistema en estas situaciones.

III. ESTUDIO EXPERIMENTAL

Para evaluar el rendimiento de la nueva propuesta utilizamos la mejor medida difusa (cardinalidad exponencial), las dos mejores funciones de disimilitud restringidas (δ_2 y δ_5) y probaremos 19 de las 21 funciones de agregación presentadas en la Tabla II de [4] (no utilizamos el producto drástico ni el mínimo nilpotente), numeradas correlativamente. Utilizamos el mismo estudio experimental que el utilizado en [1] pero nos centramos en los 44 datasets con un ratio de no balanceo alto ($IR \geq 9$).

La Tabla I resume los resultados. Cada fila corresponde a una de las 19 funciones de agregación (F) y las columnas muestran los valores obtenidos para $d = \delta_2$, ya que con $d = \delta_5$ el rendimiento fue inferior. Para cada F se muestran la media en test (Med. Tst) del F-score calculada sobre los 44 conjuntos

de datos¹, junto con el número de ellos en los que logra el mejor resultado; y el p-valor (p-val) del test de Wilcoxon que compara FARCI (rango positivo, R+) con la función F (rango negativo, R-).

Tabla I
RESULTADOS DEL ESTUDIO EXPERIMENTAL UTILIZANDO $d = \delta_2$. EN NEGRITA ESTÁ RESALTADO EL MEJOR RESULTADO.

F	Med. Tst (Vict.)	p-val (R+;R-)
1	0.604 (3)	< 0,01 (270.5;719.5)
2	0.599 (4)	< 0,01 (269.5;720.5)
3	0.602 (3)	0.04 (320.5;669.5)
4	0.609 (3)	0.199 (385.0;605.0)
5	0.598 (2)	0.02 (294.5;695.5)
6	0.617 (6)	0.784 (471.5;518.5)
7	0.612 (3)	0.302 (406.5;583.5)
8	0.603 (3)	0.434 (428.0;562.0)
9	0.609 (4)	0.149 (371.5;618.5)
10	0.622 (7)	0.866 (509.5;480.5)
11	0.603 (4)	0.134 (366.5;623.5)
12	0.611 (7)	0.636 (454.5;535.5)
13	0.607 (3)	0.08 (346.5;643.5)
14	0.600 (3)	0.165 (376.0;614.0)
15	0.600 (8)	0.241 (394.5;595.5)
16	0.607 (5)	0.121 (362.0;628.0)
17	0.606 (6)	< 0,01 (271.5;718.5)
18	0.594 (1)	0.923 (487.0;503.0)
19	0.614 (5)	-
FARCI	0.619 (6)	-

IV. CONCLUSIONES

Los resultados experimentales muestran que el uso de HM (10), F_{NA2} (19), O_{mM} (6) y de O_{RS} (12) como funciones F en las *dCF*-integrales, al usar $d = \delta_2$, tienen potencial para ser utilizadas en el proceso de inferencia de FARCI. Al utilizar $d = \delta_5$, los resultados, en general, han sido inferiores aunque algunas funciones F también pueden presentar cierto potencial. Por ello, el uso de las *dCF*-integrales podrá ser la base de futuros estudios en los que se trate de mejorar el rendimiento de este clasificador utilizando generalizaciones de la integral Choquet.

REFERENCIAS

- [1] J. Sanz, M. Sesma-Sara, H. Bustince, A fuzzy association rule-based classifier for imbalanced classification problems, *Information Sciences* 577 (2021) 265–279.
- [2] J. Alcalá-Fdez, R. Alcalá, 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* 19 (5) (2011) 857–872.
- [3] G. P. Dimuro, J. Fernández, B. Bedregal, R. Mesiar, J. A. Sanz, G. Lucca, H. Bustince, The state-of-art of the generalizations of the choquet integral: From aggregation and pre-aggregation to ordered directionally monotone functions, *Information Fusion* 57 (2020) 27–43.
- [4] J. Wierzchynski, G. Lucca, G. P. Dimuro, E. N. Borges, J. A. Sanz, T. d. C. Asmus, J. Fernández, H. Bustince, *dCF*-integrals: Generalizing *cF*-integrals by means of restricted dissimilarity functions, *IEEE Transactions on Fuzzy Systems* 31 (1) (2023) 160–173.

¹Los resultados detallados se pueden ver en https://github.com/JoseanSanz/FARCI_Choquet_integral_generalizations

An order-oriented approach to scoring hesitant fuzzy elements

Luis Merino^{*†}, Gabriel Navarro^{*†}, Carlos Salvatierra^{*}, and Evangelina Santos^{*†}

^{*}Department of Algebra, University of Granada, 18071 Granada, Spain

[†]Institute of Mathematics of Granada (IMAG), University of Granada, Spain

Abstract—Although *scores* are the preferred tool for ranking hesitant fuzzy elements (HFEs), most existing proposals are defined ad hoc and their link with the order structures that underlie comparison are not always clear. This communication lays the algebraic foundations of *order-oriented scores*. In this line we introduce the general notion of a $\preceq$ -score, i.e. a score that is monotone with respect to an explicit order $\preceq$; revisit the main orders on THFEs and clarify their mutual refinements and relationships with scores; and show how the symmetric order provides the natural ground for scores that satisfy desirable normative properties such as (weak/strong) monotonicity with respect to unions or the (weak/strong) Gärdenfors principle.

Keywords—hesitant fuzzy elements, score, orders on hesitant fuzzy elements, symmetric order, lattice structure.

I. INTRODUCTION

Hesitant fuzzy sets (HFSs), introduced by Torra [1], generalize classical fuzzy sets by allowing multiple membership degrees per element. This flexibility has enhanced their applicability in areas such as decision-making or information fusion. However, defining robust methods for ranking hesitant fuzzy elements (HFEs) remains a significant challenge.

One of the most famous ranking methods are based on the direct application of scores on HFEs. The main handicap with them is that most existing scores have been developed independently of any underlying order structure. In this work, we address this gap by introducing an *order-oriented* framework, in which scores are explicitly defined with respect to a given order. This perspective offers a clear and systematic method for verifying the meeting of normative properties of a score, such as those proposed by Alcantud *et al.* [2], through order compatibility, facilitating the study and selection of scores in diverse application domains by simply analyzing their behavior relative to the underlying order.

The main contributions of this work are as follows: (i) we analyze the main classical orders on THFEs, showing their general failure to induce lattice structures; (ii) we explicitly define scores with respect to a given order; concretely, we introduce the concepts of weak and strong $\preceq$ -scores; and (iii) we emphasize how the symmetric order, not only induces a lattice on THFEs, but also their compatible scores support core principles such as monotonicity with respect to unions and the Gärdenfors property, as discussed by Alcantud *et al.* [2].

II. ORDERS ON TYPICAL HESITANT FUZZY ELEMENTS

Let $F^*([0, 1])$ denote the set of all non-empty finite subsets of the unit interval, known as typical hesitant fuzzy elements

(THFEs) [3, Definition 3.1]. Several orders have been proposed to compare such elements. The *list order* [4] compares two THFEs of equal size by sorting their elements increasingly and comparing them component-wise. To allow comparisons between THFEs of different cardinalities, Xu and Xia [4] introduced the *pessimistic* (also known as the Xu-Xia order in works such as [5], [6]) and *optimistic orders*. We have proved that despite their intuitive appeal, these orders fail to generate a lattice structure on $F^*([0, 1])$, contrary to what was traditionally assumed in works such as [6]. For example, if we consider the pair $A = \{0.1, 0.4, 0.5, 0.7\}$ and $B = \{0.3, 0.6\}$, we show that there does not exist a greatest lower bound under the pessimistic order. Similar issues arise under the optimistic order. Other well-known orders for THFEs share this same limitation. For example, the *least-common-multiple (LCM) order* [7, Proposition 2.6], which aligns cardinalities through an LCM process, does not guarantee the existence of suprema or infima [7, Corollary 3.7]. Similarly, the *Zhang and Yang order* introduced in [8, Definition 2.4] was shown not to induce a lattice in [9, Theorem 3.8].

To address these lattice shortcomings, Jara *et al.* [10] introduced three orders on the class of all right-closed subsets of $[0, 1]$, denoted by $P^+([0, 1])$, the class of all left-closed subsets of $[0, 1]$, denoted by $P^-([0, 1])$, and the family of non-empty subsets of $[0, 1]$, denoted by $P^*([0, 1])$, respectively:

$$\begin{aligned} A \leq^r B &\iff A^+ \leq B^+ \text{ and } A < B \setminus A, \\ A \leq^l B &\iff A^- \leq B^- \text{ and } A \setminus B < B, \\ A \leq^0 B &\iff A \leq^r B \text{ and } A \leq^l B. \end{aligned}$$

Here A^+ and A^- denote the maximum and minimum of A , and $X < Y$ if for all $x \in X$, $y \in Y$, $x < y$. Orders $\leq^r$ and $\leq^l$ refine the optimistic and pessimistic ones, respectively, and their conjunction $\leq^0$ is known as the *symmetric order*.

Moreover, $(P^*([0, 1]), \leq^0)$ forms a *bounded lattice* [10], where the meet $A \wedge_0 B$ and join $A \vee_0 B$ generalize Zadeh's min and max operations- solving the open problems presented in [11, Challenge 5.3] and [12, Remark 2]. As a consequence, this structure naturally restricts to $F^*([0, 1])$, enabling $\leq^0$ to endow THFEs with a genuine lattice.

III. WEAK AND STRONG $\preceq$ -SCORES

We define a score w.r.t an order as follows:

Definition 1. Given a family $G \subseteq P^*([0, 1])$ and an order $\preceq$ on G , a function $s : G \rightarrow [0, 1]$ is a $\preceq$ -score (respectively, a

strong $\preceq$ -score) on G if, for all $A, B \in G$, $A \preceq B$ implies $s(A) \leq s(B)$ (respectively, $A \prec B$ implies $s(A) < s(B)$).

This approach allows scores to be explicitly tied to the structural properties of the order. It also ensures coherence under refinement: if $\preceq_2$ refines $\preceq_1$, then every (weak or strong) $\preceq_2$ -score is also a $\preceq_1$ -score.

Among all known scores, *symmetric scores*, that is, scores that are monotone with respect to the symmetric order, stand out. A key advantage of using scores relative to the symmetric order is that they naturally align with a well-defined lattice structure. As a consequence of this structure, we obtain the following result:

Proposition 1. *If s is a $\leq^0$ -score, then $s(A \wedge_0 B) \leq \min\{s(A), s(B)\}$ and $\max\{s(A), s(B)\} \leq s(A \vee_0 B)$.*

Furthermore, symmetric scores align precisely with some of the critical normative axioms introduced by Alcantud *et al.* [2]. Notably, while in general boundedness [2, Definition 5] implies compatibility [2, Remark 3] but not the converse, the two properties are equivalent when working with $\leq^0$ -scores:

Proposition 2. *Let $s : P^*([0, 1]) \rightarrow [0, 1]$ be a $\leq^0$ -score. Then s satisfies compatibility if and only if it satisfies boundedness.*

In addition, symmetric scores satisfy the following characterization in terms of monotonicity with respect to unions:

Proposition 3. *Let $G \subseteq P^*([0, 1])$ be a class closed under union and difference, and let $s : G \rightarrow [0, 1]$ be a map. Then:*

- i) *s is a $\leq^0$ -score if and only if it satisfies the weak monotonicity with respect to unions [WMU].*
- ii) *s is a strong $\leq^0$ -score if and only if it satisfies the strong monotonicity with respect to unions [SMU] [2, Definition 9].*

When G is the set of THFEs, symmetric scores also admit an equivalent characterization via the Gärdenfors properties:

Proposition 4. *Let $s : F^*([0, 1]) \rightarrow [0, 1]$ be a map. Then:*

- i) *s is a $\leq^0$ -score if and only if it satisfies [WMU] or, equivalently, the weak Gärdenfors property [WG].*
- ii) *s is a strong $\leq^0$ -score if and only if it satisfies [SMU] or, equivalently, the Gärdenfors property [G] [2, Definition 12].*

Another important case arises when G is the set of closed intervals of $[0, 1]$, denote by $I^c([0, 1])$. In this setting, the following result is satisfied:

Proposition 5. *Let $s : I^c([0, 1]) \rightarrow [0, 1]$ be a map. Then:*

- i) *s is a $\leq^0$ -score if and only if $a \leq c$ and $b \leq d$ implies $s([a, b]) \leq s([c, d])$.*
- ii) *s is a strong $\leq^0$ -score if and only if it satisfies the extreme monotonicity property [2, Definition 15].*

Examples of symmetric scores include the arithmetic mean, geometric mean, minimum, and maximum scores [13], [14]. In contrast, some classical scores, such as the product $P(A) = \prod_{x \in A} x$, for $A \in F^*([0, 1])$, are not.

REFERENCES

- [1] V. Torra, "Hesitant fuzzy sets," *Int. J. Intell. Syst.*, 2010.
- [2] J. C. R. Alcantud *et al.*, "Scores of hesitant fuzzy elements revisited," *Inf. Sci.*, 2023.
- [3] B. Bedregal, R. H. Santiago, H. Bustince, D. Paternain, and R. Reiser, "Typical hesitant fuzzy negations," *Int. J. Intell. Syst.*, vol. 29, pp. 525–543, 2014.
- [4] Z. Xu and M. Xia, "Distance and similarity measures for hesitant fuzzy sets," *Inf. Sci.*, vol. 181, no. 10, pp. 2128–2138, 2011.
- [5] H. Santos, B. Bedregal, R. Santiago, and H. Bustince, "Typical hesitant fuzzy negations based on Xu-Xia-partial order," in *Proc. IEEE Int. Conf. Fuzzy Syst. (FUZZ-IEEE)*, pp. 1385–1392, 2014.
- [6] H. Santos, B. Bedregal, R. Santiago, H. Bustince, and E. Barrenechea, "Construction of typical hesitant triangular norms regarding Xu-Xia-partial order," in *Proc. 16th World Congress Int. Fuzzy Syst. Assoc. (IFSA) and 9th Conf. Eur. Soc. Fuzzy Logic Technol. (EUSFLAT)*, pp. 953–959, Atlantis Press, 2015.
- [7] L. Garmendia, R. González del Campo, and J. Recasens, "Partial orderings for hesitant fuzzy sets," *Int. J. Approx. Reasoning*, vol. 84, pp. 159–167, 2017.
- [8] H. Zhang, J. Ye, and S. Yang, "Typical hesitant fuzzy rough sets," *IEEE Trans. Fuzzy Syst.*, vol. 26, no. 5, pp. 2962–2971, Oct. 2015.
- [9] Y. Xu, L. Liu, and X. Zhang, "Multilattices on typical hesitant fuzzy sets," *Inf. Sci.*, vol. 491, pp. 63–73, 2019.
- [10] P. Jara, L. Merino, G. Navarro, and E. Santos, "A lattice structure on hesitant fuzzy sets," *IEEE Trans. Fuzzy Syst.*, vol. 31, no. 6, pp. 2018–2025, 2023.
- [11] R. M. Rodríguez *et al.*, "A position and perspective analysis of hesitant fuzzy sets on information fusion in decision making. Towards high quality progress," *Inf. Fusion*, vol. 29, pp. 89–97, 2016.
- [12] H. Bustince, E. Barrenechea, M. Pagola, J. Fernández, Z. Xu, B. Bedregal, J. Montero, H. Hagra, F. Herrera, and B. De Baets, "A historical account of types of fuzzy sets and their relationships," *IEEE Trans. Fuzzy Syst.*, vol. 24, no. 1, pp. 179–194, 2016.
- [13] B. Farhadinia, "A novel method of ranking hesitant fuzzy values for multiple attribute decision-making problems," *Int. J. Intell. Syst.*, vol. 28, no. 8, pp. 752–767, 2013.
- [14] B. Farhadinia, "A series of score functions for hesitant fuzzy sets," *Inf. Sci.*, vol. 277, pp. 102–110, 2014.

On the non-falsity and threshold preserving variants of MTL logics*

Francesc Esteva
AI Research Institute (IIIA)
CSIC
Bellaterra, Spain
esteva@iiia.csic.es

Joan Gispert
Dept. Maths and Computer Science
Univ. de Barcelona
Barcelona, Spain
jgispert@ub.edu

Lluís Godo
AI Research Institute (IIIA)
CSIC
Bellaterra, Spain
godo@iiia.csic.es

Abstract—In this paper we study the definition and axiomatisation of non-falsity preserving and threshold preserving companions of several extensions of the Monoidal t-norm based fuzzy logic MTL. More in detail, we first extend some recent preliminary results on non-falsity preserving logics, and then we present a new study on threshold-preserving companions of the three most prominent fuzzy logics, Łukasiewicz, Product and Gödel logics.

Index Terms—Mathematical fuzzy logic, monoidal t-norm based logic, non-falsity preserving logics, threshold preserving logics

—This paper has been accepted to be presented at the conference EUSFLAT 2005, July 21-15, Riga, Letonia—

I. INTRODUCTION

Fuzzy logics are logics of *graded truth* and their main feature is that they allow to interpret formulas in a linearly ordered scale of truth-values. In particular, systems of fuzzy logic have been in-depth developed within the frame of mathematical fuzzy logic (MFL) [5], [9]. In deductive systems in MFL, mostly with semantics in the real unit interval $[0, 1]$, the usual notion of deduction is defined by requiring the preservation of the truth-value 1 (full *truth-preservation*), understood as absolute truth. In the last years, there have been several works studying different variants of deductive systems of fuzzy logics (see e.g. [3], [4], [6]), mainly by moving from the (full) truth-preserving paradigm to the paraconsistent *degree-preserving paradigm*, in which a conclusion follows from a set of premises if, for all evaluations, the truth degree of the conclusion is greater or equal than those of the premises, see e.g. [2]. Another way of defining a (paraconsistent) variant of a fuzzy logic is put forward in [1] for the particular case of Łukasiewicz fuzzy logic. In this approach, the notion of consequence at work is the *non-falsity preservation* principle, according to which a conclusion follows from a set of premises whenever if the premises are non-false (have a truth-value greater than 0), so must be the conclusion. Still, another parameterized family of variants of a given truth-preserving fuzzy logic is the class of consequence relations determined

by a given threshold $0 < a \leq 1$. In such variants, a conclusion follows from a set of premises whenever the premises get a truth-value at least a this must be also the case for the conclusion.

In this paper, we first review the non-falsity preserving companions of extensions of the so-called Involutive MTL logic, and second we study the syntactical characterisation of threshold-preserving logics, focusing the study on the three most prominent extensions of MTL, namely Łukasiewicz, Product and Gödel fuzzy logics.

II. PRELIMINARIES

Most well known and studied systems of mathematical fuzzy logic are the so-called *t-norm based fuzzy logics*, corresponding to formal many-valued calculi with truth-values in the real unit interval $[0, 1]$ and with a conjunction and an implication interpreted respectively by a (left-) continuous t-norm and its residuum, and thus, including e.g. the well-known Łukasiewicz, Gödel and Product infinitely-valued logics, corresponding to the calculi defined by Łukasiewicz, min and product t-norms respectively. The most general t-norm based fuzzy logic is the logic MTL (monoidal t-norm based logic) introduced in [8], whose theorems correspond to the common tautologies of all many-valued calculi defined by a left-continuous t-norm and its residuum [10].

Given a left-continuous t-norm $*$, we denote by $[0, 1]_*$ the standard MTL-algebra determined by $*$, i.e. $[0, 1]_* = ([0, 1], \min, \max, *, \rightarrow, 0, 1)$, where $\rightarrow$ is the residuum of $*$ and the negation $\neg$ is defined as $\neg x = x \rightarrow 0$. Let L be any extension of MTL complete w.r.t. the family $\mathcal{C}_L = \{[0, 1]_* \mid [0, 1]_* \text{ is a } L\text{-algebra}\}$ of standard L -algebras. Then the typical notion of truth-preserving logical consequence is the following: for every set of formulas $\Gamma \cup \{\varphi\}$:

$$\Gamma \models_L \varphi \quad \text{if, for all } [0, 1]_* \in \mathcal{C}_L \text{ and } [0, 1]_*\text{-evaluation } e, \\ \text{if } e(\psi) = 1 \text{ for all } \psi \in \Gamma, \text{ then } e(\varphi) = 1.$$

This can be generalised by considering logics that take $[a, 1]$ with $a \in (0, 1]$, or $(a, 1]$ with $a \in [0, 1)$ as sets of designated values. Then the corresponding companions of the logic L are defined respectively as follows:

*Esteva and Godo acknowledge partial support by the Spanish project LINEXSYS (PID2022-139835NB-C21), and Gispert by the Spanish project SHORE (PID2022-141529NB-C21), both funded by MCIU/AEI/10.13039/501100011033

- $\Gamma \models_L^a \varphi$ if, for all $[0, 1]_* \in \mathcal{C}_L$ and $[0, 1]_*$ -evaluation e ,
if $e(\psi) \geq a$ for any $\psi \in \Gamma$, then $e(\varphi) \geq a$.
- $\Gamma \models_L^{(a)} \varphi$ if, for all $[0, 1]_* \in \mathcal{C}_L$ and $[0, 1]_*$ -evaluation e ,
if $e(\psi) > a$ for any $\psi \in \Gamma$, then $e(\varphi) > a$.

The extreme cases are the 1-preserving logic $\models_L^1 = \models_L$, which is explosive, and the non-falsity preserving logic $\models_L^{(0)}$, which is paraconsistent w.r.t. $\neg$. Observe that the finitary versions of both logics are strongly related. Indeed:

Lemma 2.1: For every pair of formulas φ, ψ the following relation holds: $\varphi \models_L^{(0)} \psi$ iff $\neg\psi \models_L^1 \neg\varphi$.

III. NON-FALSITY PRESERVING COMPANIONS OF INVOLUTIVE MTL EXTENSIONS

Assume L is an axiomatic extension of IMTL, the extension of MTL with the involutiveness axiom (Inv) “ $\neg\neg\varphi \rightarrow \varphi$ ”, complete with respect to a class of standard algebras $\mathcal{C}_L$, and whose corresponding notion of proof is denoted $\vdash_L$.

In order to syntactically characterise $\models_L^{(0)}$, defined by the class of matrices $\mathcal{C}_L^{(0)} = \{\langle [0, 1]_*, F_{(0)} \mid \langle [0, 1]_*, F_1 \rangle \in \mathcal{C}_L \}$, the following system nf-L , called the *non-falsity preserving companion* of L , is defined in [7] as follows.

Definition 3.1: The calculus nf-L is defined by the axioms of L and the following rules:

- Rule of Adjunction: (Adj) $\frac{\varphi, \psi}{\varphi \wedge \psi}$
- Reverse Modus Ponens: (MP^r) $\frac{\neg\psi \vee \chi, \neg\varphi \vee \neg(\varphi \rightarrow \psi) \vee \chi}{\neg\varphi \vee \neg(\varphi \rightarrow \psi)}$
- Restricted Modus Ponens: (r-MP) $\frac{\varphi, \varphi \rightarrow \psi}{\psi}$, if $\vdash_L \varphi \rightarrow \psi$

The logic nf-L has been shown to be complete in [7] with respect to the intended semantics.

Theorem 3.1 (Completeness): Let L be an axiomatic extension of IMTL. The calculus nf-L is sound and complete w.r.t. to the class of matrices $\mathcal{C}_L^{(0)}$.

As a direct corollary, Definition 3.1 provides us with complete axiomatisations of non-falsity preserving companions of prominent IMTL logics like Łukasiewicz or Nilpotent Minimum logics.

IV. THRESHOLD-PRESERVING LOGICS

We consider now logics preserving lower bounds of truth-values of the form $\models_L^a$ and $\models_L^{(a)}$ for some $a \in (0, 1]$, for extensions L of MTL that are complete with respect to some class of standard L -algebras $\mathcal{C}_L$.

A. A general setting for Łukasiewicz logics

We state two general but sufficient conditions for a logic L to guarantee a finitary axiomatisation $\models_L^a$:

(C1) L satisfies a form of global Deduction Theorem in the sense that there exists a term t such that:

$$\Gamma \cup \{\varphi\} \vdash_L \psi \quad \text{iff} \quad \Gamma \vdash_L t(\varphi) \rightarrow \psi$$

(C2) The logic $\models_L^a$ is interpretable in L , that is, there exists a term r such that:

$$\varphi \models_a \psi \quad \text{iff} \quad r(\varphi) \models_L r(\psi)$$

Theorem 4.1: Let L be an extension of MTL satisfying conditions (C1) and (C2). Then the calculus L_a defined syntactically by the axioms of L and the following rules:

- the rules of L restricted to theorems of L ,
- the rule of Adjunction, and
- the restricted rule: (R_{t,r}) $\frac{\varphi, \vdash_L t(r(\varphi)) \rightarrow r(\psi)}{\psi}$

is a sound and complete axiomatisation of the finitary $\models_L^a$.

This general result can be applicable to the case of L being any finite-valued Łukasiewicz logic $\mathbb{L}_n$ and to Łukasiewicz logic $\mathbb{L}$ expanded with Baaz-Monteiro operator Δ .

B. The cases of Gödel and Product logics

Gödel and Product logics fall outside the scope of Theorem 4.1, and hence they require a specific consideration.

The case of Gödel logic: The analysis of Gödel logic G turns out to be very simple. Gödel logic, the extension of MTL with the axiom (Con) “ $\varphi \rightarrow \varphi \& \varphi$ ”, is complete with respect to the single matrix $M_1 = \langle [0, 1]_G, \{1\} \rangle$, where $[0, 1]_G$ is the standard Gödel algebra $([0, 1], \min, \max, *_G, \rightarrow_G, 0, 1)$, with $*_G = \min$ and $\rightarrow_G$ is its residuum.

Theorem 4.2: $\models_G^{(0)}$ coincides with classical logic, while for any $a \in (0, 1]$, $\models_G^a = \models_G^{(a)} = \models_G$.

The case of Product logic: Product logic Π is defined as the axiomatic extension of MTL with axioms of divisibility and contraction, see e.g. [9]. Product logic is complete with respect to the single matrix $M_1 = \langle [0, 1]_\Pi, \{1\} \rangle$, where $[0, 1]_\Pi$ is the standard product algebra $([0, 1], \min, \max, *_\Pi, \rightarrow_\Pi, 0, 1)$, with $*_\Pi$ being the product t-norm and $\rightarrow_\Pi$ is its residuum.

The intersection of all the logics $\models_\Pi^b$ for all $b \in (0, 1]$ is what is known as the degree-preserving companion of $\models_\Pi$, and is denoted as $\models_\Pi^{\leq}$, see [2]. Then, the main result regarding Product logic is the following.

Theorem 4.3:

- $\models_\Pi^{(0)}$ is classical logic.
- For every $a \in (0, 1)$, $\models_\Pi^a = \models_\Pi^{(a)} = \models_\Pi^{\leq} \subsetneq \models_\Pi$.

REFERENCES

- [1] Avron, A.: Paraconsistent fuzzy logic preserving non-falsity. *Fuzzy Sets and Systems* **292**, 75–84 (2016)
- [2] Bou, F., Esteva, F., Font, J.-M., Gil, A., Godo, L., Torrens, A., Verdú, V.: Logics preserving degrees of truth from varieties of residuated lattices. *Journal of Logic and Computation* **19**(6), 1031–1069 (2009)
- [3] Coniglio, M.-E., Esteva, F., Godo, L.: Logics of formal inconsistency arising from systems of fuzzy logic. *Logic Journal of the IGPL* **22**(6), 880–904 (2014)
- [4] Coniglio, M.-E., Esteva, F., Gispert, J., Godo, L.: Degree-preserving Gödel logics with an involution: intermediate logics and (ideal) paraconsistency. In: Arielli, O., Zamansky, A. (eds.), *Arnon Avron on Semantics and Proof Theory of Non-Classical Logics*, pp. 107–139, Springer (2021)
- [5] Cintula, P., Noguera C.: A general framework for mathematical fuzzy logic. In: Cintula, P., Hájek, P., Noguera, C. (eds.), *Handbook of Mathematical Fuzzy Logic - Volume 1*, Studies in logic, Mathematical Logic and Foundations, Vol 37, pp. 103–207, College Publications (2011)
- [6] Ertola, R.-C., Esteva, F., Flaminio, T., Godo, L., Noguera, C.: Paraconsistency properties in degree-preserving fuzzy logics. *Soft Computing* **19**(3), 531–546 (2015)
- [7] Esteva, F., Gispert, J., Godo, L.: On the paraconsistent companions of involutive fuzzy logics that preserve non-falsity. In: Lesot, M.-J. et al (eds.), *IPMU 2024, LNCS* (to appear)
- [8] Esteva, F., Godo, L.: Monoidal t-norm based logic: towards a logic for left-continuous t-norms. *Fuzzy Sets and Systems* **124**, 271–288 (2001)
- [9] Hájek, P.: *Metamathematics of Fuzzy Logic*. Vol. 4 of Trends in Logic - Studia Logica Library. Kluwer (1998)
- [10] Jenci, S., Montagna, F.: A proof of standard completeness for Esteva and Godo’s logic MTL. *Studia Logica* **70**(2), 183–192 (2002)

Una Propuesta de Función Generadora de Funciones Opuestas para Conjuntos Difusos Pareados

Javier Bonilla
Comisión Nacional del Mercado de Valores
Universidad Complutense de Madrid
Madrid, España
javierlb@ucm.es

Abstract—En este trabajo se presenta una nueva propuesta de función generadora de funciones opuestas que permite generalizar el cálculo de funciones de no pertenencia a modelos de conjuntos borrosos no solamente intuicionistas sino también pareados.

Keywords—Lógica difusa, conjuntos intuicionistas, conjuntos pareados, funciones de negación.

I. INTRODUCCIÓN

Sea X el conjunto de N observaciones con n características o variables para cada una de las observaciones $X_i = (x_{i1}, x_{i2}, \dots, x_{in})$ que se desean clasificar en K clases o conjuntos, en otras palabras, realizar una partición del conjunto de datos X . De tal forma que, según [1] se obtiene la matriz M de la partición difusa a partir de las funciones de pertenencia μ .

$$\begin{cases} \mu: \mathbb{R}^n \rightarrow [0,1]^K \\ X_i = (x_{i1} \dots x_{in}) \rightarrow \mu_k(X_i) = \mu_k(x_{i1} \dots x_{in}) = \mu_{ik} \\ M = [\mu_{ik}]_{(N \times K)} \end{cases}$$

Sin que necesariamente se cumpla que $\sum_{k=1}^K \mu_{ik} = 1$ para cada observación, condición impuesta por [2]. Esta condición implica que una observación i no pertenezca a una clase k sería el complementario de μ_{ik} , es decir, $1 - \mu_{ik}$.

Posteriormente, en [3] aparecieron los conjuntos intuicionistas, que son una generalización de los conjuntos propuestos por [1], donde se tenía en cuenta el grado de no pertenencia η . Es decir, $\{(X_i, \mu_k(X_i), \eta_k(X_i)) | X_i \in X\}$. Si $\mu_{ik} + \eta_{ik} = 1 \Rightarrow \eta_{ik} = 1 - \mu_{ik}$, en otro caso surge el grado de duda o de indeterminación π_{ik} , tal que existe la condición $\mu_{ik} + \eta_{ik} + \pi_{ik} = 1$.

Otra generalización de los conjuntos difusos serían los conjuntos paraconsistentes $\{(X_i, \mu_k(X_i), \eta_k(X_i)) | X_i \in X\}$, donde a diferencia de los conjuntos de Anassov, cuando $\mu_{ik} + \eta_{ik} > 1$ surge el grado de conflicto, es decir, $\psi_{ik} = \mu_{ik} + \eta_{ik} - 1$, como se puede ver por ejemplo en [4].

Los conjuntos intuicionistas y los conjuntos paraconsistentes se pueden agrupar en los conjuntos pareados, ver por ejemplo [4-6], donde además se añade el concepto de ambigüedad, ξ_{ik} , lo que da origen a una estructura denominada penta-valorada.

II. LA FUNCIÓN DE NEGACIÓN

Definición 1. La función de negación fuerte es una aplicación $N: [0,1] \rightarrow [0,1]$ tal que tiene las siguientes propiedades, ver por ejemplo [7-8]:

1. $N(0) = 1$.
2. $N(1) = 0$.
3. Es estrictamente decreciente, dados $x_1, x_2 \in X$ con $x_1 < x_2 \Rightarrow N(x_1) > N(x_2)$.
4. Es continua.
5. Es involutiva: $N(N(x_i)) = x_i \forall x_i \in X$.

Así, el caso más sencillo de función de negación fuerte es el complementario $\eta_{ik} = 1 - \mu_{ik}$. En la definición 2.3 de [9] y en sus observaciones 2.4 se indica que la condición de involución es una propiedad muy fuerte que implica que la negación sea biyectiva, estrictamente decreciente, continua y con un punto fijo o punto de equilibrio. Lo que implica por [9] que la Definición 1 solamente serían necesarios su propiedad 5 y la 3. Como también indica [9] hay otros autores que no imponen la involución como propiedad como por ejemplo [7].

En el caso de conjuntos intuicionistas se necesita lo que se denomina función generadora φ^{int} para obtener la función de funciones de no pertenencia, ver [10-11]. Esta nueva función es una negación fuerte más una nueva condición: $\varphi^{int}(y) \leq 1 - y, \forall y \in [0,1]$.

Ejemplos de funciones generadoras para conjuntos intuicionistas se pueden encontrar por ejemplo en [12-13]. Si se considera a α_k como un parámetro por cada clase o conjunto difuso, se tienen los siguientes ejemplos concretos:

1. Yager [14]. $\eta_{ik} = (1 - \mu_{ik}^{\alpha_k})^{\frac{1}{\alpha_k}} \forall i, k$ con $\alpha_k \in [1, \infty)$.
2. Sugeno [15]. $\eta_{ik} = \frac{(1 - \mu_{ik})}{(1 + \alpha_k \mu_{ik})} \forall i, k$ con $\alpha_k \in [0, \infty)$.

Según lo comentado en [10] para la función de negación de Yager con $\alpha_k > 1$ no se cumpliría la condición número 1 de la dedición de negación fuerte. En el caso de conjuntos paraconsistentes se puede utilizar de manera similar al intuicionista una función generadora φ^{par} , que es una

negación fuerte con la condición: $\varphi^{par}(y) \geq 1 - y, \forall y \in [0,1]$.

De manera análoga a los conjuntos intuicionistas se tienen los siguientes ejemplos concretos:

1. Yager [14]. $\eta_{ik} = (1 - \mu_{ik}^{\alpha_k})^{\frac{1}{\alpha_k}} \forall i, k \text{ con } \alpha_k \in (0, 1]$
2. Sugeno [15]. $\eta_{ik} = \frac{(1 - \mu_{ik})}{(1 + \alpha_k \mu_{ik})} \forall i, k \text{ con } \alpha_k \in (-1, 0]$

Es necesario indicar para la función de negación de Yager que $\mu_{ik} > \alpha_k$, ya que en otro caso puede dar valores mayores que uno.

Y en el caso de los conjuntos pareados se propone la siguiente función que permite generar funciones de no pertenencia o funciones opuestas a partir de las funciones de pertenencia:

$$\eta_{ik} = 1 - \frac{\mu_{ik}}{\sum_{k=1}^K \mu_{ik}} = 1 - \beta_i \mu_{ik}$$

Con la condición $1/\beta_i = \sum_{k=1}^K \mu_{ik} \geq \max_k \{\mu_{ik}\}$ cuando $\sum_{k=1}^K \mu_{ik} < 1$, ya que $\mu_{ik} + \eta_{ik} + \pi_{ik} = 1$.

Que tiene las siguientes características:

1. Si $\mu_{ik} = 0$ entonces $\eta_{ik} = 1$
2. Cuando $\sum_{k=1}^K \mu_{ik} \geq 1$ y $\mu_{ik} = 1$, entonces $\eta_{ik} \geq 0$. Este resultado no es estrictamente cero por la existencia del conflicto que hace que $\mu_{ik} + \eta_{ik} > 1$.
3. Es continua.
4. No siempre es estrictamente decreciente, ya que depende del cociente: $\frac{\beta_j}{\beta_i} = \frac{\sum_{k=1}^K \mu_{jk}}{\sum_{k=1}^K \mu_{ik}} \leq 1$ para que sea estrictamente decreciente en el caso de dos observaciones distintas i y j . Pero para una misma observación se cumple que si $\mu_{ik} \leq \mu_{ig}$ entonces $1 - \beta_i \mu_{ik} \geq 1 - \beta_i \mu_{ig}$.
5. En general no es involutiva ya que $1 - \beta_i (1 - \beta_i \mu_{ik}) = 1 - \beta_i + \beta_i^2 \mu_{ik} \neq \mu_{ik}$ excepto cuando $\beta_i = 1$, es decir cuando $\sum_{k=1}^K \mu_{ik} = 1$.
6. Si $\sum_{k=1}^K \mu_{ik} = 1$ ($\beta_i = 1$) entonces $\sum_{k=1}^K \pi_{ik} = 0$ y $\sum_{k=1}^K \psi_{ik} = 0$. Modelo estándar, genera conjuntos complementarios.
7. Si $\sum_{k=1}^K \mu_{ik} < 1$ ($\beta_i > 1$) entonces $\sum_{k=1}^K \pi_{ik} > 0$ y $\sum_{k=1}^K \psi_{ik} = 0$. Modelo intuicionista pero sin cumplir la involución. En este caso en particular se genera antónimos según el modelo pareado definido en [5,6],
8. Si $\sum_{k=1}^K \mu_{ik} > 1$ ($\beta_i < 1$) entonces $\sum_{k=1}^K \pi_{ik} = 0$ y $\sum_{k=1}^K \psi_{ik} > 0$. Modelo paraconsistente pero sin cumplir la involución. En este caso en particular se

genera antagonismos según el modelo pareado definido en [5,6].

Por último, se puede generalizar la función propuesta considerando un parámetro $\gamma \geq 1$ tal que pondere todos los μ_{ik} como μ_{ik}^γ .

III. CONCLUSIONES

La nueva función generadora propuesta en este documento permite calcular el valor de las funciones de no pertenencia, negación u opuestas, de tal forma que se puede utilizar en diversos tipos de conjuntos difusos frente a la fórmula de Yager que solamente funciona bien para conjuntos intuicionistas.

La nueva función generadora propuesta tiene la ventaja sobre la fórmula de Sugeno en que no es necesario estimar el valor del parámetro α_k , es decir, ya viene dado por las funciones de pertenencia.

Por el contrario, la nueva función propuesta no va cumplir siempre todas las condiciones de una negación fuerte, ya que es no involutiva en general, excepto cuando $\beta_i = 1$.

Referencias

- [1] L.A. Zadeh, "Fuzzy Sets," Inf. Control, 8 (3), pp. 338-353, 1965.
- [2] E.H. Ruspini, "A New Approach to Clustering," Inf. Control, 15, pp. 22-32, 1969.
- [3] K.T. Atanassov, "Intuitionistic Fuzzy Sets," Fuzzy Sets and Systems, 20, pp. 58-96, 1986.
- [4] V. Patrascu, "A New Penta-valued Based Knowledge Representation," IPMU 2008, pp. 17-22, 2008.
- [5] J.T. Rodríguez, C. Franco, D. Gómez & J. Montero, "Paried Structures and other Opposites-based Models," 16th IFSA y 9th EUSFLAT, 2015.
- [6] J. Montero, H. bustince, C. Franco, J.T. Rodríguez, D. Gómez, M. Pagola & et al., "Paried Structures in Knowledge Representation," Knowledge-Based Systems, pp. 1-9, 2016.
- [7] E. Trillas, "Sobre Funciones de Negación en la Teoría de Subconjuntos Difusos," Stochastica, III (1), pp. 47-59, 1979.
- [8] B. Bede, Mathematics of Fuzzy Sets and Fuzzy Logic, Berlín: Srpinger, 2013.
- [9] H. Bustince, M.J. Campión, L. de Miguel & E. Induráin, "Strong negations and restricted equivalence functions revisited: An analytical and topological approach," Fuzzy Sets and Systems, 441, pp. 110-129, 2022.
- [10] H. Bustince, J. Kacprzyk & V. Mohedano, "Intuitionistic Fuzzy Generator. Application to Intuitionistic Fuzzy Complementation," Fuzzy Sets and Systems, 114, pp. 485-504, 2000.
- [11] J. Montero, D. Gómez & H. Bustince, "On the Relevance of some Families of Fuzzy Sets," Fuzzy Sets and Systems, 158, pp. 2429-2442, 2007.
- [12] N.R. Pal, H. Bustince, M. Pagola, U.K. Mukherjee, D.P. Goswami & G. Beliakov, "Uncertainties with Atanassov's Intuitionistic Fuzzy Sets: Fuzziness and Lack of Knowledge," Informations Sciences (228), pp. 61-74, 2013.
- [13] E. Trillas, "On Negation Functions in Fuzzy Set Theory," Advances in fuzzy Logic, pp.31-43, 1998.
- [14] R.R. Yager, "On the Measures of fuzziness and Negation. Part II: lattices," Information and Control (44), pp. 236-260, 1980.
- [15] M. Sugeno, "Fuzzy Measures and Fuzzt Integrals: a survey," Fuzzy Automata and Decision Process, pp. 82-102, 1977.

Representativeness systems and coverage indexes

1st Inmaculada Gutiérrez

*Faculty of Statistical Studies,
UCM, Madrid, Spain
inmaguti@ucm.es
0000 - 0002 7103 393X*

2nd J. Tinguaro Rodríguez

*Faculty of Mathematics,
UCM, Madrid, Spain
jtrodrig@ucm.es
0000 - 0003 - 4989 - 2348*

3rd Xabier González

*Department of Statistics, Computer Science and Mathematics
UPNA, Pamplona, Spain
xabier.gonzalez@unavarra.es
0000 - 0001 - 7262 - 753X*

4th Daniel Gómez

*Faculty of Statistical Studies,
UCM, Madrid, Spain
dagomez@ucm.es
0000 - 0001 - 9548 - 5781*

5th Javier Montero

*Faculty of Mathematics,
UCM, Madrid, Spain
monty@ucm.es
0000 - 0001 - 8333 - 2155*

6th Humberto Bustince

*Department of Statistics, Computer Science and Mathematics
UPNA, Pamplona, Spain
bustince@unavarra.es
0000 - 0002 - 1279 - 6195*

Abstract—This work presents the key contributions of Gutiérrez et al. [EUSFLAT (2025)], which explores the challenge of assessing the quality of representativeness systems (RS). RSs aim to model scenarios where a set of sources provides representation or service to a collection of objects. Building upon this notion, the authors introduce the concept of coverage indexes as mathematical tools to evaluate the effectiveness of such systems. Several subfamilies of coverage indexes are defined within a formal framework, alongside a general construction method for their derivation. These indexes form the foundation for analyzing RSs across various contexts. For instance, in Gutiérrez et al. [EUSFLAT (2025)] these tools were applied to the assessment of cluster quality, revealing several promising features and encouraging further exploration of coverage-based evaluation methods.

Index Terms—Coverage, Representativeness Systems, Aggregation Functions

I. INTRODUCTION AND RELATED WORK

Representativeness is a fundamental concept in systems where a set of entities serve as representatives of a population of objects or individuals, hereafter referred to as a representativeness system (RS). This notion captures the extent to which objects can be identified with or covered by their representatives, based on similarity, shared interests, or the ability of representatives to meet objects' needs. Inspired by the seminal ideas in [1], representativeness can be naturally modeled as a gradable property measured through *coverage* levels, reflecting partial or full representation. In this way, a natural modeling of such representativeness could consider it as a gradable property, which can be assessed in terms of coverage levels or degrees, in the sense that, depending on how they are related, a given representative might only partially represent an object or cover its needs, e.g. because the former is not fully able to effectively provide the latter with the required service, or because a maximum quality service can not always be provided.

RSs apply in diverse contexts, such as political elections, real estate, healthcare, telecommunications, and more, where the degree of coverage varies according to relevant criteria. Given contextual variability, a universal formalization of *coverage degrees* is challenging, and operational definitions must be adapted accordingly. For instance, in some scenarios, coverage degrees might be determined based on a measure of similarity or dissimilarity between entities, whereas in other cases, they could rely on how well the services provided align with those required, the existence of links between the parties involved, and so forth. Though it may be difficult to formalize, once defined, these coverage degrees provide valuable insights into the effectiveness and dependability of representativeness systems, enabling the assessment of adequate coverage, possible redundancies, uncovered areas, and the effectiveness of clustering or assignment methods [1], [5].

Once a formal characterization is established for a specific scenario, the information provided by a coverage degree becomes significantly more valuable when processed through an information aggregation (or fusion) procedure. This process, commonly understood as the application of mathematical or computational techniques to combine and condense the available data into meaningful indices that reflect various dimensions of the system under consideration (see [1], [3], [4], [6], [9]), is a key tool for leveraging coverage information for the purposes described above. In the context of RSs, such aggregation can be broadly understood as a procedure that takes coverage data as input and produces a single summary value, whose interpretation depends on the specific characteristics of the aggregation method and may reflect different aspects of system quality.

In [8], we first proposed a formal definition of a representativeness system (RS). We then focused on evaluating the extent to which the analyzed objects are represented within such systems. To this end, we introduced the concept of a coverage index as a function satisfying a number of desirable properties, which serve as the foundation for the development of quality measures for RSs. These properties

This research has been partially funded by the Government of Spain, Grant Plan Nacional de I+D+i, PID2021-122905NB-C21 and PID2022-136627NB-I00, supported by MCIN/AEI/10.13039/501100011033/FEDER.

are designed to reflect intuitive principles such as fairness and symmetry, ensuring that the evaluation remains invariant under object permutations. Following the general idea that a good representational structure should guarantee that most objects are reasonably covered by at least one source, we characterized a quality index aimed at capturing this global perspective. This leads to the definition of the global coverage index (GCI), whose formulation is expressed in particular through the functional composition. Within this framework, grouping functions [2] are employed to model a disjunctive behavior, aligning with the ideal scenario in which each object is fully represented by a single source.

II. MAIN CONTRIBUTIONS

Next, we recall and discuss with some detail the main contributions established in [8]:

- **Formal definition of a representativeness system (RS)**, as a triplet $RS = (X, S, \mathbf{U})$, where $X = \{x_1, \dots, x_N\}$ denotes a finite set of objects, and $S = \{s_1, \dots, s_K\}$ a finite set of representatives or coverage sources. For each $i \in \{1, \dots, N\}$ and $j \in \{1, \dots, K\}$, let $\mathbf{U} = (u_{ij})$ be an $N \times K$ matrix with entries $u_{ij} \in I = [0, 1]$, where u_{ij} represents the degree of coverage of object $x_i \in X$ by representative $s_j \in S$. The value $u_{ij} = 1$ corresponds to full coverage, while $u_{ij} = 0$ denotes a complete lack of coverage. Intermediate values indicate partial degrees of representativeness. The choice of the interval $[0, 1]$ is conventional but not restrictive, as any closed interval $I = [a, b] \subseteq \mathbb{R}$ could be considered instead. The matrix $\mathbf{U} = (u_{ij})_{N \times K}$ is referred to as the *coverage matrix*. If $\mathcal{U}_{N,K}(I)$ denotes the set of all $N \times K$ matrices with entries in I , then clearly $\mathbf{U} \in \mathcal{U}_{N,K}(I)$.
- **Formal definition of coverage indexes**, as a general notion of quality index for an RS to evaluate how well objects in X are represented by sources in S from the information in the coverage matrix $\mathbf{U}$. We also provide an alternative characterization designed to satisfy a specific symmetry condition. This leads to the notion of a *fair coverage index*, which ensures that any such index remains invariant under permutations of the objects, as well as under permutations of the coverage degrees associated with each object.
- **Formal characterization of a Global Coverage Index (GCI)** as a general family of mappings representing different global conceptions of quality of an RS. We also work on a specific characterization of GCIs based on compositions of two functions, $M \circ G$. The first one, $G : \mathcal{U}_{N,K}(I) \rightarrow I^N$, is a vector function with all its components being equal, i.e. such that for any $\mathbf{U} \in \mathcal{U}_{N,K}(I)$ it is $G(\mathbf{U}) = (G(\mathbf{u}_1), \dots, G(\mathbf{u}_N)) \in I^N$, where $G : I^K \rightarrow I$. The second function, $M : I^N \rightarrow I$, produces the final aggregated value $M(G(\mathbf{u}_1), \dots, G(\mathbf{u}_N)) = (M \circ G)(\mathbf{U}) \in I$. A relevant example of GCI is the index obtained by using $G = \max$ and $M = \text{mean}$. We then focus on obtaining sufficient conditions on the functions G and M to ensure that composition $M \circ G$ captures

the intended meaning of a GCI as a quality metric that measures the extent to which *most* objects in X are *adequately* covered by the representatives in S .

We also present and demonstrate several desirable properties of these composition-based GCIs, and illustrate their potential usefulness in cluster validation analysis through a practical application: using several examples, we show how coverage indexes can be useful in capturing different aspects of a partition quality.

Thus, the described work [8] represents an advance within a well-established line of research, where many directions still warrant thorough investigation. Future research in this direction will focus on a more in-depth study of the mathematical properties of the different families of coverage indices, as well as developing their applications in various fields, such as social network analysis [7], [12]. Another research line will address the use of OWA operators [10], particularly MEOWA [11] operators, to fulfill the averaging role of the function M in GCIs.

REFERENCES

- [1] Amo, A., Montero, J., Biging, G., Cutello, V.: Fuzzy classification systems. *European Journal of Operational Research* **156**(2), 495–507 (2004)
- [2] Bustince, H., Pagola, M., Mesiar, R., Hüllermeier, E., Herrera, F.: Grouping, overlap, and generalized bintropic functions for fuzzy modeling of pairwise comparisons. *IEEE Transactions on Fuzzy Systems* (2012)
- [3] Calvo, T., Kolesárová, A., Komorníková, M., Mesiar, R.: Aggregation Operators: Properties, Classes and Construction Methods. In: *Studies in Fuzziness and Soft Computing*, pp. 3–104. Springer (2002). https://doi.org/10.1007/978-3-7908-1787-4_1
- [4] Castiblanco, F., Franco, C., Rodríguez, J.T., Montero, J.: Degree of global covering and global overlapping in solvency fuzzy classification. In: León Castro, E., Blanco Mesa, F., Gil Lafuente, A., Merigó, J., Kacprzyk, J. (eds.) *Intelligent and Complex Systems in Economics and Business*. vol. Pre Press. Springer Nature Switzerland AG (2020)
- [5] Castiblanco, F., Franco, C., Rodríguez, J.T., Montero, J.: Evaluation of the quality and relevance of a fuzzy partition. *Journal of Intelligent & Fuzzy Systems* **39**(3), 4211–4226 (2020)
- [6] Gómez, D., Rodríguez, J.T., Montero, J., Bustince, H., Barrenechea, E.: N-Dimensional overlap functions. *Fuzzy Sets and Systems* **287**, 57–75 (2016)
- [7] Gutiérrez, I., Barroso, M., Gómez, D., Castro, C.: Improving community detection algorithms in directed graphs with fuzzy measures. an application to mobility networks. *Expert Systems With Applications* **269**(126305) (2025)
- [8] Gutiérrez, I., Rodríguez, J.T., González, X., Gómez, D., Montero, J., Bustince, H.: Measuring representativeness through coverage degrees and indexes. In: 14th Conference of the European Society for Fuzzy Logic and Technology, EUSFLAT (2025), preprint
- [9] Mesiar, R., Kolesárová, A., Calvo, T., Komorníková, M.: A review of aggregation functions, vol. 20, pp. 121–144. Springer (2008)
- [10] Yager, R.R.: On ordered weighted averaging aggregation operators in multicriteria decision-making. *IEEE Transactions on Systems and Cybernetics* **18**(1), 183–190 (1988)
- [11] Yager, R.R.: Weighted maximum entropy OWA aggregation with applications to decision making under risk. *IEEE Transactions on Systems, Man, and Cybernetics - Part A: Systems and Humans* **39**(3), 555–564 (2009)
- [12] Zhang, S., Yu, J., Hu, M.: An edge server placement based on graph clustering in mobile edge computing. *Scientific Reports* **14**(1) (2024)

Funciones de agregación inspiradas en la integral Choquet

1st Humberto Bustince

Depto Estadística, Informática y Matemáticas
Universidad Pública de Navarra
Pamplona, España
bustince@unavarra.es

2nd Javier Fernández

Depto Estadística, Informática y Matemáticas
Universidad Pública de Navarra
Pamplona, España
fcojavier.fernandez@unavarra.es

3rd Julio Lafuente

Depto Estadística, Informática y Matemáticas
Universidad Pública de Navarra
Pamplona, España
lafuente@unavarra.es

4th Graçaliz Dimuro

Centro de Ciências Computacionais
Universidade Federal do Rio Grande
Rio Grande, Brasil
gracalizdimuro@furg.br

Resumen—En este trabajo, presentamos una nueva familia de funciones de agregación inspiradas en la integral Choquet, pero reemplazando la medida por funciones más generales. Es un resumen de los resultados presentados en el trabajo [3] que se encuentra bajo revisión en la revista *Information Fusion*.

Index Terms—Fusión de datos, función de agregación, integral Choquet

I. INTRODUCTION

La fusión de datos es un elemento esencial en la mayor parte de los problemas científicos y técnicos. Una de las herramientas más utilizadas para este proceso son las funciones de agregación [1], [2], [5]. Y entre estas, en los últimos tiempos han adquirido particular relevancia las agregaciones basadas en integrales Choquet y Sugeno [4], [6]

En este trabajo, y siguiendo la línea de recientes extensiones y generalizaciones de la noción de integral Choquet, proponemos una nueva familia de agregaciones cuya construcción se inspira en esta. En concreto, proponemos reemplazar la medida utilizada para la construcción de la integral Choquet por funciones crecientes, obteniendo una nueva familia de operadores que puede interpretarse como una combinación convexa de las entradas pero con pesos que dependen de todas las entradas menos de aquellas directamente influenciadas por el peso. Esta construcción permite recuperar operadores bien conocidos, como los OWA, así como algunos casos de la integral Choquet discreta.

II. PRELIMINARES

Sea $n \geq 3$. Para $i, j \in \{1, \dots, n\}$ con $i \neq j$ y $\mathbf{x} = (x_1, \dots, x_n) \in [0, 1]^n$, denotamos por $\hat{\mathbf{x}}_{i,j}$ el resultado de la proyección $P_{(i,j)} : [0, 1]^n \rightarrow [0, 1]^{n-2}$ que omite las componentes i and j .

Este trabajo está financiado por el proyecto PID2022-136627NB-I00 (MCIN/AEI/10.13039/501100011033/FEDER, UE)

Dada $(x_1, \dots, x_n) \in [0, 1]^n$, denotamos por $(\cdot)$ una permutación $\{1, \dots, n\}$ tal que:

$$x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)}.$$

Para $n \geq 2$, escribimos $N = \{1, \dots, n\}$ y 2^N denota el conjunto de subconjuntos de N .

Recordemos que una función de agregación es una función $A : [0, 1]^n \rightarrow [0, 1]$ creciente y tal que $A(0, \dots, 0) = 0$ y $A(1, \dots, 1) = 1$. Entre las funciones de agregación, cabe destacar los dos ejemplos siguientes:

- (i) ([7]) Sea $n \geq 2$ y un vector de pesos $\mathbf{w} = (w_1, \dots, w_n) \in [0, 1]^n$ (i.e., $w_1 + \dots + w_n = 1$). el operador OWA asociado a $\mathbf{w}$ es la función $F_{\mathbf{w}} : [0, 1]^n \rightarrow [0, 1]$ dada por:

$$F_{\mathbf{w}}(x_1, \dots, x_n) = w_1 x_{(n)} + w_2 x_{(n-1)} + \dots + w_n x_{(1)}.$$

- (ii) Sea $m : 2^N \rightarrow [0, 1]$ una medida difusa (no aditiva). La integral Choquet con respecto a m es la función $C_m : [0, 1]^n \rightarrow [0, 1]$ dada por:

$$C_m(\mathbf{x}) = \sum_{i=1}^n (x_{(i)} - x_{(i-1)}) \cdot m(A_{(i)}),$$

donde $A_{(i)} = \{(i), (i+1), \dots, (n)\}$ y, por definición, $x_{(0)} = 0$.

III. FUNCIONES INSPIRADAS EN LA INTEGRAL CHOQUET: DEFINICIÓN Y PROPIEDADES

El principal concepto introducido en el trabajo [3] es el siguiente:

Definición 1: Sea $n \geq 3$. Dada una función simétrica y creciente $F : [0, 1]^{n-2} \rightarrow [0, 1]$, definimos la función inspirada en la integral Choquet $\Gamma_F : [0, 1]^n \rightarrow [0, 1]$ como:

$$\Gamma_F(x_1, \dots, x_n) = x_{(1)} + \sum_{i=1}^{n-1} (x_{(i+1)} - x_{(i)}) F(\hat{\mathbf{x}}_{(i), (i+1)}). \quad (1)$$

IV. CONCLUSIONS

En [3] se presenta el concepto de función de agregación inspirada en la integral Choquet, y se estudian sus propiedades y relaciones con otros operadores de agregación bien conocidos. Consideramos que este trabajo puede ser el inicio de un nuevo campo de investigación para aquellos problemas donde la determinación de medidas para el uso de integrales difusas puede resultar complejo.

REFERENCIAS

- [1] G. Beliakov, H. Bustince, T. Calvo, A Practical Guide to Averaging Functions, Springer, Berlin, 2016.
- [2] G. Beliakov, A. Pradera, T. Calvo, Aggregation Functions: A Guide for Practitioners, Springer, Berlin, 2007.
- [3] H. Bustince, R. Mesiar, G. Dimuro, J. Fernandez, J. Lafuente, B. De Baets, Choquet-inspired aggregation functions, submitted to Information Fusion
- [4] G. Choquet, Theory of capacities, Annales de l'Institut Fourier 5 (1953–1954) 131–295.
- [5] M. Grabisch, J. Marichal, R. Mesiar, E. Pap, Aggregation Functions, Cambridge University Press, Cambridge, 2009.
- [6] M. Sugeno, Theory of Fuzzy Integrals and its Applications. Ph.D. Dissertation. Tokyo Institute of Technology, 1975.
- [7] 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.

Por ejemplo, tomemos $n = 3$. Si $F : [0, 1] \rightarrow [0, 1]$ es una función creciente, Γ_F viene dada por:

$$\Gamma_F(\mathbf{x}) = x_{(1)} + (x_{(2)} - x_{(1)})F(x_{(3)}) + (x_{(3)} - x_{(2)})F(x_{(1)}),$$

que podemos reescribir como:

$$\Gamma_F(\mathbf{x}) = x_{(1)}(1 - F(x_{(3)})) + x_{(2)}(F(x_{(3)}) - F(x_{(1)})) + x_{(3)}F(x_{(1)}).$$

En particular, si tomamos $F(x) = x$, obtenemos:

$$\Gamma_F(\mathbf{x}) = x_{(1)} + x_{(2)}x_{(3)} - x_{(1)}x_{(2)} = (1 - x_{(2)})x_{(1)} + x_{(2)}x_{(3)}.$$

Es decir, Γ_F es una combinación convexa del mínimo y el máximo, con pesos determinados por la mediana.

El principal resultado en [3] es el siguiente

Teorema 1: Sea $n \geq 3$. Para cualquier función simétrica y creciente $F : [0, 1]^{n-2} \rightarrow [0, 1]$, la función Γ_F es una función de agregación idempotente.

En el paper se analizan diversas propiedades de estas funciones. En cuanto a la relación de los nuevos operadores con funciones de agregación conocidas, por razones de espacio simplemente mencionaremos aquí los siguientes resultados sobre la conexión con los OWA

Para $k \in \{1, \dots, n-2\}$, tomemos como F el k -ésimo estadístico. Es decir:

$$F(y_1, \dots, y_{n-2}) = Q_k(y_1, \dots, y_{n-2}) = y_{(k)}.$$

Entonces se verifica que:

$$\begin{aligned} \Gamma_{Q_k}(\mathbf{x}) &= x_{(1)} + \sum_{i=1}^k (x_{(i+1)} - x_{(i)})x_{(k+2)} \\ &+ \sum_{i=k+1}^{n-1} (x_{(i+1)} - x_{(i)})x_{(k)}, \end{aligned}$$

que puede simplificarse como:

$$\begin{aligned} \Gamma_{Q_k}(\mathbf{x}) &= x_{(1)} + (x_{(k+1)} - x_{(1)})x_{(k+2)} \\ &+ (x_{(n)} - x_{(k+1)})x_{(k)} \\ &= x_{(1)}(1 - x_{(k+2)}) + x_{(k+1)}(x_{(k+2)} \\ &- x_{(k)}) + x_{(k)}x_{(n)}. \end{aligned}$$

por otro lado, tenemos el siguiente resultado.

Teorema 2: Sea $n \geq 3$ y sean $F_1, F_2 : [0, 1]^{n-2} \rightarrow [0, 1]$ dos funciones simétricas crecientes. Para todo $0 < k < 1$, se verifica que:

$$\Gamma_{kF_1 + (1-k)F_2} = k\Gamma_{F_1} + (1-k)\Gamma_{F_2}.$$

Como un operador OWA operator es una combinación convexa de estadísticos de orden, tenemos que

Corolario 1: Sea $n \geq 3$, $\mathbf{w} = (w_1, \dots, w_{n-2}) \in [0, 1]^{n-2}$ un vector de pesos y $F_{\mathbf{w}} : [0, 1]^{n-2} \rightarrow [0, 1]$ el operador OWA generado por $\mathbf{w}$. Entonces:

$$\Gamma_{F_{\mathbf{w}}}(x_1, \dots, x_n) = \sum_{k=1}^{n-2} w_k \Gamma_{Q_k}(x_1, \dots, x_n).$$

Sobre funciones de implicación basadas en órdenes admisibles en el conjunto de números borrosos discretos

M. González-Hidalgo ^{*†§}, S. Massanet ^{*†§}, A. Mir ^{*†§}, J. V. Riera ^{*†§}, L. De Miguel [¶]

^{*} Grupo de investigación en Soft Computing, Procesamiento de Imágenes y Agregación (SCOPIA)

Departamento de Ciencias Matemáticas e Informática

Universitat de les Illes Balears, 07122 Palma, España

[†] Instituto de Investigación Sanitaria de las Islas Baleares (IdISBa), 07010 Palma, España

[§] Instituto de Investigación en Inteligencia Artificial (IAIB), 07122 Palma, España

[¶] Departamento de Estadística, Informática y Matemáticas,

Universidad Pública de Navarra, Campus Arrosadia, 31006 Pamplona, Spain

E-mails: {manuel.gonzalez, s.massanet, arnau.mir, jvicente.riera}@uib.es, laura.demiguel@unavarra.es

Resumen—Este trabajo es un resumen del artículo realizado por González-Hidalgo, M., Massanet, S., Mir, A. Riera, J. Vicente y De Miguel, L. titulado "On Fuzzy Implication Functions Based on Admissible Orders on the Set of Discrete Fuzzy Numbers" y publicado en la revista *International Journal of Computational Intelligence Systems* 18, 130 (2025), <https://doi.org/10.1007/s44196-025-00874-9> para su presentación en el XXIII Congreso Español de Tecnologías y Lógica Fuzzy (ESTYLF 2025) KeyWorks.

Index Terms—Número borroso discreto, Funciones de implicación, Orden Total, Orden admisible, Ranking de números borrosos.

I. RESUMEN

En el campo de la lógica borrosa, el estudio de las funciones de implicación ha sido un tema de gran interés para los investigadores, no solo porque generalizan las implicaciones de la lógica booleana clásica, sino también por sus interesantes aplicaciones, como el razonamiento aproximado, el control difuso, la morfología matemática difusa o el procesamiento de imágenes, entre otras [1]. Desde una perspectiva teórica, se han propuesto diversas líneas de investigación, como por ejemplo, el desarrollo de nuevos métodos de construcción [1] o definir estos operadores en otros dominios de definición, considerando en lugar del intervalo unidad, cadenas finitas [2] o retículos acotados [3].

En el caso de un retículo acotado, el orden natural heredado es parcial y dos elementos del retículo no siempre son comparables. Esto puede ser un gran inconveniente, en especial en algunas aplicaciones como en problemas de toma de decisión. Los órdenes admisibles, entendidos como órdenes lineales o totales que generalizan el orden clásico, fueron introducidos por Bustince et al. [4] como una solución a este problema. Esta familia de órdenes se ha investigado en diferentes marcos, como es el caso de los conjuntos borrosos dubitativos, el conjunto de números borrosos, el conjunto de Z -números discretos [5] o el conjunto de números borrosos discretos [6]. Los números borrosos discretos y, en particular,

aquellos cuyo soporte es un subintervalo cerrado de una cadena finita $L_n = \{0, 1, \dots, n\}$, han sido ampliamente investigados en la literatura, ya que constituyen la base teórica del modelo lingüístico basado en números borrosos discretos [7]. En este contexto y considerando un orden parcial, se construyeron algunas familias de funciones de agregación y funciones de implicación basadas en operadores discretos análogos definidos en una cadena finita (para más detalles, por ejemplo, véase [8], [9]). En este artículo, nos basamos en las distintas familias de órdenes totales estudiados en [7] para presentar un novedoso método de construcción de funciones de implicación en el retículo totalmente ordenado y finito $\mathcal{L}' = (\mathcal{D}_1^{L_n \rightarrow Y_m}, \preceq_{\Delta_\delta^\dagger}, \mathbf{0}, \mathbf{1})$ (de números borrosos discretos cuyo soporte es un subintervalo cerrado de la cadena finita L_n y cuyos grados de pertenencia pertenecen a un subconjunto finito $Y_m = \{y_1 = 0, y_2, \dots, y_m = 1\}$ del intervalo $[0, 1]$), a partir de funciones de implicación definidas sobre la cadena L_k donde $k = |\mathcal{D}_1^{L_n \rightarrow Y_m}| - 1$. En concreto,

Teorema 1: Sea I una función de implicación sobre $L_k = \{0, \dots, k\}$ siendo $k = |\mathcal{D}_1^{L_n \rightarrow Y_m}| - 1$. La función

$$\begin{aligned} \mathcal{I} : \mathcal{D}_1^{L_n \rightarrow Y_m} \times \mathcal{D}_1^{L_n \rightarrow Y_m} &\longrightarrow \mathcal{D}_1^{L_n \rightarrow Y_m} \\ (A, B) &\mapsto \mathcal{I}(A, B), \end{aligned}$$

definida como $\mathcal{I}(A, B) = pos^{-1}(I(pos(A), pos(B)))$ es una función de implicación definida sobre el retículo $\mathcal{L}' = (\mathcal{D}_1^{L_n \rightarrow Y_m}, \preceq_{\Delta_\delta^\dagger}, \mathbf{0}, \mathbf{1})$ donde la función de posición pos se define como

$$\begin{aligned} pos : \mathcal{D}_1^{L_n \rightarrow Y_m} &\longrightarrow L_k \\ A &\mapsto pos(A), \end{aligned}$$

(siendo $pos(A) = |\{X \in \mathcal{D}_1^{L_n \rightarrow Y_m} \mid X \preceq_{\Delta_\delta^\dagger} A\}| - 1$ y donde $|\{X \in \mathcal{D}_1^{L_n \rightarrow Y_m} \mid X \preceq_{\Delta_\delta^\dagger} A\}|$ representa el número de elementos del conjunto $\mathcal{D}_1^{L_n \rightarrow Y_m}$ que son menores o iguales que A de acuerdo al orden total(admissible) $\preceq_{\Delta_\delta^\dagger}$).

Además, se demuestra que existe una correspondencia biyectiva entre el conjunto de funciones de implicación en el subconjunto mencionado anteriormente de números difusos discretos y el conjunto de funciones de implicación discretas definidas en la cadena discreta tal como se ve en el siguiente resultado:

Teorema 2: Consideremos $\mathcal{I}_{\mathcal{L}'}$ el conjunto de funciones de implicación definidas sobre el retículo $\mathcal{L}' = (\mathcal{D}_1^{L_n \rightarrow Y_m}, \preceq_{\Delta_\delta^\downarrow}, \mathbf{0}, \mathbf{1})$ y $\mathcal{I}_{L_k}$ el conjunto de funciones de implicación definidas sobre la cadena $L = (L_k, \leq, 0, k)$ siendo $k = |\mathcal{D}_1^{L_n \rightarrow Y_m}| - 1$.

Sea $\mathcal{H}$ la función

$$\begin{aligned} \mathcal{H}: \mathcal{I}_{\mathcal{L}'} &\longrightarrow \mathcal{I}_{L_k} \\ \mathcal{I} &\mapsto \mathcal{H}(\mathcal{I}) = I, \end{aligned}$$

donde I es una función de implicación definida como $I(u, v) = \text{pos}(\mathcal{I}(\text{pos}^{-1}(u), \text{pos}^{-1}(v)))$ para toda pareja $(u, v) \in L_k \times L_k$.

Entonces, $\mathcal{H}$ es una biyección cuya inversa es la función

$$\begin{aligned} \mathcal{H}^{-1}: \mathcal{I}_{L_k} &\longrightarrow \mathcal{I}_{\mathcal{L}'} \\ I &\mapsto \mathcal{H}^{-1}(I) = \mathcal{I}, \end{aligned}$$

donde $\mathcal{I}$ es la función de implicación construida siguiendo el procedimiento descrito en el teorema 1.

En este artículo se investigan a fondo las propiedades básicas de estas funciones de implicación, concluyendo que se conservan bajo el método de construcción propuesto. En particular, se tiene el siguiente resultado.

Teorema 3: Sea $\mathcal{I}_{\mathcal{L}'}$ el conjunto de funciones de implicación definidas sobre el retículo $\mathcal{L}' = (\mathcal{D}_1^{L_n \rightarrow Y_m}, \preceq_{\Delta_\delta^\downarrow}, \mathbf{0}, \mathbf{1})$ y sea $\mathcal{I}_{L_k}$ el conjunto de funciones de implicación definidas sobre la cadena finita $L = (L_k, \leq, 0, k)$ siendo $k = |\mathcal{D}_1^{L_n \rightarrow Y_m}| - 1$.

Sea $\mathcal{H}$ la función

$$\begin{aligned} \mathcal{H}: \mathcal{I}_{\mathcal{L}'} &\longrightarrow \mathcal{I}_{L_k} \\ \mathcal{I} &\mapsto \mathcal{H}(\mathcal{I}) = I, \end{aligned}$$

definida en el teorema 1.

- (i) Sea $I \in \mathcal{I}_{L_k}$ una función de implicación definida sobre el retículo L que verifica las propiedades más habituales (por ejemplo, el principio de neutralidad por la izquierda (**NP**), el principio de intercambio (**EP**), el principio de identidad (**IP**), el principio de ordenación (**OP**), la contrapositividad con respecto de una negación N (**CP(N)**)). Entonces $\mathcal{I} = \mathcal{H}^{-1}(I)$ es una función de implicación en el retículo $\mathcal{L}'$ que verifica las mismas propiedades. En este caso, si I verifica la contraposición con respecto a una negación discreta N definida sobre L_k , entonces $\mathcal{I}$ verificará la misma propiedad considerando la función de negación $\mathcal{N}$ definida sobre el retículo $\mathcal{L}'$ de la siguiente forma $\mathcal{N}(A) = \text{pos}^{-1}(N(\text{pos}(A)))$, $A \in \mathcal{L}'$.
- (ii) Sea $\mathcal{I}$ una función de implicación sobre el retículo $\mathcal{L}'$ que verifica las propiedades más habituales (por ejemplo, el principio de neutralidad por la izquierda (**NP**), el principio de intercambio (**EP**), el principio de identidad (**IP**), el principio de ordenación (**OP**), la contrapositividad con respecto de una negación N (**CP(N)**)). Entonces $I = \mathcal{H}(\mathcal{I})$ es una función de implicación

sobre L_k que satisface las mismas propiedades. En este caso, si $\mathcal{I}$ satisface la contraposición con respecto de la negación N definida sobre el retículo $\mathcal{L}'$, la función de implicación I verifica la misma propiedad con respecto a la función de negación N definida en L dada por $N(u) = \text{pos}(\mathcal{N}(\text{pos}^{-1}(u)))$, $u \in L_k$.

En conclusión el método de construcción presentado en este trabajo proporciona una forma sistemática de extender las funciones de implicación discretas a estructuras más complejas, al tiempo que se conservan sus propiedades subyacentes.

AGRADECIMIENTOS

M. González-Hidalgo, S. Massanet, A. Mir y J.V. Riera están parcialmente subvencionados por el proyecto R+D+i PID2024-158971NB-I00-“Aumentando la confianza en la inteligencia artificial en entornos biomédicos: incertidumbre de datos y modelos, y explicabilidad(ToTAIBiS)”, financiado por MICIU/AEI/10.13039/501100011033/FEDER.

L. De Miguel está parcialmente subvencionada por el proyecto R+D+i PID2022-136627NB-I00-“Técnicas de reducción del número de datos numéricos y datos multiatributo heterogéneos para la mejora de su fusión en problemas de inteligencia artificial” financiado por MCIN/AEI/10.13039/501100011033/.

REFERENCIAS

- [1] M. Baczyński and B. Jayaram. Fuzzy Implications, volume 231 of Studies in Fuzziness and Soft Computing. Springer Berlin Heidelberg, 2008.
- [2] Mas, M., Monserrat, M., Torrens, J.: S-implications and R-implications on a finite chain. Kybernetika 40(1), 3–20 (2004).
- [3] Wang, M., Zhang, X., Bustince, H., Fernandez, J.: Construction methods of fuzzy implications on bounded posets. Int. J. Approximate Reasoning 164, 109064 (2024).
- [4] Bustince, H., Fernandez, J., Kolesárová, A., Mesiar, R.: Generation of linear orders for intervals by means of aggregation functions. Fuzzy Sets Syst 220, 69–77 (2013). Theme: Aggregation functions.
- [5] Mir-Fuentes, A., De Miguel, L., Massanet, S., Mir, A., Riera, J.V.: On a total order on the set of z-numbers based on discrete fuzzy numbers. Computational and Applied Mathematics 43, (2024).
- [6] Riera, J.V., Massanet, S., Bustince, H., Fernandez, J.: On admissible orders on the set of discrete fuzzy numbers for application in decision making problems. Mathematics 9(1) (2021).
- [7] Massanet, S., Riera, J.V., Torrens, J., Herrera-Viedma, E.: A new linguistic computational model based on discrete fuzzy numbers for computing with words. Inf. Sci. 258, 277–290 (2014).
- [8] Riera, J.V., Torrens, J.: Aggregation of subjective evaluations based on discrete fuzzy numbers. Fuzzy Sets and Systems 191, 21–40 (2012).
- [9] Riera, J.V., Torrens, J.: Residual implications on the set of discrete fuzzy numbers. Inf. Sci. 247, 131–143 (2013).

Random choice of referring expressions with fuzzy properties based on user preferences

Nicolás Marín

Dept. Computer Science and A.I.
and CITIC-UGR
University of Granada, Spain
Email: nicm@decsai.ugr.es

Gustavo Rivas-Gervilla

Department of Statistics and Operational Research
University of Granada, Spain
Email: griger@ugr.es

Daniel Sánchez

Dept. Computer Science and A.I.
and CITIC-UGR
University of Granada, Spain
Email: daniel@decsai.ugr.es

Abstract—In this paper, we analyze the problem of referring expression generation for groups of objects (plural referring expressions) when the properties that describe the objects are affected by graduality. The proposed method — based on Fuzzy Sets, Formal Concept Analysis, and Representations by Levels — allows to randomly obtain valid referring expressions according to user preferences.

Index Terms—Referring expressions, plural referring expressions, fuzzy properties, formal concept analysis, representations by levels.

I. INTRODUCTION

One of the most relevant tasks in the field of Natural Language Generation [1] is referring expression generation [2]. Simply speaking, referring expressions can be conceived as conjunctions of properties whose aim is referring to an object or a group of objects within a dataset. In the most general case, when the target is a group of objects, they are called *plural* referring expressions. In this paper, we will consider this general case though, particularly, the reference to specific objects is also covered by considering that the target group may contain only one object.

For a machine to refer to a group of objects within a dataset, for example in the context of a dialogue with the user, an appropriate expression must be automatically generated so that the intended group of objects can be univocally identified by the user. Sometimes it is useful to be able to (randomly) vary the expression, for example, in case a previous expression had not met the objective of identifying the target group, or simply to simulate a certain naturalness in the dialogue process.

This problem is complex due to several reasons, including the fact that not all groups of objects in a dataset are referable (i.e. distinguishable from the others) and that, when they are referable, it is usually necessary to choose among different referring expressions that are valid for distinguishing the group of objects in the dataset.

One of the techniques used to determine the set of referable object groups within a dataset is Formal Concept Analysis (FCA) [3]. Within FCA, the most important notion is that of

formal concept, which biunivocally relates object groups and sets of properties that describe them. In [4], the relationship between FCA and the referring expression generation problem is introduced, with results that allow to characterize the referability (distinguishability) of a group of objects together with the space of valid referring expressions to distinguish them.

This problem of generating referring expressions for a group of objects becomes more complicated when the properties that describe the objects are gradual. In this case, the problem of symbol grounding related to the linguistic terms used to express the properties and their underlying domains must also be solved.

Zadeh's Fuzzy Set theory is a tool of proven utility for providing semantics to the linguistic terms used to express those object properties affected by graduality, establishing a way to relate these terms to an underlying measurable domain.

Furthermore, there are different approaches to FCA in the gradual setting generated by these properties (see a review in [5]). Among them, we can find RL-FCA [6], which uses the theory of Representations by Levels (RL) [7] to link the use of fuzzy sets with conventional FCA.

As we will outline in this work, the combined use of these three tools (namely, Fuzzy Sets, Representation by Levels, and FCA) allows us to obtain a space of valid referring expressions for a specific group of objects, which can additionally be nuanced according to user preferences. Once defined this space, an appropriate probability distribution can be defined to guide the automatic selection of a referring expression among those valid and preferred by the user.

II. SPACES OF REFERRING EXPRESSIONS IN A CRISP SETTING

The starting point of FCA is the definition of a formal context that describes the objects of the problem and the properties that characterize them. Formally, a context $\mathbb{C}$ is a triple $\mathbb{C} = (O, P, I)$, where O is a set of objects, P is a set of binary properties that the objects of O can satisfy, and $I \subseteq O \times P$ is such that $(o, p) \in I$ if, and only if, the object o satisfies the property p .

Once the formal context has been defined from the data set in question, the main objective of FCA is the search for what

This publication is part of the Grant PID2021-126363NB-I00 funded by MICIU/AEI/10.13039/501100011033 and by "ERDF A way of making Europe". The authors appear in alphabetical order and have contributed equally in this research work.

are known as *formal concepts*: a formal concept is nothing more than a pair of sets $(A, B) \in \mathcal{P}(O) \times \mathcal{P}(P)$, such that $A = B'$ and $B = A'$, with $A' = \{p \in P \mid \forall a \in A, (a, p) \in I\}$ and $B' = \{o \in O \mid \forall b \in B, (o, b) \in I\}$. A is called the *extension* of the concept, while B is called its *intension*.

The set $FC(\mathbb{C})$ of all formal concepts of a context $\mathbb{C}$ has the form of a lattice regarding the order relation between formal concepts induced by the inclusion between their extensions (similarly, the inclusion between intensions can also be considered).

In [5] it is shown that a set of objects $A \subseteq O$ is referable (distinguishable from the others) with properties of P if, and only if, $\exists B \subseteq P \mid (A, B) \in FC(\mathbb{C})$. Moreover, in that case, B is a referring expression for A . In fact, it is the largest referring expression that can be used to refer to A , in the sense that any other valid referring expression to A will be contained in B . The space RE_A of valid referring expressions to A can be characterized based on the neighboring of (A, B) in $FC(\mathbb{C})$, facilitating its computation from the aforementioned lattice.

III. REFERRING WITH FUZZY PROPERTIES

If any of the properties that characterize the objects in our problem are fuzzy, the membership of a pair (o, p) to I in a context is unclear, and we need tools to manage this graduality.

A way to do this is by considering formal contexts based on the theory of Representations by Levels [5]. A RL- formal context $\mathbb{C}$ allows formal contexts to be associated with levels in $(0, 1]$, considering different tolerance levels in relation to the graduality of the problem. Succinctly, we can consider, for a fuzzy property p , that (o, p) belongs to I in the context $\mathbb{C}_\alpha$ if, and only if, $\mu_p(o) \geq \alpha$, with $\alpha \in (0, 1]$. Each $\mathbb{C}_\alpha$ build in this way leads us to an associated concept lattice $FC(\mathbb{C}_\alpha)$.

Therefore, instead of a single lattice, as we had in the crisp case described in the previous section, the management of graduality in fuzzy properties leads us to an RL-concept lattice. We will call $RE_{A,\alpha}$ to the space of valid reference expressions in $\mathbb{C}_\alpha$ for each $A \subseteq O$. In particular, $RE_{A,\alpha} = \emptyset$ at those levels where A is not referable.

IV. VALID REFERRING EXPRESSIONS PREFERRED BY THE USER

From the obtained RL-concept lattice, it is possible to calculate the fuzzy set $\widetilde{RE}_A$ of *valid referring expressions* for a certain $A \subseteq O$ following the ideas proposed in [7] to obtain the fuzzy summary of any RL-system.

To do this, it is enough to take as value of $\mu_{\widetilde{RE}_A}(re)$ the probability that, taking a random $\alpha \in (0, 1]$, $re \in RE_{A,\alpha}$, for each $re \in \mathcal{P}(P)$.

Additionally, we can consider other interesting fuzzy sets defined over $\mathcal{P}(P)$, whose semantic correspond to different criteria that can be taken into account to express user preferences regarding referring expressions.

Some examples of these criteria can be expression length (for example, to determine the set of *short* expressions or the set of *long* expressions) or criteria based on the salience of the properties, among many others [2].

In general, we can consider $\widetilde{RE}_{C_1}, \dots, \widetilde{RE}_{C_n}$ as the different fuzzy sets defined on $\mathcal{P}(P)$ obtained based on the different criteria that the user employs to express her preference regarding referring expressions.

By considering $\widetilde{RE}_A$ together with $\widetilde{RE}_{C_1}, \dots, \widetilde{RE}_{C_n}$, we can build the fuzzy set of *preferred valid referring expressions* for a certain $A \subseteq O$ as follows:

$$\widetilde{prefRE}_A = \widetilde{RE}_A \otimes \widetilde{RE}_{C_1} \otimes \dots \otimes \widetilde{RE}_{C_n}$$

where $\otimes$ is a t-norm.

V. RANDOM CHOICE OF REFERRING EXPRESSIONS BASED ON USER PREFERENCES

We have commented in the introduction that, when building an automaton capable to interact with the user by means of appropriate referring expressions within a dataset, it may be interesting to randomly select referring expressions from among those preferred by the user.

One way to do this is by defining, for each $A \subseteq O$, a probability distribution $Prob_A$ over $\mathcal{P}(P)$ as follows:

$$Prob_A(re) = \frac{\mu_{\widetilde{prefRE}_A}(re)}{\sum_{re' \in \mathcal{P}(P)} \mu_{\widetilde{prefRE}_A}(re')}$$

This allows the automaton to randomly vary the expression used based on its validity and the user's preferences.

VI. CONCLUSIONS

In this work, we have outlined a way to randomly generate referring expressions to distinguish groups of objects within a dataset with fuzzy properties.

The combined use of fuzzy set theory and representations by levels allows the adequate representation and management of graduality within the framework of formal concept analysis, a useful tool for referring expression generation.

In this work, we have characterized the set of valid referring expressions preferred by the user to distinguish a given group of objects and we have shown how it can be used to induce a probability distribution on the referring expression space that allows for their consistent random selection.

REFERENCES

- [1] A. Gatt and E. Krahmer, "Survey of the state of the art in natural language generation: Core tasks, applications and evaluation," *J. Artif. Intell. Res.*, vol. 61, pp. 65–170, 2018.
- [2] K. van Deemter, *Computational Models of Referring: A Study in Cognitive Science*. The MIT Press, 2016.
- [3] B. Ganter and R. Wille, *Formal Concept Analysis - Mathematical Foundations*. Springer, 1999.
- [4] N. Marín, G. Rivas-Gervilla, M. D. Ruiz, and D. Sánchez, "Formal concept analysis for the generation of plural referring expressions," *Inf. Sci.*, vol. 579, pp. 717–731, 2021.
- [5] G. Rivas-Gervilla, *Mecanismos formales para la representación y extracción de expresiones de referencia en sistemas data-to-text*. Tesis Doctoral, 2022. [Online]. Available: <https://digibug.ugr.es/handle/10481/74613>
- [6] N. Marín, G. Rivas-Gervilla, and D. Sánchez, "RL-FCA: An alternative to fuzzy formal concept analysis based on representations by levels," *Fuzzy Sets and Systems*, in preparation.
- [7] D. Sánchez, M. Delgado, M. Vila, and J. Chamorro-Martínez, "On a non-nested level-based representation of fuzziness," *Fuzzy Sets and Systems*, vol. 192, no. 1, pp. 159–175, 2012.

Estudio de la monotonía de las medidas de similitud basadas en incrustaciones

Sara Suarez-Fernandez	Agustina Bouchet	Susana Diaz-Vazquez	Susana Montes
Dept. Estadística, I.O. y D.M.	Dept. Estadística, I.O. y D.M.	Dept. Estadística, I.O. y D.M.	Dept. Estadística, I.O. y D.M.
Universidad de Oviedo	Universidad de Oviedo	Universidad de Oviedo	Universidad de Oviedo
Oviedo, España	Oviedo, España	Oviedo, España	Oviedo, España
sfernandezsara@uniovi.es	bouchetagustina@uniovi.es	diazsusana@uniovi.es	montes@uniovi.es

Abstract—En este trabajo se estudia la monotonía de las medidas de similitud para intervalos basadas en t-normas según incrustaciones. Comúnmente se asume que dados tres conjuntos ordenados según la relación de inclusión, el grado de similitud entre el mayor y el menor debe ser inferior al valor de la similitud obtenido al comparar el que ocupa la posición central con cualquiera de los otros dos. La debilidad que presenta esta propiedad es que no exige nada cuando se consideran intervalos que no están ordenados. Por ello, se recurre a una nueva condición de monotonía basada en el ínfimo y el supremo que, aunque resulta intuitiva, no se satisface para cualquier similitud. Concretamente, se proporcionan cuatro familias de incrustaciones que definen similitudes que satisfacen dicha condición de monotonía al considerar la relación de inclusión y el orden reticular.

Index Terms—similitud, incrustación, monotonía, intervalos

I. INTRODUCCIÓN

Las medidas de similitud han sido ampliamente estudiadas en la literatura por sus aplicaciones en diversos ámbitos del conocimiento (ver, por ejemplo, [1]–[3]). En particular, las medidas de similitud pueden emplearse para comparar intervalos, que se presentan como una alternativa adecuada a los números reales cuando lo que se pretende es representar información relativa al pensamiento humano.

En este trabajo se señala que, a pesar de que la monotonía respecto a la inclusión es un axioma popularmente aceptado, puede ser demasiado débil puesto que no es posible aplicarlo a cualesquiera tres intervalos. Es por ello que, para estudiar la monotonía, se va a considerar la generalización propuesta en [10]. Esta investigación se centra en explorar qué familias de incrustaciones generan similitudes que satisfacen dicha condición.

El trabajo se organiza como sigue: en la siguiente sección se introducen algunos conceptos fundamentales. En la Sección III se desarrolla el estudio de la monotonía de las similitudes definidas mediante incrustaciones y en la sección IV se recogen las conclusiones del estudio.

II. PRELIMINARES

En esta sección se establece la notación y se introducen aquellos conceptos fundamentales para este estudio.

Los autores expresan su agradecimiento a los proyectos SEK-25-GRU-GIC-24-018, MCINN-23-PID2022-139886NB-I00, MCINN-24-PID2023-146621OB-C21 y SCIMIN-CRM-101177746.

Se considera el conjunto de todos los intervalos cerrados incluidos en el intervalo unidad $[0, 1]$ y se emplea la siguiente notación: $\mathcal{L}([0, 1]) = \{[\underline{a}, \bar{a}] | \underline{a}, \bar{a} \in [0, 1] \text{ y } \underline{a} \leq \bar{a}\}$.

Las medidas de similitud son operadores que permiten cuantificar cuán similares son dos objetos, proporcionando un grado de similitud comprendido entre 0 y 1. A pesar de que en la literatura pueden encontrarse diversas definiciones, en su mayoría presentan dos propiedades comunes (ver [5]–[7]): deben ser simétricas y el valor máximo del grado de similitud debe alcanzarse únicamente cuando los intervalos comparados coinciden. En [5] proponen adicionalmente el siguiente axioma, que se puede entender como una condición de monotonía:

$$a \subseteq b \subseteq c \Rightarrow S(a, c) \leq \min(S(a, b), S(b, c)). \quad (S3)$$

Intuitivamente, este axioma establece que si los intervalos están ordenados, el grado de similitud entre ellos también.

Las funciones de incrustación [8], denotadas por E , miden en qué medida un intervalo está contenido en otro. Estos operadores toman el valor máximo únicamente cuando el primer intervalo está totalmente incluido en el segundo, toman el mínimo cuando los intervalos no tienen ningún punto en común y satisfacen adicionalmente las dos condiciones siguientes:

$$\text{Si } b \subseteq c \text{ entonces } E(a, b) \leq E(a, c) \quad (1)$$

$$\text{Si } a \subseteq b \subseteq c \text{ entonces } E(c, a) \leq E(b, a) \quad (2)$$

En [9] se proporciona una familia de medidas de similitud tomando como punto de partida medidas de incrustación. Dada una t-norma T y una medida de incrustación E , se denomina similitud basada en la t-norma según la incrustación a la función $S_T^E : \mathcal{L}([0, 1])^2 \rightarrow [0, 1]$ definida como sigue :

$$S_T^E(a, b) = T(E(a, b), E(b, a)) \quad (3)$$

III. ESTUDIO DE LA MONOTONÍA

La idea de monotonía para las funciones de similitud expresada a través del axioma (S3) puede generalizarse considerando un orden cualquiera entre intervalos, denotado por $\preceq$. Dados $a, b, c \in \mathcal{L}([0, 1])$:

$$a \preceq b \preceq c \implies S(a, c) \leq S(a, b), S(a, c) \leq S(b, c) \quad (4)$$

La principal debilidad de esta propiedad radica en que, si se trabaja con un orden que no es total, se convierte en una condición vacía para cualquier conjuntos de tres intervalos no ordenados. En [10] se presenta una condición de monotonía que permite abordar esta limitación en el ámbito de los conjuntos difusos intervalo-valorados. Aplicándola para intervalos en $\mathcal{L}([0, 1])$, la idea de monotonía para las similitudes puede expresarse como sigue para cualesquiera $a, b, c \in \mathcal{L}([0, 1])$:

$$S(\inf_{\subseteq} (a, b, c), \sup_{\subseteq} (a, b, c)) \leq S(a, b) \quad (5)$$

Nótese que para garantizar la existencia de ínfimo y supremo es preciso que el orden considerado genere un retículo en $\mathcal{L}([0, 1])$. En el caso del orden reticular, dados $a, b, c \in \mathcal{L}([0, 1])$ el ínfimo y el supremo serán los intervalos siguientes:

$$\inf_{Lo} (a, b, c) = [\min(\underline{a}, \underline{b}, \underline{c}), \min(\bar{a}, \bar{b}, \bar{c})] \quad (6)$$

$$\sup_{Lo} (a, b, c) = [\max(\underline{a}, \underline{b}, \underline{c}), \max(\bar{a}, \bar{b}, \bar{c})] \quad (7)$$

En el caso de la relación de inclusión, se considerarán intervalos tales que $a \cap b \cap c \neq \emptyset$ para garantizar la existencia del ínfimo. Así, $\inf_{\subseteq} (a, b, c) = a \cap b \cap c$, $\sup_{\subseteq} (a, b, c) = a \cup b \cup c$.

Veamos que pueden definirse familias de incrustaciones que permiten obtener similitudes que verifican (5) para el orden reticular y la inclusión. Las implicaciones $I : [0, 1]^2 \rightarrow [0, 1]$ consideradas de ahora en adelante verifican que:

$$I(x, y) = 1 \iff x \leq y \quad (8)$$

Proposición 3.1 Sea T una t-norma e I una implicación verificando (8). El operador $E_I : \mathcal{L}([0, 1])^2 \rightarrow [0, 1]$ dado por $E_I(a, b) = 0$ si $a \cap b = \emptyset$ y $E_I(a, b) = \min(I(\underline{b}, \underline{a}), I(\bar{a}, \bar{b}))$ en otro caso, define una medida de incrustación tal que:

$$S_T^{E_I}(\inf_{Lo} (a, b, c), \sup_{Lo} (a, b, c)) \leq S_T^{E_I}(a, b) \quad (9)$$

$$S_{\min}^{E_I}(\inf_{\subseteq} (a, b, c), \sup_{\subseteq} (a, b, c)) \leq S_{\min}^{E_I}(a, b) \quad (10)$$

Algo similar ocurre al medir el grado de inclusión a partir de otras incrustaciones y se muestra en las siguientes proposiciones:

Proposición 3.2. Sea T una t-norma y $\alpha, \beta \geq 0$ tales que $\alpha + \beta \leq 1$. El operador $E_{\phi_{\alpha\beta}} : \mathcal{L}([0, 1])^2 \rightarrow [0, 1]$ dado por $E_{\phi_{\alpha\beta}}(a, b) = 1$ si $a = b$, $E_{\phi_{\alpha\beta}}(a, b) = 0$ si $a \cap b = \emptyset$ o si $w(a \cap b) = 0$ con $a \not\subseteq b$ y $E_{\phi_{\alpha\beta}}(a, b) = \alpha + \beta w(a \cap b)$ en otro caso, define una medida de incrustación tal que:

$$S_T^{E_{\phi_{\alpha\beta}}}(\inf_{Lo} (a, b, c), \sup_{Lo} (a, b, c)) \leq S_T^{E_{\phi_{\alpha\beta}}}(a, b) \quad (11)$$

$$S_{\min}^{E_{\phi_{\alpha\beta}}}(\inf_{\subseteq} (a, b, c), \sup_{\subseteq} (a, b, c)) \leq S_{\min}^{E_{\phi_{\alpha\beta}}}(a, b) \quad (12)$$

Proposición 3.3 Sea I una implicación satisfaciendo (8). El operador $E_{\phi_I} : \mathcal{L}([0, 1])^2 \rightarrow [0, 1]$ dado por $E_{\phi_I}(a, b) = 0$ si $a \cap b = \emptyset$ y $E_{\phi_I}(a, b) = I(w(a), w(a \cap b))$ en otro caso, define una medida de incrustación tal que:

$$S_{\min}^{E_{\phi_I}}(\inf_{Lo} (a, b, c), \sup_{Lo} (a, b, c)) \leq S_{\min}^{E_{\phi_I}}(a, b) \quad (13)$$

$$S_{\min}^{E_{\phi_I}}(\inf_{\subseteq} (a, b, c), \sup_{\subseteq} (a, b, c)) \leq S_{\min}^{E_{\phi_I}}(a, b) \quad (14)$$

Un caso particular de la familia E_{ϕ_I} que permite considerar una t-norma arbitraria viene caracterizado por la siguiente función de implicación:

$$I_{T_{\min}}(x, y) = \sup\{z \in [0, 1] \mid \min(x, z) \leq y\} \quad (15)$$

Proposición 3.4 Sea T una t-norma. La similitud según T basada en el operador $E_{\phi_{I_{T_{\min}}}}$ cumple que:

$$S_T^{E_{\phi_{I_{T_{\min}}}}}(\inf_{Lo} (a, b, c), \sup_{Lo} (a, b, c)) \leq S_T^{E_{\phi_{I_{T_{\min}}}}}(a, b) \quad (16)$$

$$S_T^{E_{\phi_{I_{T_{\min}}}}}(\inf_{\subseteq} (a, b, c), \sup_{\subseteq} (a, b, c)) \leq S_T^{E_{\phi_{I_{T_{\min}}}}}(a, b) \quad (17)$$

En aquellos casos en los que se ha recurrido a la t-norma del mínimo, puede probarse que la propiedad no es cierta al considerar una t-norma cualquiera.

IV. CONCLUSIONES

En este trabajo se ha explorado la monotonía de algunas familias de similitudes obtenidas a partir de medidas de incrustación. Particularmente, se han proporcionado cuatro familias de incrustaciones tales que, al considerar el orden reticular o la inclusión, el valor de similitud obtenido al comparar el ínfimo y el supremo de tres intervalos cualesquiera es menor que el grado de similitud entre los pares de intervalos implicados. A pesar de que la condición de monotonía empleada resulta coherente e intuitiva, no se verifica para cualquier función de similitud entre intervalos. Es por ello que, en algunos casos ha sido preciso restringirse a la t-norma del mínimo para garantizar la propiedad de monotonía.

Como trabajo futuro, nos proponemos estudiar esta nueva propiedad de monotonía en un contexto más general, en particular, para familias más generales de medidas de incrustación.

REFERENCIAS

- [1] Basseville, M.: Distance measures for signal processing and pattern recognition. *Signal Processing* 18(4), 349–369 (1989).
- [2] Bustince, H., Marco-Detchart, C., Fernandez, J., Wagner, C., Garibaldi, J., Takáč, Z.: Similarity between interval-valued fuzzy sets taking into account the width of the intervals and admissible orders. *Fuzzy Sets and Systems* 390, 23–47 (2020).
- [3] Dice, L. R. (1945). Measures of the amount of ecologic association between species. *Ecology*, 26(3), 297–302.
- [4] Torres-Manzanera, E., Bouchet, A., Díaz-Vazquez, S., Montes, S.: On inclusion and order relation of interval-valued fuzzy sets. In: *Actas del XXI Congreso Español sobre Tecnologías y Lógica Fuzzy* (2022)
- [5] Kabir, S., Wagner, C., Havens, T.C., Anderson, D.T.: A similarity measure based on bidirectional subethood for intervals. *IEEE Transactions on Fuzzy Systems* 28(11), 2890–2904 (2020)
- [6] McCulloch, J., Wagner, C. and Aickelin, U.: Analysing fuzzy sets through combining measures of similarity and distance. in *Proc. IEEE Int. Conf. Fuzzy Syst.*, 2014, pp. 155–162.
- [7] Ontañón, S. An overview of distance and similarity functions for structured data. *Artificial Intelligence Review*, 53(7), 5309–5351. (2020)
- [8] Bouchet, A., Sesma-Sara, M., Ochoa, G., Bustince, H., Montes, S., Díaz, I.: Measures of embedding for interval-valued fuzzy sets. *Fuzzy Sets and Systems* 467, 108505 (2023).
- [9] Bouchet, A., Díaz-Vazquez, S., Díaz, I., Montes, S.: Similitudes basadas en medidas de incrustación. In: *Proceedings of the XX Conferencia de la Asociación Española para la Inteligencia Artificial* (2024)
- [10] Goguen, J. A. (1967). L-fuzzy sets. *Journal of mathematical analysis and applications*, 18(1), 145–174.

Convexidad en conjuntos difusos *typical hesitant* aplicada a la toma de decisiones

1º Pedro Huidobro
Dept. de Estadística e I.O.
Universidad de Oviedo
Oviedo, España
huidobropedro@uniovi.es

2º Pedro Alonso
Dept. de Matemáticas
Universidad de Oviedo
Oviedo, España
palonso@uniovi.es

3º Vladimir Janiš
Dept. of Mathematics
Matej Bel University
Banska Bystrica, Eslovaquia
vladimir.janis@umb.sk

4ª Susana Montes
Dept. de Estadística e I.O.
Universidad de Oviedo
Oviedo, España
montes@uniovi.es

Abstract—Los conjuntos difusos *typical hesitant* (THF) permiten representar la incertidumbre asociada al hecho de que diferentes expertos asignen múltiples grados de pertenencia a un mismo elemento. Esta flexibilidad los hace especialmente adecuados para problemas de decisión, pero también plantea desafíos en la comparación de alternativas y la preservación de propiedades matemáticas fundamentales. En este trabajo proponemos un marco de decisión basado en una nueva noción de convexidad para conjuntos THF, con el objetivo de garantizar rigor en los procesos de optimización y toma de decisiones. Para ello, desarrollamos operaciones de intersección, unión y α -cortes compatibles con esta nueva convexidad, y demostramos cómo aplicarlas en problemas reales de decisión. La propuesta, publicada en [5], permite trabajar con restricciones imprecisas, respetando la estructura original de los datos.

Index Terms—Conjuntos difusos *typical hesitant*, convexidad, toma de decisiones, orden admisible.

I. PRELIMINARES Y MOTIVACIÓN

En determinados contextos de decisión, es habitual que distintos expertos proporcionen valoraciones imprecisas sobre las alternativas disponibles. Los conjuntos difusos introducidos por Zadeh son una herramienta útil para trabajar con esa imprecisión inherente a los datos [3], [14]. Una de sus extensiones, los conjuntos *typical hesitant* (THF) [1], [10], permite asignar a cada elemento un conjunto finito y no vacío de grados de pertenencia, capturando la diversidad de opiniones en situaciones de incertidumbre. Formalmente:

Definición 1.1: [1] Sea $\mathbb{H} = \{S \subseteq [0, 1] : S \text{ es finito y } S \neq \emptyset\}$. Un conjunto difuso *typical hesitant* A sobre un universo X se define como $A = \{\langle x, h_A(x) \rangle : x \in X\}$, donde $h_A : X \rightarrow \mathbb{H}$.

Denotamos el cardinal del elemento THF $h_A(x)$ como $\#h_A(x)$. Sin embargo, la comparación entre elementos *hesitant* plantea dificultades cuando los conjuntos tienen distinta cardinalidad. Para dar solución a este problema se han planteado distintas alternativas, como por ejemplo, funciones de *score* [4], [12] y normalizaciones [9], pero muchas de estas soluciones pierden información.

P. Huidobro, P. Alonso and S. Montes are supported by the Spanish Ministry of Science and Innovation projects PID2022- 139886NB-I00 and by the FICYT project SEK-25-GRU-GIC-24-018. S. Montes is also supported by the European Project UE-23-AI4RA-101132914. V. Janiš thanks for the support to the grant 1/0124/24 by Slovak grant agency VEGA.

En este trabajo se propone una alternativa consistente en utilizar relaciones de orden total sobre $\mathbb{H}$. Wang y Xu [11] formalizaron el concepto de *orden admisible* en $\mathbb{H}$, y Matzenauer et al. [8] lo ampliaron y dieron un método constructivo.

Definición 1.2: [8] Un orden $\trianglelefteq$ sobre $\mathbb{H}$ es *admisible* si refina el orden parcial $\trianglelefteq_{RH}$ definido como:

$$\begin{aligned} h_A(x) \trianglelefteq_{RH} h_B(x) \Leftrightarrow & (h_A(x) = \mathbf{0}_{\mathbb{H}}) \\ & \text{o } (h_B(x) = \mathbf{1}_{\mathbb{H}}) \\ & \text{o } (\#h_A(x) = \#h_B(x) = n \\ & \text{y } h_A^{(i)} \leq h_B^{(i)}, \forall i \in \mathbb{N}_n) \end{aligned}$$

Teorema 1.1: [8] Las relaciones $\trianglelefteq_{Lex1}$ y $\trianglelefteq_{Lex2}$ sobre $\mathbb{H}$ se definen, respectivamente, de la siguiente manera, para cualesquiera elementos difusos *typical hesitant* $h_A(x)$ y $h_B(x)$ con $m = \#h_A(x)$ y $n = \#h_B(x)$:

- $h_A(x) \trianglelefteq_{Lex1} h_B(x)$ si y sólo si:

$$\begin{aligned} \exists i \in \mathbb{N}_{\min\{m,n\}} : & h_A(x)^i < h_B(x)^i \text{ y} \\ & h_A(x)^j = h_B(x)^j, \forall j < i, \\ & \text{o bien } m \leq n \text{ y } h_A(x)^j = h_B(x)^j, \forall j \in \mathbb{N}_m. \end{aligned}$$

- $h_A(x) \trianglelefteq_{Lex2} h_B(x)$ si y sólo si:

$$\begin{aligned} \exists i \in \mathbb{N}_{\min\{m,n\}} : & h_A(x)^i < h_B(x)^i \text{ y} \\ & h_A(x)^j = h_B(x)^j, \forall j > i, \\ & \text{o bien } m \leq n \text{ y } h_A(x)^j = h_B(x)^j, \forall j \in \mathbb{N}_m. \end{aligned}$$

Ambas son órdenes admisibles.

II. MARCO PROPUESTO PARA LA TOMA DE DECISIONES

Antes de abordar el modelo de toma de decisiones, presentamos dos conceptos clave para ello, estos son la intersección y la convexidad, así como la relación que existe entre ellos.

A. Intersección y convexidad

En esta subsección, proponemos redefinir la operación de intersección entre conjuntos THF:

Definición 2.1: Sean A y B dos conjuntos difusos *typical hesitant* sobre X , y $\trianglelefteq_o$ un orden total sobre $\mathbb{H}$. La o -intersección de A y B se define como:

$$(A \cap_o B)(x) = \min_{\trianglelefteq_o} \{h_A(x), h_B(x)\}$$

Esta definición generaliza la intersección clásica y permite preservar propiedades estructurales clave como la convexidad. Inspirándonos en convexidades previamente definidas [2], [6] y [7], introducimos la siguiente noción:

Definición 2.2: Sea X un conjunto ordenado y $\trianglelefteq_o$ una relación de orden sobre $\mathbb{H}$. Un conjunto THF A es o -convexo si para cualesquiera $x < y < z$ en X se cumple:

$$h_A(x) \trianglelefteq_o h_A(y) \quad \text{o} \quad h_A(z) \trianglelefteq_o h_A(y)$$

Cuando no se dé la igualdad, diremos estrictamente o -convexo.

Proposición 2.1: Sean A y B conjuntos THF o -convexos sobre X . Si $A \cap_o B \neq \emptyset$, entonces $A \cap_o B$ también es o -convexo.

B. Toma de decisiones

Inspirándonos en el enfoque clásico de Bellman y Zadeh [2], así como en la extensión condicional propuesta por Yager y Basson [13], proponemos un modelo de decisión en el que tanto los objetivos como las restricciones se expresan mediante conjuntos difusos *typical hesitant*.

Definición 2.3: Sea $X = \{x_1, \dots, x_n\}$ el conjunto de alternativas. Sean $G_1, \dots, G_p$ los objetivos y $C_1, \dots, C_m$ las restricciones, todos ellos definidos como conjuntos THF sobre X . Dados estos elementos y un orden total $\trianglelefteq_o$ sobre $\mathbb{H}$, la *decisión resultante* se define como:

$$D_o = G_1 \cap_o \dots \cap_o G_p \cap_o C_1 \cap_o \dots \cap_o C_m$$

Este modelo permite evitar la agregación forzada de las valoraciones *hesitant*, manteniendo la información aportada por distintos expertos. Además, el uso de órdenes admisibles garantiza que las operaciones realizadas preservan la estructura interna de los conjuntos y sus propiedades.

El marco puede extenderse a restricciones formuladas en otros espacios (utilizando el principio de extensión de Zadeh [14]) o a restricciones condicionales mediante conjuntos difusos condicionales como los que propone Yager [13].

A continuación, presentamos un resultado clave sobre la preservación de la convexidad y sus implicaciones en la unicidad del conjunto de decisiones óptimas:

Proposición 2.2: Sea X un conjunto ordenado, $\trianglelefteq_o$ un orden total sobre $\mathbb{H}$, y D_o el conjunto decisión definido como la intersección de los objetivos G_i y restricciones C_j . Entonces:

- Si todos los G_i y C_j son conjuntos THF o -convexos, entonces D_o es o -convexo, y el conjunto de alternativas que maximizan D_o es un conjunto convexo clásico.
- Si todos los G_i y C_j son conjuntos THF estrictamente o -convexos, entonces D_o es estrictamente o -convexo y el número de alternativas óptimas es a lo sumo dos.

III. CONCLUSIONES

Este trabajo desarrolla un marco teórico coherente y sólido para abordar problemas de toma de decisiones en entornos inciertos, modelados mediante conjuntos difusos *typical hesitant*. A partir de la noción de órdenes admisibles sobre

el conjunto de elementos *hesitant*, hemos definido operaciones fundamentales como la intersección y la convexidad, respetando las nociones clásicas.

La redefinición de la intersección a través de un orden total garantiza una mayor consistencia interna, evitando la pérdida de información causada por agregaciones o normalizaciones. Gracias a ello, se logra integrar de forma natural las preferencias múltiples o imprecisas de varios expertos, sin comprometer la lógica del modelo.

Además, se ha demostrado que la convexidad se preserva en este nuevo contexto y que su aplicación a problemas de decisión permite identificar soluciones óptimas con propiedades deseables, como la unicidad.

Este marco no solo aporta fundamentos matemáticos rigurosos, sino que también sienta las bases para futuras aplicaciones en escenarios donde la incertidumbre y la imprecisión son intrínsecas, abriendo nuevas vías para el diseño de sistemas de decisión más robustos.

REFERENCES

- [1] B. Bedregal, R. Reiser, H. Bustince, C. Lopez-Molina, and V. Torra, "Aggregation functions for typical hesitant fuzzy elements and the action of automorphisms," *Information Sciences*, vol. 255, pp. 82–99, 2014.
- [2] R. E. Bellman and L. A. Zadeh, "Decision-making in a fuzzy environment," *Management Science*, vol. 17, no. 4, pp. B-141–B-164, 1970.
- [3] E. Czogała and H. J. Zimmermann, "Decision making in uncertain environments," *European Journal of Operational Research*, vol. 23, no. 2, pp. 202–212, 1986.
- [4] B. Farhadinia, "A series of score functions for hesitant fuzzy sets," *Information Sciences*, vol. 277, pp. 102–110, 2014.
- [5] P. Huidobro, P. Alonso, V. Janiš, and S. Montes, "Convexity of typical hesitant fuzzy sets applied to decision-making problems," *Computational and Applied Mathematics*, vol. 44, 304, 2025.
- [6] P. Huidobro, P. Alonso, V. Janiš, and S. Montes, "Convexity of hesitant fuzzy sets based on aggregation functions," *Computer Science and Information Systems*, vol. 18, no. 1, pp. 213–230, 2021.
- [7] G. Klir and B. Yuan, *Fuzzy Sets and Fuzzy Logic: Theory and Applications*, Prentice Hall, 1995.
- [8] M. Matzenauer, R. Reiser, H. Santos, B. Bedregal, and H. Bustince, "Strategies on admissible total orders over typical hesitant fuzzy implications applied to decision making problems," *International Journal of Intelligent Systems*, vol. 36, no. 5, pp. 2144–2182, 2021.
- [9] T. Rashid and I. Beg, "Convex hesitant fuzzy sets," *Journal of Intelligent and Fuzzy Systems*, vol. 30, no. 5, pp. 2791–2796, 2016.
- [10] V. Torra, "Hesitant fuzzy sets," *International Journal of Intelligent Systems*, vol. 25, no. 6, pp. 529–539, 2010.
- [11] H. Wang and Z. Xu, "Admissible orders of typical hesitant fuzzy elements and their application in ordered information fusion in multi-criteria decision making," *Information Fusion*, vol. 29, pp. 98–104, 2016.
- [12] M. Xia and Z. Xu, "Distance and similarity measures for hesitant fuzzy sets," *Information Sciences*, vol. 181, no. 11, pp. 2128–2138, 2011.
- [13] R. Yager and D. Basson, "Decision making with fuzzy sets," *Decision Sciences*, vol. 6, no. 3, pp. 590–600, 1975.
- [14] L. A. Zadeh, "Fuzzy sets," *Information and Control*, vol. 8, no. 3, pp. 338–353, 1965.

Eficiencia normalizada de conjuntos de reglas de decisión

Fernando Chacón-Gómez
Departamento de matemáticas
Universidad de Cádiz
Puerto Real, España
fernando.chacon@uca.es

M. Eugenia Cornejo
Departamento de matemáticas
Universidad de Cádiz
Puerto Real, España
mariaeugenia.cornejo@uca.es

Jesús Medina
Departamento de matemáticas
Universidad de Cádiz
Puerto Real, España
jesus.medina@uca.es

Resumen—En el marco de la teoría de conjuntos rugosos (RST, del inglés Rough Set Theory), las bases de datos relacionales son descritas por medio de reglas de decisión. Estas descripciones son analizadas por la noción de eficiencia, que cuantifica la proporción de objetos clasificados correctamente por dichas reglas. En este trabajo, se introduce una definición de eficiencia normalizada en el entorno de la teoría de conjuntos rugosos difusos (FRST, del inglés Fuzzy Rough Set Theory), y se muestran distintas propiedades de esta noción.

Palabras clave—Teoría de conjuntos rugosos difusos, reglas de decisión, eficiencia.

I. INTRODUCCIÓN

Pawlak introdujo RST [5], [6] como una herramienta formal para el manejo, en términos matemáticos, de información incompleta contenida en conjuntos de datos relacionales. Posteriormente, este marco de trabajo se extendió al entorno difuso, mediante la introducción de FRST [3], [4], permitiendo abordar estudios de mayor complejidad. En ambos enfoques, los conjuntos de datos se representan mediante tablas de decisión, que están formadas por un conjunto de objetos y los atributos que los describen. A partir de estas tablas, se definen reglas de decisión para facilitar la extracción de conocimiento. Los conjuntos de reglas seleccionados son analizados por la noción de eficiencia [7], introducida por Pawlak en RST, que muestra la capacidad de estas reglas para clasificar correctamente los objetos presentes en la tabla. Este trabajo introduce una definición de eficiencia en el marco de FRST, con el objetivo de analizar conjuntos de reglas de decisión desde la perspectiva difusa. A diferencia de propuestas anteriores, como en [1], esta noción solo toma valores en el intervalo unidad, lo que facilita su interpretación.

II. TEORÍA DE CONJUNTOS RUGOSOS DIFUSOS

Esta sección recoge nociones preliminares de FRST [2]. En primer lugar, se introducen las tablas de decisión, que representan los conjuntos de datos relacionales en términos matemáticos.

Parcialmente subvencionado por el proyecto PID2022-137620NB-I00 financiado por MICIU/AEI/10.13039/501100011033 y por FEDER, UE y por el proyecto TED2021-129748B-I00 financiado por MICIU/AEI/10.13039/501100011033 y por la Unión Europea NextGenerationEU/PRTR.

Definición 1. Sean U y $\mathcal{A}$ conjuntos no vacíos de objetos y atributos, respectivamente. Una tabla de decisión es una tupla $S = (U, \mathcal{A}_d, \mathcal{V}_{\mathcal{A}_d}, \overline{\mathcal{A}_d})$ con $\mathcal{A}_d = \mathcal{A} \cup \{d\}$ y $d \notin \mathcal{A}$, $\mathcal{V}_{\mathcal{A}_d} = \{V_a \mid a \in \mathcal{A}_d\}$, donde V_a es el conjunto de valores asociados al atributo a sobre U , y $\overline{\mathcal{A}_d} = \{\bar{a}: U \rightarrow V_a \mid a \in \mathcal{A}_d\}$.

La información almacenada en las tablas de decisión puede ser expresada en términos lógicos mediante la noción de fórmula. En la siguiente definición se presentan varios conceptos relacionados con las fórmulas.

Definición 2. Sean $S = (U, \mathcal{A}_d, \mathcal{V}_{\mathcal{A}_d}, \overline{\mathcal{A}_d})$ una tabla de decisión, $T = \{T_a: V_a \times V_a \rightarrow [0, 1] \mid a \in \mathcal{A}_d\}$ una familia de relaciones de tolerancia difusas separables y $B \subseteq \mathcal{A}_d$.

- El conjunto de fórmulas asociadas a B , denotado como $For(B)$, se construye a partir de pares atributo valor (a, v) , donde $a \in B$ y $v \in V_a$, por medio del conectivo lógico de conjunción $\wedge$.
- La aplicación $\|\cdot\|_S^T: For(B) \rightarrow [0, 1]^U$ se define como:

$$\|\Phi\|_S^T(x) = T_a(\bar{a}(x), v)$$

para todo $x \in U$ y $\Phi = (a, v) \in For(B)$. Por tanto, $\|\Phi\|_S^T(x)$ es el grado de satisfacibilidad del objeto x con respecto a la fórmula Φ .

- Para cada $\Phi, \Psi \in For(B)$, el grado de satisfacibilidad de la conjunción $\Phi \wedge \Psi$ se define, para todo $x \in U$, como:

$$\|\Phi \wedge \Psi\|_S^T(x) = \min\{\|\Phi\|_S^T(x), \|\Psi\|_S^T(x)\}$$

- Dadas $\Phi, \Phi' \in For(\mathcal{A}_d)$, tales que:

$$\begin{aligned} \Phi &= (a_1, v_1) \wedge \dots \wedge (a_n, v_n) \\ \Phi' &= (a'_1, w_1) \wedge \dots \wedge (a'_m, w_m) \end{aligned}$$

la relación de F -indiscernibilidad $R_{Fd}: For(\mathcal{A}_d) \times For(\mathcal{A}_d) \rightarrow [0, 1]$ se define como:

$$R_{Fd}(\Phi, \Phi') = \begin{cases} \bigwedge_{i \in \{1, \dots, n\}} T_{a_i}(v_i, w_i) & \text{si } n = m \text{ y } a_i = a'_i \text{ para todo } i \in \{1, \dots, n\} \\ 0 & \text{en caso contrario} \end{cases}$$

- Dado $\varepsilon \in [0, 1]$, se define el R_{Fd} - ε -bloque, para cada $\Phi \in For(\mathcal{A}_d)$, como:

$$[\Phi]_\varepsilon = \{\Phi' \in For(\mathcal{A}_d) \mid \varepsilon \leq R_{Fd}(\Phi, \Phi')\}$$

Las tablas de decisión permiten analizar la dependencia entre los atributos de condición $\mathcal{A}$ y el atributo de decisión d . Esta relación de dependencia es estudiada por medio de reglas de decisión, y diversos indicadores de relevancia asociados a ellas. Entre ellos, se destaca el T -soporte, que muestra la representatividad de las reglas en la tabla de decisión.

Definición 3. Sean S una tabla de decisión, T una familia de relaciones de tolerancia difusas separables y $C \subseteq \mathcal{A}$.

- Una regla de decisión en S es una expresión $\Phi \rightarrow \Psi$, donde $\Phi \in \text{For}(C)$, $\Psi \in \text{For}(\{d\})$ son el antecedente y consecuente de la regla de decisión, respectivamente.
- El T -soporte de la regla de decisión $\Phi \rightarrow \Psi$ es el valor:

$$\text{supp}_S^T(\Phi, \Psi) = \text{card}_F(\|\Phi \wedge \Psi\|_S^T)$$

donde $\text{card}_F(\cdot)$ denota el cardinal de un conjunto difuso.

III. ε -EFICIENCIA NORMALIZADA

Antes de presentar la definición de ε -eficiencia normalizada, es necesario introducir un nuevo indicador de relevancia, llamado ε - T -fuerza. Es un valor en el intervalo unidad que representa la proporción de objetos de la tabla de decisión que satisface la regla estudiada, y depende de un umbral ε . Este hecho permite adoptar diferentes criterios de similitud entre reglas, tanto para su agregación como para el cálculo de este indicador.

Definición 4. Sea S una tabla de decisión, T una familia de relaciones de tolerancia difusas separables y $\text{Dec}(S)$ un conjunto de reglas de decisión. Dado $\varepsilon \in [0, 1]$, se denomina ε - T -fuerza de la regla $\Phi \rightarrow \Psi$ con respecto a $\text{Dec}(S)$ a:

$$\sigma_S^{\varepsilon, T}(\Phi, \Psi) = \frac{\sum_{\{\Phi' \rightarrow \Psi' \in \text{Dec}(S) \mid \Phi' \in [\Phi]_\varepsilon, \Psi' \in [\Psi]_\varepsilon\}} \text{supp}_S^T(\Phi', \Psi')}{\sum_{\Phi'' \rightarrow \Psi'' \in \text{Dec}(S)} \text{supp}_S^T(\Phi'', \Psi'')}$$

La ε -eficiencia normalizada de un conjunto de reglas $\text{Dec}(S) = \{\Phi_i \rightarrow \Psi_i \mid i \in \{1, \dots, m\}\}$ se define a partir de las reglas más representativas similares a cada antecedente Φ_i . Con este fin, deben construirse los siguientes conjuntos, para todo $i \in \{1, \dots, m\}$:

$$\begin{aligned} B_{\Phi_i, \Psi_i} &= \{\Phi' \rightarrow \Psi' \in \text{Dec}(S) \mid \eta_{\Phi_i}^{N_\varepsilon}(\text{Dec}(S)) = \sigma_S^{\varepsilon, T}(\Phi', \Psi'), \Phi' \in [\Phi_i]_\varepsilon\} \\ \text{con } \eta_{\Phi_i}^{N_\varepsilon}(\text{Dec}(S)) &= \max\{\sigma_S^{\varepsilon, T}(\Phi, \Psi) \mid \Phi \rightarrow \Psi \in \text{Dec}(S), \Phi \in [\Phi_i]_\varepsilon\} \\ C_{\Phi_i, \Psi_i} &= \{\Phi'' \rightarrow \Psi'' \in \text{Dec}(S) \mid \Phi'' \in [\Phi_k]_\varepsilon, \Psi'' \in [\Psi_k]_\varepsilon, \\ &\quad k = \min\{l \in \{1, \dots, m\} \mid \Phi_l \rightarrow \Psi_l \in B_{\Phi_i, \Psi_i}\}\} \\ D_{\Phi_i, \Psi_i} &= C_{\Phi_i, \Psi_i} \setminus \{\Phi'' \rightarrow \Psi'' \in D_{\Phi_j, \Psi_j} \mid j < i\} \end{aligned}$$

que comentamos a continuación. En primer lugar, el conjunto B_{Φ_i, Ψ_i} recoge las reglas con máxima ε - T -fuerza y cuyos antecedentes son similares a Φ_i . Por su parte, el conjunto C_{Φ_i, Ψ_i} contiene a las reglas involucradas en el cálculo de la ε - T -fuerza de la primera regla del conjunto B_{Φ_i, Ψ_i} . Finalmente, el conjunto D_{Φ_i, Ψ_i} selecciona las reglas del conjunto C_{Φ_i, Ψ_i} que aún no han sido tenidas en cuenta. De esta forma, cada regla es considerada como máximo en una ocasión. La definición de ε -eficiencia normalizada se introduce a continuación.

Definición 5. Sean S una tabla de decisión, T una familia de relaciones de tolerancia difusas separables y $\text{Dec}(S) = \{\Phi_i \rightarrow \Psi_i \mid i \in \{1, \dots, m\}\}$ un conjunto de reglas de decisión. Dado $\varepsilon \in [0, 1]$, se denomina ε -eficiencia normalizada, denotada como N_ε -eficiencia, de $\text{Dec}(S)$ al valor:

$$\eta^{N_\varepsilon}(\text{Dec}(S)) = \frac{\sum_{i=1}^m \sum_{\Phi'' \rightarrow \Psi'' \in D_{\Phi_i, \Psi_i}} \text{supp}_S^T(\Phi'', \Psi'')}{\sum_{\Phi''' \rightarrow \Psi''' \in \text{Dec}(S)} \text{supp}_S^T(\Phi''', \Psi''')}$$

La ε -eficiencia normalizada satisface diversas propiedades, entre las que se destacan dos. La primera de ellas muestra que generaliza la noción de RST dada por Pawlak. La segunda presenta los valores entre los que se encuentra acotada, que coinciden con las cotas de RST.

Proposición 6. Sean S una tabla de decisión, T una familia de relaciones de tolerancia difusas separables y $\text{Dec}(S) = \{\Phi_i \rightarrow \Psi_i \mid i \in \{1, \dots, m\}\}$ un conjunto de reglas de decisión.

1. Si T es una familia de relaciones de tolerancia booleana y todo objeto de S satisface alguna regla de $\text{Dec}(S)$, entonces:

$$\eta^{N_\varepsilon}(\text{Dec}(S)) = \eta(\text{Dec}(S))$$

donde $\eta(\text{Dec}(S))$ es la eficiencia introducida en RST.

2. Se satisface la siguiente desigualdad:

$$\max\{\sigma_S^{\varepsilon, T}(\Phi, \Psi) \mid \Phi \rightarrow \Psi \in \text{Dec}(S)\} \leq \eta^{N_\varepsilon}(\text{Dec}(S)) \leq 1$$

para todo $\varepsilon \in [0, 1]$.

IV. CONCLUSIONES Y TRABAJO FUTURO

La eficiencia normalizada permite analizar la calidad de la clasificación basada en las reglas utilizadas para describir tablas de decisión. Además, el hecho de tomar valores en el intervalo unidad, junto con el resto de propiedades que satisface esta noción, facilita una interpretación óptima de la eficiencia obtenida. Como trabajo futuro, se continuará estudiando esta noción para relacionarla con la consistencia de las tablas de decisión. También, se comparará la eficiencia de diferentes conjuntos de reglas extraídas de la misma tabla con el fin de relacionar la eficiencia con la reducción de atributos.

REFERENCIAS

- [1] F. Chacón-Gómez, M. E. Cornejo, J. Medina, and E. Ramírez-Poussa. Efficiency of fuzzy rough set decision algorithms. *Studies in Computational Intelligence*, 1127:17–24, 2024.
- [2] F. Chacón-Gómez, M. E. Cornejo, J. Medina, and E. Ramírez-Poussa. Rough set decision algorithms for modeling with uncertainty. *Journal of Computational and Applied Mathematics*, 437:115413, 2024.
- [3] D. Dubois and H. Prade. Rough fuzzy sets and fuzzy rough sets. *International Journal of General Systems*, 17, 06 1990.
- [4] A. Nakamura. Fuzzy rough sets. *Note on Multiple-Valued Logic in Japan*, 9:1–8, 1988.
- [5] Z. Pawlak. Rough sets. *International Journal of Computer and Information Science*, 11:341–356, 1982.
- [6] Z. Pawlak. *Rough Sets: Theoretical Aspects of Reasoning About Data*. Kluwer Academic Publishers, Norwell, MA, USA, 1992.
- [7] Z. Pawlak. Rough sets, decision algorithms and bayes' theorem. *European Journal of Operational Research*, 136(1):181–189, 2002.

Ecuaciones bipolares de relaciones difusas con negación suelo

David Lobo Palacios
Departamento de Matemáticas
Universidad de Cádiz
Cádiz, España
david.lobouca@uca.es

Pablo López Molina
Departamento de Matemáticas
Universidad de Cádiz
Cádiz, España
pablo.lopezmolina@uca.es

Resumen—Una negación suelo se define como una modificación de un operador de negación, de tal forma que la negación suelo coincide con el operador de negación por debajo de un cierto umbral y toma el valor 0 en otro caso. Tras introducir el concepto de negación suelo, en este trabajo se estudia la resolución de ecuaciones bipolares de relaciones difusas con negación suelo, caracterizando la resolubilidad y la existencia de soluciones maximales y minimales de tales ecuaciones.

Palabras clave—ecuación bipolar de relaciones difusas, t-norma producto, negación involutiva, análisis forense digital

I. INTRODUCCIÓN

Una ecuación de relaciones difusas (FRE) es una expresión en la que la composición de dos relaciones difusas, siendo una de ellas desconocida, da lugar a una tercera relación difusa [1], [5], [8]. Las ecuaciones bipolares de relaciones difusas (BFRE), por su parte, se definen como una generalización de las anteriores donde la relación desconocida aparece junto a su negación lógica, modelando la ausencia de la misma [2]–[4], [6]. La definición original de BFRE se fundamenta en la idea de que la presencia y la ausencia de una relación desconocida afectan simultáneamente al término independiente de la ecuación.

En este artículo se considera una versión debilitada de la afirmación anterior, de forma que la ausencia de la relación desconocida únicamente afecta al término independiente de la ecuación cuando la relación desconocida es inferior a un cierto umbral. Para ello, se introduce el concepto de negación suelo como una modificación de un operador de negación tal que coincide con este por debajo de un determinado valor y se anula por encima de dicho valor. Además, se estudia la resolución de BFRE definidas con la composición máximo producto y con una negación suelo asociada a una negación involutiva, proporcionando una caracterización de su resolubilidad y algunos resultados sobre la existencia y la expresión analítica de soluciones maximales y minimales. En particular, se caracteriza la existencia de una solución máxima y de una solución mínima.

II. LA NEGACIÓN SUELO

Una ecuación bipolar de relaciones difusas (BFRE) es una expresión de la forma:

$$\bigvee_{j=1}^m (a_j^+ * x_j) \vee (a_j^- * \neg x_j) = b \quad (1)$$

donde $a_j^+, a_j^-, b \in [0, 1]$ para cada $j \in \{1, \dots, m\}$, las variables $x_1, \dots, x_m$ son desconocidas y $\neg: [0, 1] \rightarrow [0, 1]$ es un operador de negación, es decir, una aplicación no creciente tal que $\neg 0 = 1$ y $\neg 1 = 0$.

Nótese que tanto la presencia de las variables $x_1, \dots, x_m$ como la ausencia de las mismas, modeladas como $\neg x_1, \dots, \neg x_m$, afectan al valor del término independiente de la expresión (1). Aún más, en general, el valor de $\neg x_1, \dots, \neg x_m$ tiene un efecto permanente en el miembro de la derecha de (1).

Sin embargo, cabe considerar ciertas situaciones en las que la ausencia de las variables implicadas en una BFRE no tienen una influencia continua en el término independiente de la misma, sino que lo hacen únicamente por debajo de un umbral. Para modelar este tipo de situaciones, se introduce un nuevo tipo de negación denominada α -negación suelo.

Definición 1. Sean $\neg: [0, 1] \rightarrow [0, 1]$ un operador de negación y $\alpha \in [0, 1]$. La α -negación suelo de $\neg$ es la aplicación $\neg_\alpha: [0, 1] \rightarrow [0, 1]$ dada por:

$$\neg_\alpha(x) = \begin{cases} \neg(x) & \text{si } x \leq \alpha \\ 0 & \text{si } x > \alpha \end{cases}$$

Nótese que la α -negación suelo de un operador de negación es, asimismo, un operador de negación.

III. ECUACIONES BIPOLARES DE RELACIONES DIFUSAS CON NEGACIÓN SUELO

En esta sección se presenta una primera aproximación a la resolución de BFRE con negación suelo. En particular, se estudiará el caso de BFRE definidas con la composición máximo producto y una α -negación suelo asociada a una negación involutiva.

Formalmente, analizaremos la resolución de una BFRE de la forma

$$\bigvee_{j=1}^m (a_j^+ * x_j) \vee (a_j^- * \neg_{\alpha} x_j) = b \quad (2)$$

siendo $\neg: [0, 1] \rightarrow [0, 1]$ una negación involutiva y $(*, \leftarrow_P)$ el par residuado formado por la t-norma producto $*$: $[0, 1] \times [0, 1] \rightarrow [0, 1]$ y su implicación residuada $\leftarrow_P: [0, 1] \times [0, 1] \rightarrow [0, 1]$, definida como:

$$z \leftarrow_P y = \begin{cases} \frac{z}{y} & \text{si } z < y \\ 1 & \text{si } y \leq z \end{cases}$$

Destacar que la propiedad de adjunción resulta esencial en los resultados mostrados a continuación [7].

El siguiente resultado presenta una sencilla caracterización de resolubilidad de (2) en términos de las siguientes tuplas:

$$\begin{aligned} g^+ &= (b \leftarrow_P a_1^+, \dots, b \leftarrow_P a_m^+) \\ \neg g^- &= (\neg(b \leftarrow_P a_1^-), \dots, \neg(b \leftarrow_P a_m^-)) \end{aligned}$$

Teorema 2. *La BFRE (2) es resoluble si y solo si g^+ o $\neg g^-$ es una solución.*

Atendiendo a la estructura algebraica del conjunto de soluciones de (2), se muestra una caracterización de sus soluciones maximales, así como de la existencia de una solución máxima.

Teorema 3. *Sea (2) una BFRE resoluble.*

- a) Si g^+ es solución de (2), entonces es la solución máxima.
- b) Si g^+ no es solución de (2), entonces el conjunto de soluciones maximales no es vacío y viene dado por:

$$\{(1, \dots, 1, \neg g_k^-, 1, \dots, 1) \mid k \in K^-\}$$

$$\text{donde } K^- = \{j \in \{1, \dots, m\} \mid a_j^- \geq b, \neg g_j^- \leq \alpha\}.$$

Nótese que siempre existen soluciones maximales de (2) cuando esta es resoluble. Por su parte, la existencia de soluciones minimales no está garantizada en general.

Teorema 4. *Sea (2) una BFRE resoluble.*

- a) Si $\neg g^- \leq (\alpha, \dots, \alpha)$ y $\neg g^-$ es solución de (2), entonces $\neg g^-$ es la solución mínima.
- b) Si $\neg g^- \leq (\alpha, \dots, \alpha)$ y $\neg g^-$ no es solución de (2), entonces el conjunto de soluciones minimales no es vacío y viene dado por:

$$\{(0, \dots, 0, g_k^+, 0, \dots, 0) \mid a_k^+ \geq b\}$$

- c) Si existe un único $k \in \{1, \dots, m\}$ tal que $\neg g_k^- > \alpha$:
 - Si $a_k^+ \geq b$ y $a_j^- < b$ para cada $j \in \{1, \dots, m\}$, $j \neq k$, la única solución minimal es $(0, \dots, 0, g_k^+, 0, \dots, 0)$.
 - Si $a_k^+ < b$ o existe algún $k' \neq k$ tal que $a_{k'}^- \geq b$, entonces no existen soluciones minimales.
- d) Si existen $k, k' \in \{1, \dots, m\}$ con $k \neq k'$ tales que $\neg g_k^-, \neg g_{k'}^- > \alpha$, no existen soluciones minimales.

IV. APLICACIÓN AL ANÁLISIS FORENSE DIGITAL

En esta sección se presenta un ejemplo de aplicación de BFRE con negación suelo en un contexto de análisis forense digital. Supongamos que queremos determinar si un homicidio ha sido premeditado, para lo cual se tiene en cuenta la falta de testigos, denotada $x_1 \in [0, 1]$.

Parece natural que una baja o nula cantidad de testigos, i.e. un alto valor x_1 , estará directamente relacionado con la premeditación del crimen. Podemos modelar este hecho como $a^+ * x_1$ con $a^+ > 0$. Por otro lado, una cantidad excesivamente alta de testigos, i.e. un valor de x_1 por debajo de un cierto umbral $\alpha > 0$, podría inducir a pensar que el crimen también es premeditado debido a la posibilidad de distracciones entre la multitud. Nótese que el término $a^- * (1 - x_1)$ no modela correctamente esto último, puesto que la evidencia de premeditación solo se tiene cuando $x_1 < \alpha$. Una forma razonable de modelar la planificación del asesinato en función de la elevada presencia de testigos es la expresión $a^- * \neg_{\alpha} x_1$, donde $\neg_{\alpha}$ es la negación suelo asociada a la negación estándar $\neg x = 1 - x$. Con lo cual, el efecto del número de testigos en la premeditación del crimen puede modelarse como

$$(a^+ * x) \vee (a^- * \neg_{\alpha} x)$$

Si b denota el nivel de premeditación de un crimen, razonamientos similares para variables como la presencia de ADN y la accesibilidad al arma del crimen, denotadas $x_2 \in [0, 1]$ y $x_3 \in [0, 1]$, darían lugar a una expresión de la forma (2).

Asumiendo que el valor de b es conocido, por ejemplo inferido de la decisión de un jurado popular, la resolución de la correspondiente BFRE con negación suelo permite concluir si la decisión es coherente de acuerdo con las soluciones de (2).

V. CONCLUSIONES

El estudio sobre la resolución de BFRE con negación suelo muestra que el comportamiento de tales ecuaciones es similar al de las BFRE definidas con una negación involutiva [2]. Se pretende como trabajo futuro profundizar en esta similitud.

REFERENCIAS

- [1] R. Bělohlávek, "Fuzzy Relational Systems: Foundations and Principles," Kluwer Academic Publishers, 2002.
- [2] M. E. Cornejo, D. Lobo, J. Medina and B. De Baets, "Bipolar equations on complete distributive symmetric residuated lattices: The case of a join-irreducible right-hand side," *Fuzzy Sets and Systems*, vol. 442, pp. 92–108, 2022.
- [3] M. E. Cornejo, D. Lobo and J. Medina, "Bipolar fuzzy relation equations systems based on the product t-norm," *Mathematical Methods in the Applied Sciences*, vol. 42, no. 17, pp. 5779–5793, 2019.
- [4] M. E. Cornejo, D. Lobo and J. Medina, "Optimization of partially monotonic functions subject to bipolar fuzzy relation equations," *Information Sciences*, vol. 648, 2023.
- [5] A. Di Nola, E. Sanchez, W. Pedrycz and S. Sessa, "Fuzzy relation equations and their applications to knowledge engineering," Kluwer Academic Publishers, 1989.
- [6] S. Freson, B. De Baets and H. De Meyer, "Linear optimization with bipolar max–min constraints," *Information Sciences*, vol. 234, pp. 3–15, 2013.
- [7] E. P. Klement, R. Mesiar and E. Pap, "Triangular norms," Kluwer Academic Publishers, 2000.
- [8] E. Sanchez, "Resolution of composite fuzzy relation equations," *Information and Control*, vol. 30, no. 1, pp. 38–48, 1976.

On direct systems of implications with graded attributes

Manuel Ojeda-Hernández

Depto. de Álgebra, Geometría y Topología
Universidad de Málaga
Málaga, Spain
manuojeda@uma.es

Domingo López-Rodríguez

Depto. de Matemática Aplicada
Universidad de Málaga
Málaga, Spain
dominlopez@uma.es

Abstract—This keywork (based on the paper “On Direct Systems of Implications with Graded Attributes” published in EUSFLAT’25) investigates the concept of direct systems of implications within the fuzzy setting, where attributes can be satisfied to varying degrees. The directness property of such systems is advantageous as it allows for the rapid computation of closure operators, a relevant task in applications like Fuzzy Formal Concept Analysis (FFCA). The study provides a novel algebraic characterization of directness for these fuzzy implications, connecting it to principles from Simplification Logic. Furthermore, this paper outlines an iterative algorithmic approach, grounded in this characterization, for constructing equivalent direct systems from general implicational systems.

Index Terms—Fuzzy Formal Concept Analysis, Bases of implications, Simplification Logic, Direct Systems

I. INTRODUCTION

Implicational rules, such as functional dependencies or Horn clauses, are fundamental in various mathematical and computational fields, including relational databases and logic programming. In Formal Concept Analysis (FCA), introduced by Wille [5], and its fuzzy extension (FFCA), implications represent knowledge extracted from data tables. FFCA, initiated by Burusco and Fuentes-González using lattices as truth values [4], and later developed by Pollandt [10], Bělohávek [12] and others using complete residuated lattices $\mathbb{L} = (L, \wedge, \vee, \otimes, \rightarrow, 0, 1)$, accommodates attributes that can be satisfied to a certain degree.

The computation of an element’s closure under a system of implications is typically an iterative procedure until a fixed point is reached [7]. Research to optimize this has followed two main lines: minimizing the number of implications in a basis (e.g., through pseudointents [8], [9]), or developing systems that compute closures in a single iteration, known as direct systems [11]. While direct systems might be larger than minimal bases, their ability to compute closures quickly makes them suitable when numerous closure computations are required. This keywork, based on [1], extends the notion of direct systems to the fuzzy setting where attribute implications

$A \rightarrow B$ involve fuzzy sets $A, B \in L^M$ (M being the set of attributes of a fuzzy formal context $\mathbb{K} = (G, M, I)$ with $I : G \times M \rightarrow L$). This approach differs from some prior work on directness in fuzzy settings which considered implications between crisp sets with an associated degree of validity [6]. We assume L is a finite numerical chain throughout this paper.

II. FUZZY DIRECT SYSTEMS AND THEIR CHARACTERIZATION

A fuzzy formal context is a triplet $\mathbb{K} = (G, M, I)$, where G and M are sets of objects and attributes, and $I : G \times M \rightarrow L$ is a fuzzy relation. For fuzzy sets of attributes $A, B \in L^M$, an implication $A \rightarrow B$ is an expression where A is the premise and B is the consequence. Let $\Sigma = \{A_i \rightarrow B_i\}$ be a system of such fuzzy implications. The operator $\pi_\Sigma : L^M \rightarrow L^M$ is defined for $X \in L^M$ and $m \in M$ as:

$$\pi_\Sigma(X)(m) = X(m) \vee \bigvee \{S(A, X) \otimes B(m) \mid A \rightarrow B \in \Sigma\} \quad (1)$$

Here, $S(A, X)$ represents the degree to which fuzzy set A is contained in fuzzy set X (with $S(A, X) = \bigwedge_{k \in M} (A(k) \rightarrow X(k))$ [12], [13]), and $\otimes$ is the t-norm of the residuated lattice L . This operator is inflationary and isotone. The corresponding closure operator, φ_Σ , is obtained by iterating π_Σ until a fixed point is reached:

$$\varphi_\Sigma(X) = \pi_\Sigma(X) \cup \pi_\Sigma^2(X) \cup \pi_\Sigma^3(X) \cup \dots \quad (2)$$

where $\pi_\Sigma^k(X)$ is the k -th application of π_Σ to X .

Definition II.1. A system of implications Σ is said to be *direct* if $\varphi_\Sigma(X) = \pi_\Sigma(X)$ for all $X \in L^M$.

This means that for a direct system, the closure is computed in one step via π_Σ . The main theoretical result is the characterization of direct systems via the fuzzy exchange condition. The assumption that L is a finite chain ensures that if $\pi_\Sigma(X)(m) > X(m)$, then there exists a single implication $G \rightarrow H \in \Sigma$ such that $\pi_\Sigma(X)(m) = S(G, X) \otimes H(m)$.

Definition II.2. An implicational system Σ is said to satisfy the fuzzy exchange condition if for all pairs of implications $A \rightarrow B, C \rightarrow D \in \Sigma$, for all $m, n \in M$ and for all $c, d \in L$: if $c \otimes B(m) > c \otimes A(m)$, $d \otimes D(n) > c \otimes A(n) \vee d \otimes C(n)$,

This research is partially supported by the State Agency of Research (AEI), the Spanish Ministry of Science, Innovation, and Universities (MCIU) and the European Social Fund (FEDER) through the research projects with reference PID2021-127870OB-I00 (MCIU/AEI/FEDER, UE) and the VALID research project (PID2022-140630NB-I00 funded by MCIN/AEI/ 10.13039/501100011033).

and $d \otimes C(m) < c \otimes B(m)$, then there exists an implication $G \rightarrow H \in \Sigma$ such that

$$S(G, c \otimes A \cup (d \otimes C \setminus c \otimes B)) \otimes H(n) \geq d \otimes D(n). \quad (3)$$

The operation $X \setminus Y$ (pointwise $a \setminus b$) is defined such that $a \setminus b \leq k \iff a \leq b \vee k$. Juxtaposition XY denotes $X \cup Y$.

Theorem II.3. *An implicational system Σ is direct if and only if it satisfies the fuzzy exchange condition.*

The proof of this theorem relies on showing that if the fuzzy exchange condition holds, then $\pi_\Sigma^2(X) \subseteq \pi_\Sigma(X)$ for any $X \in L^M$, which, given π_Σ is inflationary, implies $\pi_\Sigma^2(X) = \pi_\Sigma(X)$ and thus directness. Also, it shows that any supposed counterexample to directness can be iteratively reduced using the fuzzy exchange condition until a contradiction is reached. The linearity and finiteness of L are important in this result.

III. ALGORITHMIC CONSTRUCTION

Building upon the fuzzy exchange condition, this work proposes an algorithm, termed `DirectSystem` (Algorithm 1 in the original work), to convert an implicational system Σ_{in} into an equivalent direct system Σ_d . The algorithm first ensures the input system is *reduced*, meaning for each $A \rightarrow B \in \Sigma_{in}$, it is transformed into $A \rightarrow B \setminus A$ using the [DeEq] equivalence from Fuzzy Attribute Simplification Logic (FASL) [13]. For such reduced systems, the fuzzy exchange condition can be stated more simply.

The core of the algorithm involves iteratively generating “derived implications”. Proposition 1 of the summarized work establishes that for any two valid implications $A \rightarrow B, C \rightarrow D$ and scalars $c, d \in L$, the implication $G' \rightarrow H'$ where $G' = c \otimes A \cup (d \otimes C \setminus c \otimes B)$ and $H' = d \otimes D \setminus c \otimes A$ is also valid and satisfies the consequent of the fuzzy exchange condition if $d \otimes D(n) > c \otimes A(n)$ for some n . The `DirectSystem` algorithm repeatedly calls two main helper functions: `AddDerived` and `Combine`.

- `AddDerived`: For each pair of implications from the current system Σ and each pair of scalars $c, d \in L$, it constructs all such necessary derived implications $G' \rightarrow H'$ and collects them in a set $\mathcal{D}$.
- `Combine`: This function takes the current system Σ and the set $\mathcal{D}$ of derived implications. It integrates $\mathcal{D}$ into Σ . If a derived $G' \rightarrow H'$ has a premise G' already in Σ as $G' \rightarrow H_{old}$, the implication is updated to $G' \rightarrow H_{old} \cup H'$ (using FASL rule [UnEq] [13]). If G' is new, $G' \rightarrow H'$ is added.

The main loop continues until the system satisfies the fuzzy exchange condition and is thus direct. It can be proved that the `DirectSystem` algorithm terminates (due to finite M and L) and outputs a direct system equivalent to the input.

IV. CONCLUSION AND FUTURE WORK

This work has formally defined direct systems of implications in the fuzzy setting where attributes are graded, providing a distinct approach from previous fuzzy implication models. The primary contribution is the algebraic characterization of

these systems via the fuzzy exchange condition, which is a non-trivial extension from the crisp case. Based on this, an iterative algorithm (`DirectSystem`) has been designed and proven to transform any valid fuzzy implicational system into an equivalent direct one, facilitating faster closure computations. The algorithm’s correctness relies on systematically adding derived implications until the fuzzy exchange condition is met globally.

Future research avenues include a thorough analysis of the computational complexity of the proposed algorithm, which involves $|L|^2$ potential derivations per pair of implications in each iteration. Additionally, an investigation into *direct-optimal* systems is planned, aiming to find the most concise and interpretable direct systems without loss of inferential power, enhancing practical applicability.

ACKNOWLEDGMENT

The authors would like to thank Dr. Kira Adaricheva who inspired us to study this topic.

REFERENCES

- [1] M. Ojeda-Hernández and D. López-Rodríguez. “On direct systems of implications with graded attributes”. In: 14th Conference of the European Society for Fuzzy Logic and Technology. Riga, 2025.
- [2] R. Bělohávek and V. Vychodil, “Fuzzy attribute implications: computing non-redundant bases using maximal independent sets,” In: Australasian Joint Conference on Artificial Intelligence. pp. 1126–1129. Springer (2005).
- [3] R. Bělohávek and V. Vychodil, “Attribute implications in a fuzzy setting,” In: Formal Concept Analysis: 4th International Conference, ICFCA 2006, Dresden, Germany, February 13–17, 2006. Proceedings. pp. 45–60. Springer (2006).
- [4] A. Burusco and R. Fuentes-González, “The study of the L-fuzzy concept lattice,” *Mathware and soft computing* 3, 209–218 (1994).
- [5] R. Wille, “Restructuring lattice theory: an approach based on hierarchies of concepts,” In: Ordered sets. pp. 445–470. Springer Netherlands (1982).
- [6] P. Cordero, M. Enciso, and A. Mora, “Directness in fuzzy formal concept analysis,” In: International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems. pp. 585–595. Springer (2018).
- [7] B. Ganter and S. Obiedkov, “Conceptual exploration,” Springer (2016).
- [8] J. Guigues and V. Duquenne, “Familles minimales d’implications informatives résultant d’une table de données binaires,” *Mathématiques et Sciences Humaines* 95, 5–18 (1986).
- [9] M. Ojeda-Hernández, I.P. Cabrera, and P. Cordero, “Quasi-closed elements in fuzzy posets,” *Journal of Computational and Applied Mathematics* 404, 113390 (2022).
- [10] S. Pollandt, “Fuzzy Begriffe: Formale Begriffsanalyse von unscharfen Daten,” Springer-Verlag, Berlin–Heidelberg (1997).
- [11] E. Rodríguez-Lorenzo, K. Bertet, P. Cordero, M. Enciso, and A. Mora, “Direct-optimal basis computation by means of the fusion of simplification rules,” *Discret. Appl. Math.* 249, 106–119 (2018).
- [12] R. Bělohávek, “Fuzzy Relational Systems,” Springer (2002).
- [13] R. Bělohávek, P. Cordero, M. Enciso, A. Mora, V. Vychodil: Automated prover for attribute dependencies in data with grades. *Int. J. Approximate Reasoning* 70, 51–67 (2016).

Teorías, modelos y bases de implicaciones de atributos en retículos de conceptos multiadjuntos

M. Eugenia Cornejo
Departamento de Matemáticas
Universidad de Cádiz
Cádiz, España
mariaeugenia.cornejo@uca.es

Jesús Medina
Departamento de Matemáticas
Universidad de Cádiz
Cádiz, España
jesus.medina@uca.es

Francisco José Ocaña
Departamento de Matemáticas
Universidad de Cádiz
Cádiz, España
franciscojose.ocana@uca.es

Resumen—Este trabajo recopila resultados recientes sobre la obtención de bases de implicaciones de atributos en el marco de los retículos de conceptos multiadjuntos con intensificadores que se han presentado en el artículo científico *Theories, models and bases of attribute implications in multi-adjoint concept lattices with hedges*, 2025.

Palabras clave—Implicaciones de atributos, teoría, modelos, bases, retículos de conceptos multiadjuntos.

I. INTRODUCCIÓN

El modelado de sistemas de conocimiento representa hoy en día uno de los grandes desafíos para la extracción de información relevante a partir de grandes volúmenes de datos. En este contexto, el Análisis de Conceptos Formales (FCA, por sus siglas en inglés) [1] ofrece una herramienta eficaz para la obtención de reglas, conocidas como implicaciones de atributos [1], [2], que facilitan el tratamiento y la gestión del conocimiento. Determinar de manera eficiente un conjunto óptimo de implicaciones de atributos válidas constituye un desafío significativo en FCA, debido al elevado número de implicaciones válidas que pueden derivarse de un conjunto de datos, así como la presencia de implicaciones redundantes. Siguiendo la línea de investigación desarrollada en [4], en este trabajo abordamos el desarrollo teórico de las nociones y propiedades relacionadas con los conceptos de teoría y modelo, lo que nos permitirá elaborar la noción de base de implicaciones de atributos en el marco de los retículos de conceptos multiadjuntos con intensificadores.

II. IMPLICACIONES DE ATRIBUTOS. TEORÍA Y MODELOS

El ambiente natural de este trabajo es el análisis de conceptos formales en el marco de los retículos de conceptos multiadjuntos con intensificadores [3], [4]. Por ello, a fin de presentar formalmente los resultados obtenidos, necesitamos considerar un marco multiadjunto $(L_1, L_2, P, \preceq_1, \preceq_2, \preceq, \&_1, \preceq^1, \preceq_1, \dots, \&_n, \preceq^n, \preceq_n)$, un triple adjunto $(\&, \preceq, \preceq)$

Parcialmente subvencionado por el proyecto PID2022-137620NB-I00 financiado por MICIU/AEI/10.13039/501100011033 y por FEDER, UE, por el proyecto TED2021-129748B-I00 financiado por MICIU/AEI/10.13039/501100011033 y por la Unión Europea NextGenerationEU/PRTR, y por el contrato predoctoral industrial PU/EPIF-FPI-GRUPOENERGETICOPUERTOREAL/CP/2022-051, correspondiente al Programa de Fomento e Impulso de la Investigación y la Transferencia de la Universidad de Cádiz 2018/2019.

externo al marco con respecto a $(L_1, \preceq_1)$, dos familias arbitrarias $*_A$ y $*_B$ de intensificadores de verdad, y el contexto $(A*_A, B*_B, R, \sigma)$ junto con el correspondiente retículo de conceptos multiadjunto $(\mathcal{M}_*, \preceq)$. Siguiendo la misma filosofía que en el marco residuado, se presentará una generalización de las definiciones, propiedades y resultados introducidos en [5] sobre teorías y modelos de implicaciones de atributos al marco multiadjunto.

Definición 1: Sean $f_1, f_2 \in L_1^A$.

- Una *implicación de atributos* sobre A es una expresión $f_2 \Leftarrow f_1$. Al conjunto formado por todas las implicaciones sobre A lo denotaremos como $\mathcal{I}_A$.
- Una *teoría* es un subconjunto difuso de implicaciones de atributos sobre A , esto es $\mathcal{T}: \mathcal{I}_A \rightarrow L_1$. Denotaremos como $\mathcal{T}(f_2 \Leftarrow f_1)$ al grado en el que la implicación de atributos $f_2 \Leftarrow f_1$ pertenece a $\mathcal{T}$.
- Un *modelo* de $\mathcal{T}$ es un subconjunto difuso de atributos $M \in L_1^A$ satisfaciendo que $\mathcal{T}(f_2 \Leftarrow f_1) \preceq_1 \|f_2 \Leftarrow f_1\|_M$, para todo $f_1, f_2 \in L_1^A$.
- Una *teoría clásica* es un subconjunto de implicaciones de atributos sobre A .
- Un *modelo* de una teoría clásica $\mathcal{T}_{cr}$ es un subconjunto difuso de atributos $M \in L_1^A$ satisfaciendo que $\|f_2 \Leftarrow f_1\|_M = \top_1$, para todo $f_1, f_2 \in L_1^A$ tal que $f_2 \Leftarrow f_1 \in \mathcal{T}_{cr}$.

El siguiente teorema presenta un resultado clave en este trabajo para la construcción de bases, el cual nos permite expresar cualquier teoría como una teoría clásica, y en consecuencia con los mismos modelos asociados.

Teorema 1: Dada una teoría $\mathcal{T}$, dos subconjuntos difusos de atributos $f_1, f_2 \in L_1^A$ y la teoría clásica $\mathcal{T}_{cr}$ definida como:

$$\mathcal{T}_{cr} = \{f \Leftarrow f_1 \mid f_1, f_2 \in L_1^A \text{ and } f \neq f_{\perp_1}\}$$

donde $f = f_2 \& \mathcal{T}(f_2 \Leftarrow f_1)$, y $f_{\perp_1}: A \rightarrow L_1$ es el subconjunto difuso definido como $f_{\perp_1}(a) = \perp_1$, para todo $a \in A$. Si el conjuntor $\&$ del triple adjunto $(\&, \preceq, \preceq)$ es asociativo y satisface que $x \& \top_1 = x$, para todo $x \in L_1$, entonces se cumple que:

- (1) $Mod(\mathcal{T}) = Mod(\mathcal{T}_{cr})$
- (2) $\|f_2 \Leftarrow f_1\|_{\mathcal{T}} = \|f_2 \Leftarrow f_1\|_{\mathcal{T}_{cr}}$

A continuación, nos centraremos en dos propiedades fundamentales de las teorías de implicaciones de atributos: la

completitud y la no redundancia. Estas propiedades resultan esenciales para la definición de bases de implicaciones de atributos, tal y como mostraremos en detalle más adelante.

Definición 2: Una teoría $\mathcal{T}$ se dice *completa* si para toda implicación de atributos $f_2 \Leftarrow f_1$ se satisface que:

$$\|f_2 \Leftarrow f_1\|_{\mathcal{T}} = \|f_2 \Leftarrow f_1\|_{\mathcal{M}_*}$$

A partir de ahora, para introducir formalmente la definición de no redundancia, consideraremos $\mathcal{T}$ como una teoría clásica. Este hecho puede asumirse sin pérdida de generalidad, debido a los resultados establecidos en el Teorema 1.

Definición 3: Una teoría clásica $\mathcal{T}_{cr}$ se dice *no redundante* si para toda implicación de atributos $f_2 \Leftarrow f_1 \in \mathcal{T}_{cr}$ se satisface que $\|f_2 \Leftarrow f_1\|_{\mathcal{T}_{cr} \setminus \{f_2 \Leftarrow f_1\}} \neq \top_1$.

III. BASES DE IMPLICACIONES DE ATRIBUTOS

El objetivo principal de esta sección es presentar una extensión de la noción de base de implicaciones de atributos al marco multiadjunto. En particular, estudiaremos una de las bases más usuales consideradas en FCA, denominada base de Duquenne-Guigues. A partir de una teoría completa podemos definir la noción de base de implicaciones de atributos, como se muestra a continuación.

Definición 4: Una teoría $\mathcal{T}_{cr}$ se denomina base de $(\mathcal{M}_*, \preceq)$ si es completa y ningún subconjunto propio de $\mathcal{T}_{cr}$ es completo.

El siguiente teorema muestra que la noción de base puede caracterizarse en términos de la completitud y la no redundancia. La completitud garantiza que todas las demás implicaciones válidas puedan deducirse de este conjunto. La no redundancia asegura que las implicaciones de la base no pueden deducirse unas de otras.

Teorema 2: Una teoría $\mathcal{T}_{cr}$ es una base de $(\mathcal{M}_*, \preceq)$ si y solo si se satisfacen las siguientes condiciones:

- (1) $\mathcal{T}_{cr}$ es completa.
- (2) $\mathcal{T}_{cr}$ es no redundante.

Como mencionamos previamente, estamos interesados en definir la base de Duquenne-Guigues en el marco multiadjunto con intensificadores de verdad. Esta base depende en gran medida de la noción de pseudointensión [1], cuya definición presentamos a continuación.

Definición 5: Un subconjunto difuso de atributos $\mathcal{P} \subseteq L_1^A$ se dice que es un *sistema de pseudointensiones* si, dado $P \in L_1^A$, para todo $Q \in \mathcal{P}$ tal que $Q \neq P$ se tiene que:

$$P \in \mathcal{P} \text{ si y solo si } P \neq P^{\downarrow\uparrow*} \text{ y } \|Q^{\downarrow\uparrow*} \Leftarrow Q\|_P = \top_1,$$

La importancia de la noción de sistema de pseudointensiones radica en su aplicación directa para la extracción de bases de implicaciones de atributos. En particular, permite extender la base de Duquenne-Guigues al marco multiadjunto, siendo clave para definir las implicaciones de atributos de dicha base, como se muestra en el siguiente teorema.

Teorema 3: Si el conjuntor $\&$ del triple adjunto $(\&, \swarrow, \searrow)$ es asociativo, $x \& \top_1 = x$, $\top_1 \& y = y$, para todo $x, y \in L_1$, y $\swarrow$ es una implicación forzada. Dado un retículo de conceptos multiadjunto $(\mathcal{M}_*, \preceq)$, si $\mathcal{P}$ es un sistema de pseudointensiones de $(\mathcal{M}_*, \preceq)$, entonces la teoría definida como:

$$\mathcal{T}_{\mathcal{P}} = \{P^{\downarrow\uparrow*} \Leftarrow P \mid P \in \mathcal{P}\}$$

es una base de $(\mathcal{M}_*, \preceq)$.

Finalmente presentamos un ejemplo ilustrativo sobre cómo podemos aplicar el resultado anterior para obtener bases de implicaciones de atributos a partir de un contexto dado.

Ejemplo 1: Consideremos el siguiente marco multiadjunto $([0, 1]_2, \leq, (\&_{DP}, \swarrow_{DP}^{\text{DP}}, \searrow_{DP}), (\&_{DL}, \swarrow_{DL}^{\text{DL}}, \searrow_{DL}))$, el contexto $(A_{*A}, B_{*B}, R, \sigma)$ donde $[0, 1]_2$ es una partición del intervalo unidad en dos trozos, $A = \{a_1, a_2\}$, $B = \{b_1, b_2\}$, la relación R se define como: $R(b_1, a_1) = 0.5$, $R(b_1, a_2) = 1$, $R(b_2, a_1) = 0$, $R(b_2, a_2) = 0.5$; $(\&_{DP}, \swarrow_{DP}^{\text{DP}}, \searrow_{DP})$ y $(\&_{DL}, \swarrow_{DL}^{\text{DL}}, \searrow_{DL})$ son los triples adjuntos definidos a partir de la discretización de la t-norma producto y Łukasiewicz en $[0, 1]_2$, $*_A$ y $*_B$ son familias de intensificadores identidad, y σ asigna el triple adjunto correspondiente a $\&_{DL}$ para b_1 y a $\&_{DL}$ para b_2 .

A partir de la Definición 5, el único sistema de pseudointensiones que se obtiene es el siguiente:

$$\mathcal{P} = \{\{\}, \{1/a_2\}, \{0.5/a_1, 0.5/a_2\}\}$$

De esta manera, aplicando el Teorema 3, se tiene que la teoría definida como $\mathcal{T}_{\mathcal{P}} = \{P^{\downarrow\uparrow*} \Leftarrow P \mid P \in \mathcal{P}\}$ constituye una base de $(\mathcal{M}_*, \preceq)$, compuesta por las siguientes implicaciones de atributos:

$$\begin{aligned} \mathcal{T}_{\mathcal{P}} = \{ & \{0.5/a_2\} \Leftarrow \{\}, \{0.5/a_1, 1/a_2\} \Leftarrow \{1/a_2\}, \\ & \{0.5/a_1, 1/a_2\} \Leftarrow \{0.5/a_1, 0.5/a_2\} \} \end{aligned}$$

Normalmente, las bases se representan eliminando la información redundante en el consecuente, con el objetivo destacar la información más significativa. Tras dicha simplificación, la base anterior queda expresada como:

$$\begin{aligned} \mathcal{T}_{\mathcal{P}} = \{ & \{0.5/a_2\} \Leftarrow \{\}, \{0.5/a_1\} \Leftarrow \{1/a_2\}, \\ & \{1/a_2\} \Leftarrow \{0.5/a_1, 0.5/a_2\} \} \end{aligned}$$

IV. CONCLUSIONES Y TRABAJO FUTURO

En este artículo se ha presentado un desarrollo teórico de las nociones y resultados necesarios para extender la noción de base de implicaciones al marco multiadjunto. Como trabajo futuro, nos centraremos en el desarrollo de mecanismos eficientes para la extracción de implicaciones válidas de atributos, a fin optimizar la construcción de las bases propuestas en este trabajo.

REFERENCIAS

- [1] B. Ganter and R. Wille, *Formal Concept Analysis: Mathematical Foundation*. Springer Verlag, 1999.
- [2] J.-L. Guigues and V. Duquenne. Familles minimales d'implications informatives résultant d'un tableau de données binaires. *Mathématiques et Sciences Humaines*, 95:5–18, 1986.
- [3] J. Medina, M. Ojeda-Aciego, and J. Ruiz-Calviño, Formal concept analysis via multi-adjoint concept lattices, *Fuzzy Sets and Systems*, vol. 160, no. 2, pp. 130–144, 2009.
- [4] M. E. Cornejo, J. Medina, and F. J. Ocaña, Attribute implications in multi-adjoint concept lattices with hedges, *Fuzzy Sets and Systems*, vol. 479, p. 108854, 2024.
- [5] R. Belohlávek and V. Vychodil. Attribute dependencies for data with grades I.. *International Journal of General Systems*, 45(7-8):864–888, 2016.

Suma de contextos basada en enlaces multiadjuntos*

Roberto G. Aragón, Jesús Medina, Samuel Molina-Ruiz

Departamento de Matemáticas, Universidad de Cádiz, España

Correo: {roberto.aragon, jesus.medina, samuel.molina}@uca.es

Abstract—Los enlaces dentro del análisis de conceptos formales son una herramienta formal que relaciona distintos contextos y, entre otras características, proporciona un método para su agregación. En este trabajo, introducimos una generalización de esta noción al ambiente multiadjunto y estudiamos cómo estos enlaces pueden usarse para agregar los retículos de conceptos mediante el producto cartesiano y la suma horizontal.

Index Terms—Análisis de conceptos formales, retículo de conceptos, enlaces, marco multiadjunto

I. INTRODUCCIÓN

En la realidad digital actual saturada de datos y caracterizada por un crecimiento exponencial de la información, es fundamental desarrollar herramientas rigurosas para su comprensión y tratamiento. Las bases de datos no solo almacenan información, sino que resultan ser la fuente principal para la extracción de conocimiento y la toma de decisiones. Estas tareas requieren una base matemática sólida que garantice su correcto funcionamiento y la transparencia y reproducibilidad de los procesos.

El análisis de conceptos formales (abreviado como FCA por sus siglas en inglés) fue introducido como una teoría matemática para estudiar la estructura jerárquica interna que existe en la información contenida en las bases de datos relacionales. De esta forma, el FCA trata estas bases de datos como un contexto: una terna compuesta de un conjunto de objetos, un conjunto de atributos y una relación binaria entre ellos. Además, los conceptos, que representan la unidad elemental de información dentro de un contexto, tienen la estructura de un retículo completo, lo que permite describir la organización jerárquica de la información de forma rigurosa.

Entre las generalizaciones de esta teoría al ambiente difuso queremos resaltar la basada en álgebras multiadjuntas [10], [11]. Esta teoría se basa en la noción de marco multiadjunto, una tupla formada por dos retículos completos, un conjunto parcialmente ordenado, y una lista de triples adjuntos (una terna compuesta por una conjunción y dos implicaciones difusas satisfaciendo la propiedad de adjunción [3], [4]) que determinan distintas formas de tratar la información. Los contextos se definen asociados a un marco multiadjunto y constan de dos conjuntos no vacíos de objetos y atributos,

una relación difusa entre ellos y una aplicación que asocia a cada par de objeto-atributo un triple adjunto del marco multiadjunto considerado. La elección de triples adjuntos para distintos pares es una de las características que otorga a esta teoría una mayor flexibilidad y versatilidad con respecto a otras extensiones difusas.

La agregación de información es fundamental a la hora de obtener conclusiones coherentes a partir de diferentes fuentes. Los enlaces se introdujeron dentro del FCA en [7] como un método de agregación de contextos. Varias generalizaciones se han propuesto para las diversas extensiones del FCA al ambiente difuso, incluyendo el basado en retículos residuados [2], [8], [9]. En este trabajo, introduciremos una generalización de esta noción, así como su aplicación para la agregación de contextos, dentro del ambiente basado en álgebras multiadjuntas.

II. ENLACES EN EL MARCO MULTIADJUNTO

Comenzamos esta sección introduciendo la noción de enlace entre dos contextos disjuntos (aquellos que no comparten ningún objeto ni atributo) en el ambiente multiadjunto. Observar que la notación $\phi_{x,t} \in L_2^{O_1}$ se refiere al conjunto difuso que toma el valor t en x , y $\perp_2$ para los demás valores. El conjunto difuso $\phi_{a,s} \in L_1^{P_2}$ se define de manera análoga.

Definición 1. Dados dos contextos disjuntos $\mathbb{M}_1$ y $\mathbb{M}_2$ asociados al mismo marco multiadjunto, un *enlace multiadjunto* de $\mathbb{M}_1$ a $\mathbb{M}_2$ es un contexto $\mathbb{M}_{12} = (O_1, P_2, R_{12}, \sigma_{12})$, donde $R_{12}: O_1 \times P_2 \rightarrow Q$ es una relación Q -difusa, $\sigma_{12}: O_1 \times P_2 \rightarrow \{1, \dots, n\}$ es una aplicación que asigna triples adjuntos, y tal que

- $\phi_{x,t}^{\uparrow \sigma_{12}}$ es una intensión de $\mathbb{M}_2$ y
- $\phi_{a,s}^{\downarrow \sigma_{12}}$ es una extensión de $\mathbb{M}_1$,

para todo objeto $x \in O_1$, atributo $a \in P_2$, y valores $t \in L_2$, $s \in L_1$.

El siguiente resultado muestra cómo la definición anterior extiende las definiciones de enlaces clásica y dentro de los retículos residuados.

Proposición 2. Sean $\mathbb{M}_1$ y $\mathbb{M}_2$ dos contextos disjuntos asociados al mismo marco multiadjunto. El contexto $\mathbb{M}_{12}$ es un *enlace multiadjunto* si y solo si las extensiones de $\mathbb{M}_{12}$ son extensiones de $\mathbb{M}_1$ y las intensiones de $\mathbb{M}_{12}$ son intensiones de $\mathbb{M}_2$.

*Partially supported by the project PID2022-137620NB-I00 funded by MICIU/AEI/10.13039/501100011033 and FEDER, UE, by the grant TED2021-129748B-I00 funded by MCIN/AEI/10.13039/501100011033 and European Union NextGenerationEU/PRTR, and by the project PR2023-009 funded by the University of Cádiz.

El resultado anterior indica que un enlace multiadjunto contiene solamente información de los contextos originales, luego la preserva. Debido a la condición que debe satisfacer el contexto $\mathbb{M}_{12}$ para ser enlace multiadjunto, es importante estudiar condiciones en las relaciones y aplicaciones que forman enlaces multiadjuntos. En este trabajo nos centramos en dos clases particulares de enlaces, que definimos a continuación, para las que asumiremos que el conjunto parcialmente ordenado Q del marco multiadjunto considerado está acotado inferiormente por $\perp$ y superiormente por $\top$.

Dados dos contextos disjuntos $\mathbb{M}_1$ y $\mathbb{M}_2$ asociados al mismo marco multiadjunto, podemos definir los contextos $\mathbb{M}_{12}^\top = (O_1, P_2, R_{12}^\top, \sigma_{12})$ y $\mathbb{M}_{12}^\perp = (O_1, P_2, R_{12}^\perp, \sigma_{12})$, donde $R_{12}^\top(x, a) = \top$ y $R_{12}^\perp(x, a) = \perp$, para todo $(x, a) \in O_1 \times P_2$, y $\sigma_{12}: O_1 \times P_2 \rightarrow \{1, \dots, n\}$ es una aplicación arbitraria. El contexto $\mathbb{M}_{12}^\top$ siempre es un enlace multiadjunto entre los contextos, mientras que para $\mathbb{M}_{12}^\perp$ no es cierto en general. Existen condiciones suficientes que se basan en conjuntos sin divisores de cero para garantizar que este contexto sea un enlace multiadjunto [1].

III. SUMA DE CONTEXTOS BASADA EN ENLACES MULTIADJUNTOS

En esta sección, fijaremos un marco multiadjunto $(L_1, L_2, Q, \&_1, \dots, \&_n)$, donde Q estará acotado por $\perp$ y $\top$, y dos contextos $\mathbb{M}_1 = (O_1, P_2, R_1, \sigma_1)$, $\mathbb{M}_2 = (O_2, P_2, R_2, \sigma_2)$ asociados a este marco multiadjunto. Es claro que aunque hayamos definido los enlaces multiadjuntos de $\mathbb{M}_1$ a $\mathbb{M}_2$, naturalmente podemos también considerar enlaces multiadjuntos de $\mathbb{M}_2$ a $\mathbb{M}_1$. De esta forma, consideraremos también dos contextos adicionales $\mathbb{M}_{12} = (O_1, P_2, R_{12}, \sigma_{12})$ y $\mathbb{M}_{21} = (O_2, P_1, R_{21}, \sigma_{21})$. La siguiente definición introduce el método de agregación de contextos basado en enlaces multiadjuntos.

Definición 3. La suma de los contextos $\mathbb{M}_1$ y $\mathbb{M}_2$ vía $\mathbb{M}_{12}$ y $\mathbb{M}_{21}$ es el contexto $\mathbb{M} = (O, P, R, \sigma)$, donde $O = O_1 \cup O_2$, $P = P_1 \cup P_2$ y la relación R y la aplicación σ vienen definidas por la Tabla I.

La definición anterior admite la posibilidad de que los contextos $\mathbb{M}_{12}$ y $\mathbb{M}_{21}$ no sean enlaces multiadjuntos. Sin embargo, si pretendemos que la información de los contextos originales no se altere, debemos considerar la suma de contextos solo cuando $\mathbb{M}_{12}$ y $\mathbb{M}_{21}$ sean enlaces multiadjuntos. En particular, dado que $\mathbb{M}_{12}^\top$ y $\mathbb{M}_{21}^\top$ son siempre enlaces multiadjuntos, podemos considerar la suma de los contextos usando estos enlaces multiadjuntos, de lo que se obtiene el producto cartesiano de sus retículos de conceptos [5], como enunciamos en el siguiente resultado.

Teorema 4. Si $\mathbb{M}$ es la suma de los contextos $\mathbb{M}_1$ y $\mathbb{M}_2$ vía $\mathbb{M}_{12}^\top$ y $\mathbb{M}_{21}^\top$, entonces

$$\mathcal{M} \cong \mathcal{M}_1 \times \mathcal{M}_2$$

donde $\mathcal{M}$, $\mathcal{M}_1$ y $\mathcal{M}_2$ son los retículos de conceptos de $\mathbb{M}$, $\mathbb{M}_1$ y $\mathbb{M}_2$, respectivamente.

R	P_1	P_2	σ	P_1	P_2
O_1	R_1	R_{12}	O_1	σ_1	σ_{12}
O_2	R_{21}	R_2	O_2	σ_{21}	σ_2

TABLE I

LA RELACIÓN R Y LA APLICACIÓN σ DE LA SUMA DE CONTEXTOS EN LA DEFINICIÓN 3

Por otro lado, los contextos $\mathbb{M}_{12}^\perp$ y $\mathbb{M}_{21}^\perp$ no siempre son enlaces multiadjuntos. No obstante, cuando lo son, como por ejemplo bajo las condiciones presentadas en [1], la suma de los contextos coincide con la suma horizontal de los retículos de conceptos originales [6].

Teorema 5. Sea $\mathbb{M}$ la suma de los contextos $\mathbb{M}_1$ y $\mathbb{M}_2$ vía $\mathbb{M}_{12}^\perp$ y $\mathbb{M}_{21}^\perp$. Si $\mathbb{M}_1$ y $\mathbb{M}_2$ tienen al menos dos conceptos cada uno y $\mathbb{M}_{12}^\perp$ y $\mathbb{M}_{21}^\perp$ son enlaces multiadjuntos donde σ_{12} y σ_{21} asignan conjuntos sin divisores de cero, entonces

$$\mathcal{M} \cong \mathcal{M}_1 \oplus \mathcal{M}_2$$

donde $\mathcal{M}$, $\mathcal{M}_1$ y $\mathcal{M}_2$ son los retículos de conceptos de $\mathbb{M}$, $\mathbb{M}_1$ y $\mathbb{M}_2$, respectivamente.

IV. CONCLUSIONES Y TRABAJO FUTURO

Los resultados expuestos en este trabajo reflejan dos métodos de agregación de contextos mediante enlaces multiadjuntos que se basan en el producto cartesiano y la suma horizontal de retículos de conceptos. Aunque solamente hemos tratado dos clases de contextos, definidos por las relaciones constantemente $\perp$ y constantemente $\top$, para contextos arbitrarios suelen existir otros enlaces multiadjuntos que dan lugar a otras construcciones distintas para la agregación de los contextos. En el futuro estudiaremos cómo la elección arbitraria de enlaces multiadjuntos influye en la estructura del retículo de conceptos asociado al contexto agregado, así como métodos para obtener y elegir los enlaces multiadjuntos en función de información externa.

REFERENCES

- [1] R. G. Aragón, J. Medina, and S. Molina-Ruiz. The notion of bond in the multi-adjoint concept lattice framework. In *Advances in Artificial Intelligence*, pages 243–253, Cham, 2024. Springer Nature Switzerland.
- [2] R. Bělohávek. Lattice generated by binary fuzzy relations (extended abstract). In *4th Intl Conf on Fuzzy Sets Theory and Applications*, page 11, 1998.
- [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. Multi-adjoint algebras versus non-commutative residuated structures. *International Journal of Approximate Reasoning*, 66:119–138, 2015.
- [5] B. A. Davey and H. A. Priestley. *Introduction to Lattices and Order*. Cambridge University Press, 2 edition, 2002.
- [6] M. El-Zekey, J. Medina, and R. Mesiar. Lattice-based sums. *Information Sciences*, 223:270–284, 2013.
- [7] B. Ganter and R. Wille. *Formal Concept Analysis: Mathematical Foundations*. Springer, 2 edition, 2024.
- [8] J. Konecny and O. Krídlo. On biconcepts in formal fuzzy concept analysis. *Information Sciences*, 375:16–29, 2017.
- [9] O. Krídlo, S. Krajčí, and M. Ojeda-Aciego. L-bonds vs extents of direct products of two L-fuzzy contexts. In *Proceedings of Concept Lattices and Applications (CLA)*, pages 70–79, 2010.
- [10] J. Medina and M. Ojeda-Aciego. Multi-adjoint t-concept lattices. *Information Sciences*, 180(5):712–725, 2010.
- [11] 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.

Contextos formales fibrados

Víctor Ramos-González
Dpto. CC Computación e IA
Universidad de Sevilla
Sevilla, España
vramos1@us.es

Joaquín Borrego-Díaz
Dpto. CC Computación e IA
Universidad de Sevilla
Sevilla, España
jborrego@us.es

Fco. Félix Lara-Martín
Dpto. CC de la Computación e IA
Universidad de Sevilla
Sevilla, España
fflara@us.es

Abstract—Este trabajo introduce las denominadas extensiones fibradas de contextos formales difusos en el marco del Análisis de Conceptos Formales (FCA). Dicha estructura que generaliza la FCA fuzzy tradicional al incorporar factores de atenuación para objetos y atributos, permitiendo definir conceptos atenuados, sin prefijar umbrales de corte. La motivación principal es abordar el desafío de alinear los grandes modelos de lenguaje Grandes (LLMs) con estructuras ontológicas humanas.

Index Terms—Análisis de conceptos formales fuzzy, IA generativa, ontologías, alineación, IA generativa

I. INTRODUCTION

El Análisis de Conceptos Formales Difusos (FFCA, del inglés *Fuzzy Formal Concept Analysis*) constituye una extensión madura del Análisis de Conceptos Formales (FCA) [1] clásico, en el que los atributos se modelan como predicados difusos y la relación objeto-atributo está inducida por sus funciones de pertenencia [2]–[4]. Existen diversas propuestas sobre qué podría considerarse un contexto o concepto formal difuso [2], [5], [6], y también desarrollos orientados a aplicaciones específicas (e.g. [7]–[9]) y el análisis de su robustez para el aprendizaje de conceptos [10].

A pesar de la diversidad de enfoques, la gran mayoría modelos comparten una característica común en su vertiente algorítmica: dependen de umbrales discretos aplicados sobre las funciones de pertenencia, lo cual implica una pérdida inevitable de información granular. En dominios donde los valores de pertenencia varían dinámicamente —como en entornos de revisión de creencias— resulta necesario disponer de modelos más generales, capaces de preservar la fineza semántica inherente a dichas variaciones. Es en este contexto en el que proponemos un marco unificado mediante una construcción que introduce dos mecanismos complementarios: (1) factores de atenuación para objetos, que modulan su contribución relativa, y (2) truncamiento de atributos, permitiendo su adaptación a objetos atenuados. Formalmente, la extensión fibrada genera un nuevo contexto donde tanto objetos como atributos adquieren dimensiones adicionales que preservan la continuidad del espacio difuso original.

La motivación detrás de esta propuesta proviene del estudio de un problema crítico en IA generativa: la evaluación de la desalineación semántica entre el conocimiento emergente en modelos de lenguaje masivos (LLMs) y las estructuras ontológicas consolidadas. Nuestro método lo aborda en tres fases:

Introspección guiada: Inducción de un contexto formal difuso a partir de las respuestas del LLM, cuantificando sus grados de certeza sobre pares objeto-atributo.

Extensión fibrada: Construcción del contexto extendido y extracción de sus conceptos atenuados, que capturan relaciones semánticas graduales.

Validación ontológica: Comparación sistemática con una ontología de referencia (en nuestro caso, tomaremos la de Sowa [11]), utilizando métricas de similitud estructural.

Este enfoque permite identificar discrepancias conceptuales no triviales. Ofrece, de este modo, un marco cuantitativo para el diagnóstico sistemático de desalineaciones semánticas entre el conocimiento representado en los LLMs y las ontologías formales de referencia. Así, se propone una etapa adicional centrada en el refinamiento de las creencias del modelo —mediante técnicas de *prompting*—, orientada a favorecer la emergencia de una estructura conceptual más coherente con la ontología considerada.

II. ANÁLISIS DE CONCEPTOS FORMALES (DIFUSOS)

El formato de información utilizado en FCA está organizado en lo que se conoce como *Contexto Formal*, un conjunto de tres elementos $\mathbb{K} = (G, M, I)$, donde G es un conjunto de *objetos*, M es un conjunto no vacío de *atributos* y $I \subseteq G \times M$. Un contexto formal induce dos operadores. Dados $A \subseteq G$ y $B \subseteq M$:

$$A^\downarrow := \{m \in M \mid \forall g \in A, (g, m) \in I\}$$

$B^\uparrow := \{g \in G \mid \forall m \in B, (g, m) \in I\}$ Un concepto formal es un par $C = (A, B)$ de conjuntos de objetos y atributos, (extensión e intención) con $A^\downarrow = B$ y $B^\uparrow = A$. El conjunto de conceptos de $\mathbb{K}$ tiene una estructura de retículo completo mediante la relación de *subconcepto*, $\leq$ (cf. [1]).

Un contexto formal difuso $\mathbb{K} = (G, M, I)$ se diferencia fundamentalmente en que I es una relación difusa sobre $G \times M$ (con función de pertenencia μ_I). La relación difusa I induce una función de pertenencia difusa para cada atributo $m \in M$, definida por $\mu_m(g) := \mu_I(g, m)$. Existen varias formas de definir conceptos en FFCA (e.g. [7], [8]).

Dado $t \in [0, 1]$ un umbral, el operador de derivación difuso sobre objetos se define como $A^\downarrow = \{m \in M : \forall g \in A, \mu_I(g, m) \geq t\}$, y sobre atributos como $B^\uparrow = \{g \in G \mid \forall m \in B : \mu_I(g, m) \geq t\}$. Un concepto formal difuso C es un par $(A, \varphi(B))$ tal que $A^\downarrow = B$ y $B^\uparrow = A$, donde $\varphi(B)$ es el predicado difuso cuya función de pertenencia es $\mu_{\varphi(B)}(g) := \min_{m \in B} \mu_m(g)$. De esta manera, se obtiene una estructura análoga al retículo de conceptos para contextos formales clásicos.

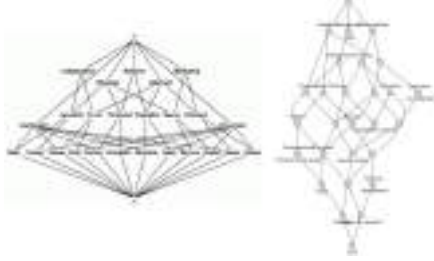


Fig. 1. Izquierda: *KR Ontology*, Derecha: retículo de conceptos asociado al contexto para el umbral 0.51

III. EXTENSIÓN FIBRADA

La extensión fibrada de un contexto formal difuso $\mathbb{K} = (G, M, I)$ es un contexto $\mathbb{K}_f = (G_f, M_f, I_f)$ donde $G_f = G \times [0, 1]$, $M_f = M \times [0, 1]$ y $\mu_{(m, \alpha)}(g, \beta) = \min(\beta, \alpha \cdot \mu_m(g))$.

A los objetos de esta extensión se les denomina objetos atenuados y a los atributos, atributos truncados, mientras que al parámetro correspondiente se le llama grado de atenuación. El conjunto G se identifica con $G \times \{1\}$, y M con $M \times \{1\}$.

Todo objeto atenuado $(g, \beta) \in G_f$, es un predicado difuso sobre M_f , con pertenencia $\mu_{(g, \beta)}(m, \alpha) := \mu_{(m, \alpha)}(g, \beta)$.

Notaremos con X, Y (posiblemente con subíndices) conjuntos de objetos o atributos de la extensión fibrada. Un conjunto se dice *representativo* si no contiene elementos con atenuación 0 y no hay dos elementos con el mismo objeto de G . Trabajaremos siempre con conjuntos representativos; usaremos el operador $*$ para eliminar de un conjunto todos los elementos con atenuación 0. La *contención fibrada* entre objetos/atributos de $\mathbb{K}_f$ se define como:

$$X \subseteq_f Y \iff \forall (c, \gamma) \in X \exists (c, \delta) \in Y \text{ tal que } \gamma \leq \delta.$$

En la extensión fibrada, se definen los operadores extensión/intención sobre conjuntos representativos; para $X \subseteq G_f$, y para $Y \subseteq M_f$:

$$X^\downarrow = \left\{ (m, \sup_{(g, \beta) \in X} \left\{ \frac{\beta}{\mu_m(g)} \right\}) \mid \forall (g_i, \beta_i) \in X : \mu_m(g_i) \geq \beta_i \wedge \mu_m(g_i) > 0 \right\}$$

$$Y^\uparrow = \left\{ (g, \inf_{(m, \alpha) \in Y} \{\mu_m(g) \cdot \alpha\}) : g \in G \right\}$$

Este operador satisface las propiedades básicas de una conexión de Galois respecto a $\subseteq_f$. Se puede definir un concepto fibrado (X, Y) a partir de esos operadores.

IV. DESALINEACIÓN SEMÁNTICA EN LLM

Para resolver el problema de la alineación humano-LLM, existen múltiples aproximaciones: guardarrailes, aprendizaje por refuerzo con apoyo humano, o inducción de procesos introspectivos [12], [13]. Nuestro objetivo es extraer una organización conceptual del conocimiento implícito en el modelo.

Con tal fin, deseamos evaluar la fidelidad de la información introspectiva del modelo con respecto a una ontología bien conocida: en nuestro caso, la ontología de alto nivel *KR Ontology*



Fig. 2. Tres prompts utilizados en la introspección con chatgpt

de J.F. Sowa, O_S , [11] (Fig. 1). Al tratarse de una ontología de alto nivel, utilizaremos un conjunto de objetos heterogéneos, pero no representativos de todas las clases de O_S (en la Fig. 2 se muestran algunos de los *prompts* utilizados). Un concepto fibrado es

$\{(Fot., 1.0), (Lápiz, 0.4), \dots, (Todo, 1.0)\}, \{(PC\text{-}Indep, 1.0)\}$ lo que muestra que el modelo no respeta la organización ontológica de O_S

AGRADECIMIENTOS

Proyecto PID2023-147198NB-I00 financiado por MICIU/AEI/10.13 039/501100011033 y por FEDER, UE.

REFERENCES

- [1] B. Ganter and R. Wille, *Formal Concept Analysis: Mathematical Foundations*, 1st ed. Secaucus, NJ, USA: Springer-Verlag New York, Inc., 1997.
- [2] R. Bělohlávek, "Fuzzy Galois connections," *Mathematical Logic Quarterly*, vol. 50, no. 5-6, pp. 391–406, 2004, doi or other optional fields can be added here.
- [3] J. Poelmans, D. I. Ignatov, S. O. Kuznetsov, and G. Dedene, "Fuzzy and rough formal concept analysis: a survey," *Int. Journal of General Systems*, vol. 43, no. 2, pp. 105–134, 2014.
- [4] P. K. Singh and C. A. Kumar, "A note on computing the crisp order context of a fuzzy formal context for knowledge reduction," *J. Inf. Process. Syst.*, vol. 11, no. 2, pp. 184–204, 2015.
- [5] P. Butka, J. Pócs, and J. Pócs, "Representation of fuzzy concept lattices in the framework of classical FCA," *J. Appl. Math.*, vol. 2013, pp. 236725:1–236725:7, 2013.
- [6] J. Macko, "User-friendly fuzzy FCA," in *Formal Concept Analysis, 11th Int. Conf., ICFCA 2013, Proceedings*, ser. LNCS, P. Cellier, F. Distel, and B. Ganter, Eds., vol. 7880. Springer, 2013, pp. 156–171.
- [7] Q. T. Tho, S. C. Hui, A. C. M. Fong, and Tru Hoang Cao, "Automatic fuzzy ontology generation for Semantic Web," *IEEE Trans. Knowl. Data Eng.*, vol. 18, no. 6, pp. 842–856, 2006.
- [8] Y. Zhai and D. Li, "Knowledge structure preserving fuzzy attribute reduction in fuzzy formal context," *Int J Approx Reason*, vol. 115, pp. 209 – 220, 2019.
- [9] Wu, X. et al., "A fuzzy formal concept analysis-based approach to uncovering spatial hierarchies among vague places extracted from user-generated data," *Int. J. Geogr. Inf. Sci.*, vol. 33, no. 5, pp. 991–1016, 2019.
- [10] G. A. Aranda-Corral, J. Borrego-Díaz, and J. Galán-Páez, "Concept learning consistency under three-way decision paradigm," *International Journal of Machine Learning and Cybernetics*, vol. 13, pp. 2977–2999, 2022.
- [11] J. F. Sowa, *Knowledge Representation: Logical, Philosophical, and Computational Foundations*. Brooks Cole Publishing Co., 2000.
- [12] I. M. Comsa and M. Shanahan, "Does it make sense to speak of introspection in large language models?" 2025. [Online]. Available: <https://arxiv.org/abs/2506.05068>
- [13] Y. Jiang, Y. Xiong, Y. Yuan, C. Xin, W. Xu, Y. Yue, Q. Zhao, and L. Yan, "Pag: Multi-turn reinforced llm self-correction with policy as generative verifier," 2025. [Online]. Available: <https://arxiv.org/abs/2506.10406>

Definición de sistemas de inferencia basados en la f -inclusión

Carolina Díaz-Montarroso
Dept. de Matemáticas
Universidad de Cádiz
Cádiz, España
carolina.diaz@uca.es

Nicolás Madrid
Dept. de Matemáticas
Universidad de Cádiz
Cádiz, España
nicolas.madrid@uca.es

Eloísa Ramírez-Poussa
Dept. de Matemáticas
Universidad de Cádiz
Cádiz, España
eloisa.ramirez@uca.es

Resumen—En trabajos recientes se ha demostrado que el f -índice de inclusión puede emplearse para modelizar un Modus Ponens Generalizado. En este trabajo se define un Sistema de Inferencia Difuso basado en el f -índice de inclusión. Este sistema está fundamentado en una lógica formal, lo cual permite obtener una serie de resultados que culminan en un teorema de correctitud. Finalmente, se presenta un resultado de aproximación universal de funciones continuas empleando este sistema de inferencia.

Palabras Clave—lógica difusa, f -índice de inclusión, sistemas de inferencia difusos, medidas de inclusión.

I. INTRODUCCIÓN

Los sistemas de inferencia difusos (SIF) son mecanismos ampliamente utilizados en la ingeniería que abarcan desde problemas de control hasta la representación de conocimiento. Los SIF propuestos hasta la actualidad se dividen en dos vertientes principales. Por un lado, se encuentran los sistemas difusos basados en relaciones difusas, entre los que se sitúa el clásico sistema de inferencia de Mamdani [1] y aquellos que usan la regla composicional de inferencia [2]. Por otro lado, están los sistemas de inferencia formados por reglas que generan valores *crisp* que luego se combinan mediante operadores de agregación para obtener una inferencia final. Dentro de este grupo, los sistemas más conocidos y utilizados son el de Takagi–Sugeno [3] y el de Tsukamoto [4].

A pesar de que estos sistemas destacan por su utilidad práctica, la mayoría presenta una característica contraproducente desde un punto de vista teórico: no está basada en un sistema lógico formal. Estos SIF, a pesar del empleo que hacen de los operadores lógicos, carecen de una sintaxis y una semántica que permitan establecer una definición rigurosa de consecuencia lógica dentro de ellos, lo cual impide demostrar la correctitud de los procesos de inferencia empleados.

En [5] se propone un SIF fundamentado en una lógica formal inspirada en la Lógica Descriptiva Difusa (LDD) [6]. Dentro de este sistema formal, podemos definir el concepto

de consecuencia lógica y definir un Modus Ponens Generalizado (MPG) basado en la f -inclusión [7] entre conjuntos difusos [8], permitiendo realizar inferencias cuando la entrada del SIF no coincida con el antecedente de ninguna de las reglas. Los resultados mostrados en [9] y [10] justifican el uso de la f -inclusión como una implicación para modelar una regla de inferencia difusa y como operador para representar la similaridad entre dos conjuntos difusos.

En este trabajo recordamos la construcción de este tipo de SIF y presentamos un resultado que permite afirmar que el SIF propuesto puede utilizarse como un aproximador universal de funciones continuas.

En la Sección II se presenta el SIF basado en la f -inclusión, prestando especial atención a la construcción de su motor de inferencia, y se finaliza con un resultado de aproximación de funciones satisfecho por este SIF. En la Sección III se resumen las conclusiones de este trabajo y las líneas de investigación futuras.

II. SISTEMA DE INFERENCIA DIFUSO BASADO EN LA f -INCLUSIÓN

II-A. Construcción del sistema de inferencia difuso

Para definir el sistema de inferencia, se parte de un marco, denotado generalmente por $(X, Y, \mathcal{A}_X, \mathcal{B}_Y)$, donde X es el universo de entrada, Y el universo de salida y $\mathcal{A}_X = \{A_i\}_{i \in I}$ y $\mathcal{B}_Y = \{B_j\}_{j \in J}$ son dos particiones difusas definidas respectivamente sobre cada universo. Dentro de este marco, se define la base de conocimiento del sistema de inferencia, denotada por Γ , que está compuesta por reglas de la forma:

$$\langle A_i \rightarrow B_j \ ; f_{ij} \rangle, \quad (1)$$

donde $A_i \in \mathcal{A}_X$, $B_j \in \mathcal{B}_Y$ y $f_{ij}: [0, 1] \rightarrow [0, 1]$ es una función que satisface las siguientes propiedades:

- es monótona,
- $f_{ij}(a) \leq a$ para todo $a \in [0, 1]$ y
- existe una función (única) $\bar{f}_{ij}$ tal que $(f_{ij}, \bar{f}_{ij})$ forma un par adjunto.

El conjunto de funciones en $[0, 1]^{[0, 1]}$ que satisfacen estas condiciones se denota por $\mathcal{G}$. La semántica se define de la siguiente forma.

Partially supported by the projects PID2022-140630NB-I00 and PID2022-137620NB-I00 funded by MICIU/AEI/10.13039/501100011033 and FEDER, UE, by the grant TED2021-129748B-I00 funded by MCIN/AEI/10.13039/501100011033 and European Union NextGenerationEU/PRTR. Partially supported by Plan Propio-UCA 2025-2027.

Definición 2.1: [5] Una interpretación es un subconjunto $I \subseteq X \times Y$. Decimos que una interpretación satisface una regla de $\langle A_i \rightarrow B_j \ ; \ f_{ij} \rangle \in \Gamma$ si y solo si

$$f_{ij}(A_i(x)) \leq B_j(y)$$

para todo $(x, y) \in I$. Diremos que una interpretación es un modelo de Γ si satisface todas las reglas de Γ .

En este SIF, las entradas aparecen en forma de conjuntos difusos definidos sobre uno de los universos que conforman el marco; por simplicidad, asumiremos que todas las entradas son conjuntos difusos sobre el universo X .

Definición 2.2: [5] Sea Γ una base de conocimiento y $A \in \mathcal{F}(X)$ una entrada en el marco $(X, Y, \mathcal{A}_X, \mathcal{B}_Y)$. Un modelo de $\Gamma \cup \{A\}$ es un conjunto difuso $M \in \mathcal{F}(X \times Y)$ tal que para todo $\alpha > 0$, el α -corte de M , $M^\alpha = \{(x, y) \in X \times Y \mid M(x, y) \geq \alpha\}$, satisface:

- M^α es un modelo de Γ y
- $M^\alpha \subseteq \{(x, y) \in X \times Y \mid A(x) \geq \alpha\}$.

De manera similar, se puede definir el modelo de $\Gamma \cup \{B\}$ cuando $B \in \mathcal{F}(Y)$. Esto nos permite establecer la definición de consecuencia lógica en nuestro marco a partir de una base de conocimiento y de una entrada.

Definición 2.3: [5] Sea Γ una base de conocimiento y $A \in \mathcal{F}(X)$ una entrada en un marco $(X, Y, \mathcal{A}_X, \mathcal{B}_Y)$. Se dice que un conjunto difuso es consecuencia lógica de Γ y de A si y solo si todo modelo de $\Gamma \cup \{A\}$ es modelo de $\Gamma \cup \{B\}$.

Se ha demostrado que el conjunto de consecuencias tiene estructura de retículo completo con el orden de Zadeh entre conjuntos difusos. Dentro de esta estructura, se pretende buscar el mínimo elemento contenido en este conjunto.

II-B. Obtención de consecuencias a través del MPG

La aplicación del MPG basado en la f -inclusión [8] permite el uso del siguiente motor de inferencia: dada una entrada $A \in \mathcal{F}(X)$, calculamos el f -índice de inclusión restringido a $\mathcal{G}$ de A en cada $A_i \in \mathcal{A}_X$, denotado por $\text{Inc}_{\mathcal{G}}(A, A_i)$ (ver [9]), y definimos el conjunto difuso en $\mathcal{F}(Y)$, al que llamaremos salida de $\Gamma \cup \{A\}$, de la siguiente forma:

$$B(y) = \inf_{\substack{A_i \in \mathcal{A}_X, \\ B_j \in \mathcal{B}_Y}} \{ \overline{\text{Inc}}_{\mathcal{G}}(A, A_i) \circ \bar{f}_{ij}(B_j(y)) \}, \quad (2)$$

donde $\overline{\text{Inc}}_{\mathcal{G}}(A, A_i)$ es el adjunto por la derecha de $\text{Inc}_{\mathcal{G}}(A, A_i)$. El siguiente resultado muestra que esta salida es una consecuencia lógica de $\Gamma \cup \{A\}$.

Corolario 2.1: Sea Γ una base de conocimiento y A una entrada en un marco $(X, Y, \mathcal{A}_X, \mathcal{B}_Y)$. Entonces el conjunto difuso $B \in \mathcal{F}(Y)$ dado en (2) es una consecuencia de $\Gamma \cup \{A\}$.

Recientemente hemos demostrado el siguiente resultado teórico, que permite afirmar que podemos aproximar funciones continuas con la precisión deseada a través de SIF basados en la f -inclusión.

Teorema 2.1: Sea $f: D \subseteq \mathbb{R}^n \rightarrow \mathbb{R}^m$ una función continua en todo su dominio. Para todo $\varepsilon > 0$ existe un marco $(\mathbb{R}^n, \mathbb{R}^m, \mathcal{A}_{\mathbb{R}^n}, \mathcal{B}_{\mathbb{R}^m})$ y una base de conocimiento Γ tal que para toda entrada $A = \{x_0\}$, con $x_0 \in D$, se tiene

$$\|\widetilde{x_0} - f(x_0)\| < \varepsilon, \quad (3)$$

donde $\widetilde{x_0}$ es cualquier valor del soporte de la salida del sistema de inferencia obtenida al aplicar el motor de inferencia a Γ y a la entrada A (ver (2)).

III. CONCLUSIONES Y TRABAJO FUTURO

En este trabajo se ha recordado la construcción de un sistema de inferencia difuso basado en la f -inclusión entre conjuntos difusos. La principal ventaja que presenta este sistema de inferencia desde el punto de vista teórico es su fundamento basado en una lógica formal, lo cual permite demostrar la correctitud del sistema. Asimismo, se puede demostrar la validez de este SIF para aproximar funciones continuas. No obstante, quedan pendientes de estudio una serie de líneas de investigación relacionadas con este SIF:

- En primer lugar, desde un punto de vista teórico, resulta de especial interés investigar resultados de completitud para este sistema de inferencia.
- En segundo lugar, tenemos interés en aplicar este SIF a problemas reales, lo que implica abordar diferentes tareas, como el desarrollo de un proceso de aprendizaje para reglas y un estudio de sensibilidad a los procesos de fuzzificación y defuzzificación.

REFERENCIAS

- [1] Mamdani, E.; Assilian, S. An experiment in linguistic synthesis with a fuzzy logic controller. *Int. J.-Man-Mach. Stud.* **1975**, *7*, 1–13. [https://doi.org/10.1016/S0020-7373\(75\)80002-2](https://doi.org/10.1016/S0020-7373(75)80002-2).
- [2] Štěpnička, M.; Jayaram, B.; Su, Y. A short note on fuzzy relational inference systems. *Fuzzy Sets Syst.* **2018**, *338*, 90–96. <https://doi.org/10.1016/j.fss.2017.08.006>.
- [3] Sugeno, M. *Industrial Applications of Fuzzy Control*; Elsevier Science Ltd.: Amsterdam, The Netherlands, 1985.
- [4] Gupta, P.K.; Muhuri, P.K. Extended Tsukamoto's inference method for solving multi-objective linguistic optimization problems. *Fuzzy Sets Syst.* **2019**, *377*, 102–124. <https://doi.org/10.1016/j.fss.2019.02.022>.
- [5] Díaz-Montaroso, C.; Madrid, N.; Ramírez-Poussa, E. Correctness of Fuzzy Inference Systems Based on f -Inclusion. *Mathematics* **2025**, *13*, 1897. <https://www.mdpi.com/2227-7390/13/11/1897>.
- [6] Straccia, U. A fuzzy description logic. In *Proceedings of the AAAI/IAAI*, 1998, Madison, WI, USA, 26–30 July 1998; pp. 594–599.
- [7] Madrid, N.; Ojeda-Aciego, M.; Perfiljeva, I. f -inclusion indexes between fuzzy sets. In *Proceedings of the Conference of the International Fuzzy Systems Association and the European Society for Fuzzy Logic and Technology*, Riga, Latvia, 21–25 July 2025; pp. 1528–1533. <https://doi.org/10.2991/ifs-eusflat-15.2015.217>.
- [8] Díaz-Montaroso, C.; Madrid, N.; Ramírez-Poussa, E. Towards a Generalized Modus Ponens Based on the f -Index of Inclusion. In *Proceedings of the Conceptual Knowledge Structures: First International Joint Conference, CONCEPTS 2024*, Cádiz, Spain, 9–13 September 2024; pp. 36–48. https://doi.org/10.1007/978-3-031-67868-4_3.
- [9] Madrid, N.; Ojeda-Aciego, M. The f -index of inclusion as optimal adjoint pair for fuzzy modus ponens. *Fuzzy Sets Syst.* **2023**, *466*, 108474. <https://doi.org/10.1016/j.fss.2023.01.009>.
- [10] Madrid, N.; Ojeda-Aciego, M. Functional degrees of inclusion and similarity between L -fuzzy sets. *Fuzzy Sets Syst.* **2020**, *390*, 1–22. <https://doi.org/10.1016/j.fss.2019.03.018>.

Influencia probabilística del número de votantes en el muestreo mediante el modelo de Mallows

1st Mario Villar
Departamento de Informática
Universidad de Oviedo
Oviedo, España
villarmario@uniovi.es

2nd Noelia Rico
Departamento de Informática
Universidad de Oviedo
Oviedo, España
noeliarico@uniovi.es

3rd Irene Díaz
Departamento de Informática
Universidad de Oviedo
Oviedo, España
sirene@uniovi.es

Abstract—La agregación de rankings es fundamental en la toma de decisiones, y existen métodos teóricos bien definidos cuya implementación no siempre puede realizarse de forma eficiente. Un aspecto crucial y frecuentemente desatendido en el desarrollo de algoritmos en el campo de la Teoría de la Elección Social Computacional es la evaluación justa de las implementaciones. Para evaluarlos, es habitual crear un conjunto de datos de preferencias sintético mediante el modelo de Mallows, que permite generar votaciones artificiales con distintos números de votantes y alternativas, usando un parámetro de dispersión para simular diferentes niveles de acuerdo entre los votantes. Varios estudios han mostrado experimentalmente que, a igual dispersión, variar los parámetros relativos al número de alternativas y votantes afecta a las propiedades de las votaciones generadas, lo que lleva a una variación también en la dificultad que requiere a los algoritmos resolver la agregación de las preferencias. A pesar de estos resultados experimentales, no existen estudios que establezcan un marco teórico para detallar cómo influye en la dificultad de resolver la agregación el hecho de variar el número de votantes durante la generación de perfiles. Actualmente estamos trabajando en cuantificar el impacto de la modificación de los parámetros en las votaciones, para lograr un mismo nivel de desacuerdo y dificultad de resolución en la obtención de votaciones sintéticas en forma de rankings con el modelo de Mallows. El objetivo es determinar las condiciones que permitan crear conjuntos de datos de votaciones comparables para evaluar sobre ellos algoritmos de agregación de rankings.

Index Terms—agregación de rankings, modelo de Mallows, probabilidad combinatoria

I. INTRODUCCIÓN

La agregación de preferencias permite tomar decisiones a partir de las opiniones individuales de varios votantes sobre un conjunto de alternativas. Una forma habitual de expresarlas es mediante rankings, donde cada votante ordena las opciones de mayor a menor preferencia. El conjunto de estos rankings conforma un perfil, que sirve de base para aplicar un método de agregación con el objetivo de alcanzar un consenso. El campo de la Teoría de la Elección Social se encarga de estudiar las técnicas que permiten la agregación de estos rankings [1], [2], cuyo objetivo es que el ranking consenso obtenido refleje los individuales de la mejor forma posible.

This research has been funded by the Government of Spain through project MCINN-23-PID2022-139886NB-I00 and by grant Project PAPI-24-TESIS-14 of the University of Oviedo.

La parte computacional de la Teoría de la Elección Social [3] estudia entre otras cosas la aplicación de los métodos de agregación de preferencias desde una perspectiva tecnológica, centrándose en la viabilidad del uso de los algoritmos de agregación de preferencias en entornos reales. Una línea dentro de este campo es el desarrollo de algoritmos eficientes, así como su evaluación sobre datos de elecciones con distintas características. Para realizar esta evaluación, existen trabajos en la literatura que recurren a utilizar conjuntos de datos reales, pero la carencia de estos lleva en muchas ocasiones a la generación de conjuntos de datos de votaciones sintéticos para poder incrementar la batería de pruebas sobre los algoritmos desarrollados [4], [5]. En este contexto, los dos modelos más utilizados para la generación de estas votaciones sintéticas son Impartial Culture y el modelo de Mallows [6], aunque el primero puede ser replicado como un caso particular (dispersión máxima) a través del segundo.

Una práctica habitual consiste en generar conjuntos de datos con diferentes parámetros del modelo de Mallows, mediante muestreos probabilísticos. Sin embargo, la mayoría de trabajos no discuten la posible influencia de variar estos parámetros sobre la dificultad de realizar una agregación de preferencias. Experimentalmente, sí hay varios estudios que han demostrado que estas variaciones pueden tener efecto en el rendimiento de los algoritmos [7], [8]. Por tanto, analizar cómo influye el número de votantes en el muestreo de datos es crucial para evaluar objetivamente los algoritmos de agregación de preferencias.

En lo restante de este documento, la Sección II detalla los conceptos teóricos base y la Sección III describe la línea de investigación que estamos desarrollando para estudiar cómo variar el número de votantes afecta a las probabilidades de muestreo de elecciones. El objetivo de esta investigación es obtener una metodología para crear un conjunto de datos en el que evaluar los algoritmos de agregación de preferencia de forma justa y la incertidumbre de estos escenarios.

II. EL MODELO DE MALLOWS

El modelo de Mallows [6] es un modelo probabilístico que puede ser utilizado para muestrear rankings mediante el modelo de inserción repetida [9]. Se basa en dos parámetros: la dispersión, que fija el nivel de acuerdo entre los votantes,

y un ranking central. Se entiende intuitivamente que, a menor dispersión, más cercanos al ranking central son los rankings muestreados.

El proceso para generar un ranking consiste en insertar, de forma iterativa, las alternativas del ranking central en un nuevo ranking inicialmente vacío. En cada paso, se toma una alternativa del ranking base y se inserta en una posición aleatoria del nuevo ranking. La probabilidad de inserción de una alternativa en cada posición del nuevo ranking depende de la dispersión del modelo, del orden de la alternativa en el ranking central y de la distancia entre dicha posición y la posición original en el ranking central.

Un perfil de rankings se obtiene tras usar el anterior proceso de forma independiente para tantos rankings como votantes se consideren. Aunque la probabilidad de obtener un ranking no depende del número de votantes, modificar el valor de este parámetro sí que experimentalmente parece afectar a la dificultad de obtener consenso en las votaciones y algunas otras propiedades [8], [10]. No obstante, no hemos encontrado estudios que definan de forma teórica las probabilidades de obtener un perfil determinado.

III. ESTUDIO SOBRE GENERACIÓN DE PERFILES

Debido al proceso explicado, el hecho de que cada ranking sea generado de forma independiente, habitualmente provoca que los autores asuman que el número de votantes utilizado no tendrá impacto en las propiedades de los perfiles generados. Sin embargo, como se ha mencionado anteriormente, resultados experimentales sí muestran variaciones [8].

Con el objetivo de determinar la influencia del número de votantes, es necesario definir la probabilidad de obtener un perfil de rankings concretos en base a sus rankings. Obtener esta probabilidad permitiría determinar la distribución de perfiles según se varía el número de votantes. Un aspecto relevante a analizar sería comparar, a medida que se incrementa el número de votantes, la probabilidad de obtener perfiles donde la mayoría de votantes seleccionan el mismo ranking con la probabilidad de obtener perfiles donde la mayoría de preferencias son diferentes.

Para el análisis mencionado en el anterior párrafo es necesario considerar los perfiles específicos. El número de posibles perfiles, dado un número de votantes n y un número de alternativas m , es el número de multiconjuntos de n elementos tomados de los posibles $m!$ rankings:

$$\binom{m!}{n} = \binom{n + m! - 1}{n}$$

Esto provoca que, en elecciones con números de votantes y alternativas elevados, el número de perfiles existentes no permita realizar un análisis pormenorizado de su distribución. Por ello, otro punto importante sería la definición de la probabilidad de muestrear perfiles que cumplan una cierta propiedad.

En este sentido, la proporción de rankings de un perfil en los que una alternativa aparece en una posición específica es una propiedad interesante. Por ejemplo, el rendimiento de

los métodos basados en Condorcet [2] podría beneficiarse si una alternativa tiende a aparecer en primera posición de los rankings más habitualmente al incrementar el número de votantes. Esto provocaría que la evaluación de algoritmos que implementen estos métodos esté sesgada por el número de votantes utilizado. Para evitarlo, sería necesario analizar la evolución, al aumentar el número de votantes, de la probabilidad de muestrear un perfil en el que un porcentaje de sus rankings presenten una alternativa determinada en una posición específica.

Para avanzar en el análisis de las preferencias generadas, nos centramos también en estudiar otras propiedades globales de los perfiles, más allá de la posición de una alternativa concreta. En particular, estamos analizando métricas relacionadas con los márgenes de mayoría en las comparaciones por pares, que reflejan el grado de consenso o conflicto entre las preferencias individuales. Estas propiedades permiten observar cómo evoluciona la estructura interna de los perfiles al modificar el número de votantes y podrían influir directamente en la dificultad que tienen los algoritmos de agregación para alcanzar un resultado.

Con todo esto, buscamos establecer una base más sólida para generar perfiles sintéticos con propiedades controladas, que permitan una evaluación más justa y homogénea de los algoritmos. Estudiar cómo varían estas propiedades al cambiar el número de votantes en la generación es indispensable para garantizar que las pruebas experimentales reflejen de forma adecuada la dificultad real de resolver los problemas de agregación de preferencias.

REFERENCES

- [1] W. Gaertner, *A Primer in Social Choice Theory*, revised edition Edition, Oxford University Press, 2009.
- [2] P. C. Fishburn, *The theory of social choice*, Princeton University Press, 1973.
- [3] F. Brandt, V. Conitzer, U. Endriss, J. Lang, A. Procaccia, *Handbook of Computational Social Choice*, Cambridge University Press, 2016. doi:10.1017/CBO9781107446984.
- [4] N. Boehmer, N. Schaar, Collecting, classifying, analyzing, and using real-world ranking data, in: *Proceedings of the 2023 International Conference on Autonomous Agents and Multiagent Systems*, 2023, pp. 1706–1715. doi:10.48550/arXiv.2204.03589.
- [5] N. Boehmer, P. Faliszewski, Łukasz Janeczko, A. Kaczmarczyk, G. Lisowski, G. Pierczyński, S. Rey, D. Stolicke, S. Szufa, T. Was, Guide to numerical experiments on elections in computational social choice, in: *Proceedings of the Thirty-Third International Joint Conference on Artificial Intelligence*, 2024, pp. 7962–7970. doi:10.24963/ijcai.2024/881.
- [6] C. L. Mallows, Non-null ranking models. i, *Biometrika* 44 (1957) 114–130. doi:10.2307/2333244.
- [7] N. Boehmer, P. Faliszewski, S. Kraicz, Properties of the mallows model depending on the number of alternatives: A warning for an experimentalist, in: *Proceedings of Machine Learning Research*, Vol. 202, 2023, pp. 2689–2711.
- [8] M. Villar, N. Rico, I. Díaz, Understanding data properties in the mallows model: impact of voter count variability, in: *14th Conference of the European Society for Fuzzy Logic and Technology*, 2025, accepted for publication.
- [9] J.-P. Doignon, A. Pekeč, M. Regenwetter, The repeated insertion model for rankings: Missing link between two subset choice models, *Psychometrika* 69 (2004) 33–54. doi:10.1007/BF02295838.
- [10] M. Villar, N. Rico, I. Díaz, Exploring condorcet properties in mallows model-generated preferences, in: *20th Conference of the Spanish Association for Artificial Intelligence*, 2024, pp. 65–70.

Generalizing a construction of non-strong fuzzy metrics from metrics and studying their induced topology

1st Olga Grigorenko

*Institute of Mathematics and Computer Science
University of Latvia
Riga, Latvia
olga.grigorenko@lu.lv*

2nd Juan-José Miñana

*Instituto de Investigación para la Gestión Integrada de Zonas Costeras
Universitat Politècnica de València
Gandia, Spain
juamiapr@mat.upv.es*

3rd Simona Talia

*Departamento de Matemática Aplicada
Universitat Politècnica de València
València, Spain
stalia1@doctor.upv.es*

Abstract—The problem of obtaining new examples of fuzzy metrics has gained interest in the literature, as this type of measurement has proven to be very useful in engineering problems. In this context, different works have addressed the problem of deriving fuzzy metrics from classical ones. This paper is devoted to generalize a construction of non-strong fuzzy metrics from metrics already provided in the literature, both for continuous Archimedean t -norms and for the minimum t -norm. Moreover, we explore the conditions required to adapt this generalized method to obtain fuzzy metrics in the sense of George and Veeramani. Furthermore, we study the relationship between the topology induced by the fuzzy metric obtained through these methods and the topology induced by the classical metric from which it is constructed. Several examples are provided to support and illustrate our study.

Index Terms—Fuzzy metric, Archimedean t -norm, non-strong fuzzy metric, additive generator, minimum t -norm.

I. EXTENDED ABSTRACT

Kramosil and Michalek introduced a notion of fuzzy metric based on probabilistic metrics in [8], where proximity between points is expressed as a value in $[0, 1]$ depending on a parameter t . Following the arguments exposed in [9], a fuzzy metric space can be defined as an ordered triple $(\mathcal{X}, \mathcal{M}, *)$ such that $\mathcal{X}$ is a (non-empty) set, $*$ is a continuous t -norm and $\mathcal{M} : \mathcal{X} \times \mathcal{X} \times \mathbb{R}^+ \rightarrow [0, 1]$ is a mapping satisfying the following conditions, for all $x, y, z \in \mathcal{X}$ and $t, s \in \mathbb{R}^+$.

- (KM1) $\mathcal{M}(x, y, t) = 1$ for all $t \in \mathbb{R}^+$ if and only if $x = y$;
- (KM2) $\mathcal{M}(x, y, t) = \mathcal{M}(y, x, t)$;
- (KM3) $\mathcal{M}(x, y, t) * \mathcal{M}(y, z, s) \leq \mathcal{M}(x, z, t + s)$;
- (KM4) The function $\mathcal{M}_{xy} : \mathbb{R}^+ \rightarrow [0, 1]$ is left-continuous, where $\mathcal{M}_{xy}(t) = \mathcal{M}(x, y, t)$ for each $t \in \mathbb{R}^+$.

Later, George and Veeramani refined this notion in [2] with the aim of ensuring that each fuzzy metric induces a Hausdorff topology, and showed that a fuzzy metric can be obtained from a classical one, which induces the same

topology. Subsequently, several works were dedicated to the topological study of these fuzzy metrics. A significant result in this direction was proved in [5], which establishes that the topology induced by a fuzzy metric is metrizable.

While classical and fuzzy metrics are topologically equivalent, they differ significantly in purely metric contexts such as fixed point theory and completion. Moreover, fuzzy metrics have proven useful in various engineering applications, including image processing, color perception, robust model estimation, and clustering (see, for instance, [1], [3], [4], [10]–[13]). This has led to growing interest in constructing new fuzzy metrics to broaden the available tools for handling problems involving uncertainty.

The aim of the paper is to generalize a construction of (pseudo-)fuzzy metrics from classical (pseudo-)metrics provided in [9], which is recalled below.

Theorem 1.1: Let (X, d) be a pseudo-metric space and let $*$ be a continuous t -norm with additive generator f_* . Then, $(M, *)$ is a fuzzy pseudo-metric on X , where M is the fuzzy set defined on $X \times X \times]0, \infty[$ as follows:

$$M(x, y, t) = f_*^{(-1)}(\max\{d(x, y) - t, 0\}),$$

for all $x, y \in X$ and for all $t \in]0, \infty[$. Furthermore, $(M, *)$ is a fuzzy metric on X if and only if d is a metric on X .

To achieve the aforementioned generalization, a positive real-valued function is introduced in the construction to manage the parameter t , following the ideas of a method for constructing fuzzy metrics from classical metrics, as proposed in [6] and recalled below.

Theorem 1.2: Let (X, d) be a pseudo-metric space, let $\varphi :]0, \infty[\rightarrow]0, \infty[$ be an increasing and left-continuous function, and let $*$ be a continuous Archimedean t -norm and let f_* be an additive generator of $*$. Then $(X, M, *)$ is a strong fuzzy

pseudo-metric space, where the fuzzy set $M : X \times X \times]0, \infty[$ is given, for each $x, y \in X$ and $t \in]0, \infty[$, by

$$M(x, y, t) = f_*^{(-1)} \left(\frac{d(x, y)}{\varphi(t)} \right). \quad (1)$$

Moreover, M is a strong fuzzy metric if and only if d is a metric.

To adapt the previous result to the construction presented in Theorem 1.1, it is necessary to impose an additional condition on the function φ , namely, superadditivity.

Moreover, we also focus on adapting the second method presented in [6] to the construction given in Theorem 1.1, in which the minimum t -norm is involved. It is recalled below.

Theorem 1.3: Let (X, d) be a pseudo-metric space, let $\varphi :]0, \infty[\rightarrow]0, \infty[$ be an increasing, left-continuous and superadditive function, and let $g : [0, \infty] \rightarrow [0, 1]$ be a decreasing and left-continuous function, such that $g(0) = 1$. Then $(X, M, *_M)$ is a fuzzy pseudo-metric space, where the fuzzy set $M : X \times X \times]0, \infty[$ is given, for each $x, y \in X$ and $t \in]0, \infty[$, by

$$M(x, y, t) = g \left(\frac{d(x, y)}{\varphi(t)} \right). \quad (2)$$

Moreover if, in addition, $g^{-1}(1) = \{0\}$ then M is a fuzzy metric if and only if d is a metric.

Finally, we study whether the topology is preserved under both new methods for constructing fuzzy (pseudo-)metrics from classical (pseudo-)metrics. That is, whether the topology induced by the constructed M coincides with the topology induced by the (pseudo-)metric from which it is derived.

The aforementioned results are included in the paper [7].

ACKNOWLEDGMENT

Financiado con Ayuda a Primeros Proyectos de Investigación (PAID-06-24), Vicerrectorado de Investigación de la Universitat Politècnica de València (UPV). This research is part of projects PID2022-139248NB-I00 and PID2022-140189OB-C21 funded by MICIU/AEI/10.13039/501100011033 and ERDF/EU, and CIAICO/2022/051 funded by Generalitat Valenciana.

REFERENCES

- [1] J. G. Camarena, V. Gregori, S. Morillas, and A. Sapena, "Fast detection and removal of impulsive noise using peer groups and fuzzy metrics," *J. Vis. Commun. Image Represent.*, vol. 19, pp. 20–29, 2008.
- [2] A. George and P. Veeramani, "On some results in fuzzy metric spaces," *Fuzzy Sets Syst.*, vol. 64, no. 3, pp. 395–399, 1994.
- [3] S. Grečova and S. Morillas, "Perceptual similarity between color images using fuzzy metrics," *J. Vis. Commun. Image Represent.*, vol. 34, pp. 230–235, 2016.
- [4] V. Gregori, J. J. Miñana, and S. Morillas, "Some questions in fuzzy metric spaces," *Fuzzy Sets Syst.*, vol. 204, pp. 71–85, 2012.
- [5] V. Gregori and S. Romaguera, "Some properties of fuzzy metric spaces," *Fuzzy Sets Syst.*, vol. 115, pp. 485–489, 2000.
- [6] O. Grigorenko, J. J. Miñana, and O. Valero, "Two new methods to construct fuzzy metrics from metrics," *Fuzzy Sets Syst.*, vol. 467, Art. no. 108483, pp. 1–18, 2023.
- [7] O. Grigorenko, J. J. Miñana, and S. Talia, "Generalizing a construction of non-strong fuzzy metrics from metrics and studying their induced topology," *Mathematics*, submitted.
- [8] I. Kramosil and J. Michalek, "Fuzzy metric and statistical metric spaces," *Kybernetika*, vol. 11, pp. 336–344, 1975.
- [9] J. J. Miñana and O. Valero, "A duality relationship between fuzzy metrics and metrics," *Int. J. Gen. Syst.*, vol. 47, no. 6, pp. 593–612, 2018.
- [10] S. Morillas, V. Gregori, and G. Peris-Fajarnés, "New adaptive vector filter using fuzzy metrics," *J. Electron. Imaging*, vol. 16, no. 3, Art. no. 033007, pp. 1–15, 2007.
- [11] A. Ortíz, E. Ortiz, J. J. Miñana, and O. Valero, "On the use of fuzzy metrics for robust model estimation: a RANSAC-based approach," in **Advances in Computational Intelligence (IWANN 2021)**, I. Rojas, G. Joya, and A. Català, Eds., *Lect. Notes Comput. Sci.*, vol. 12861, Cham: Springer, 2021, pp. 165–177.
- [12] N. Ralević, M. Delić, and N. Nedović, "Aggregation of fuzzy metrics and its application in image segmentation," *Iran. J. Fuzzy Syst.*, vol. 19, no. 3, pp. 19–37, 2022.
- [13] A. Stamenković, N. Milosavljević, and N. Ralević, "Application of fuzzy metrics in clustering problems of agricultural crop varieties," *Ekon. Poljopriv.*, vol. 71, no. 1, pp. 121–134, 2024.

Índice de autores

- Alonso, Pedro, 69
Alonso-Moral, Jose M., 21
Andrés, Miguel, 7
Andreu-Perez, Javier, 4
Angulo Sánchez-Herrera, Eusebio, 30
Aragón, Roberto G., 79
Arellano, Joaquin, 40
- Bedregal, B., 19
Bernabe-Moreno, J., 15
Bonilla, Javier, 57
Borrego-Díaz, Joaquín, 81
Bouchet, Agustina, 67
Bugarín-Diz, Alberto, 21
Bustince, Humberto, 19, 23, 40, 59, 61
- Cabrerizo, Francisco Javier, 13, 15
Carrasco, Ramón Alberto, 13
Catala, Alejandro, 21
Chacón-Gómez, Fernando, 71
Cobo, M.J., 15
Cordón, Oscar, 27
Cornejo, M. Eugenia, 71, 77
Cruz, A., 19
Cuevas Martínez, Juan Carlos, 49
- De Miguel, L., 19, 63
Díaz, Irene, 17, 85
Díaz-Jiménez, David, 32, 36, 47, 49
Díaz-Montarroso, Carolina, 83
Díaz-García, Jose A., 45
Díaz-Vázquez, Susana, 67
Dimuro, Graçaliz, 61
Dutta, Bapi, 1
- El Amraoui-Leghzali, Hossam, 41
Espinilla, Macarena, 32, 36, 47, 49
- Esteva, Francesc, 55
- Febbraio, Ferdinando, 17
Fernández, Francisco Javier, 23, 40, 61
Fernández-Basso, Carlos, 32, 41
Fernández-Mora, Arturo, 25
Ferrero-Jaurrieta, Mikel, 19, 23
Figueira, José Rui, 1
Flores-Vidal, Pablo, 34
- Gaitán-Guerrero, Juan F., 36, 47
García-Zamora, Diego, 1
Gispert, Joan, 55
Godo, Lluís, 55
Gómez, Daniel, 34, 59
González, Xabier, 59
González-Hidalgo, M., 63
González-Quesada, Juan Carlos, 13
Grigorenko, Olga, 87
Guerra, Carlos, 40
Guerrero-Sosa, Jared D.T., 43
Gutiérrez, Inmaculada, 59
- Herrera-Viedma, Enrique, 13, 15
Herrerros-Bodalo, Juan-Luis, 49
Huidobro, Pedro, 69
- Ibáñez, Oscar, 27
- Janiš, Vladimir, 69
Ju, Yanbing, 11
Junquera Álvarez, Enol, 17
- Kahana, Tzipi, 27
- Lafuente, Julio, 61
Lara-Martín, Fco. Félix, 81

- Liu, Yaya, 9
Lobo Palacios, David, 73
López, José L., 32, 36, 47
López, Victoria, 7
López-Gómez, Julio Alberto, 3, 30
Lopez-Joya, Salvador, 45
Lopez-Molina, C., 19
López Molina, Pablo, 73
López-Rodríguez, Domingo, 75
- Madrid, Nicolás, 83
Magdalena, Luis, 34
Marco-Detchart, Cedric, 40
Marín, Nicolás, 65
Márquez Hernández, Francisco Alfredo, 38
Martín Moreno, Javier, 38
Martin-Bautista, Maria J., 41, 45
Martínez, Luis, 1, 4, 9, 25
Martínez-Alonso, Raúl, 41
Martínez-Cruz, Carmen, 36, 47
Massanet, S., 63
Mata, Francisco, 13
Medina, Jesús, 71, 77, 79
Merino, Luis, 53
Miñana, Juan José, 87
Mir, A., 63
Molina-Ruiz, Samuel, 79
Montero, Javier, 59
Montes, Susana, 67, 69
Montoro-Montarroso, Andrés, 43
Morales-Garzón, Andrea, 41
Moreno-Garcia, Juan, 3, 25
Muñoz-Exposito, Jose-Enrique, 49
Muñoz-Valero, David, 3, 25
- Navarro, Gabriel, 53
- Ocaña, Francisco José, 77
Ojeda-Hernández, Manuel, 75
Olivas, Jose A., 43
- Paiva, R., 19
Peregrín Rubio, Antonio, 38
Pérez, Ignacio Javier, 13, 15
- Perez-Ferreiro, Pablo Miguel, 21
Plaza, Guadalupe, 43
Porcel, Carlos, 11
- Quesada-Real, Francisco J., 4
- R-García, Guillermo, 27
Ramírez-Poussa, Eloísa, 83
Ramos-Cruz, Bruno, 4
Ramos-González, Víctor, 81
Remeseiro, Beatriz, 17
Rico, Noelia, 85
Riera, J. V., 63
Riofrío-Luzcando, Diego, 7
Rivas-Gervilla, Gustavo, 65
Rodríguez, J. Tinguaro, 59
Rodríguez, Rosa M., 9
Rodriguez-Martinez, Iosu, 23
Romero, Francisco P., 30, 43
Ruiz, M. Dolores, 45
- Salvatierra, Carlos, 53
Sánchez, Daniel, 65
Santos, Evangelina, 53
Santos, Matilde, 7
Sanz, José Antonio, 51
Serrano-Guerrero, Jesus, 43
Sesma-Sara, Mikel, 51
Špírková, Jana, 23
Suarez-Fernandez, Sara, 67
- Takáč, Z., 19
Talia, Simona, 87
Triviño-Cabrera, Alicia, 49
- Villar, Mario, 85
Villarrubia-Martin, Enrique Adrian, 3
Villegas-Yeguas, Antonio D., 27
- Wang, Han, 11
- Yao, Yiyu, 9
Yuste-Delgado, Antonio-Jesús, 49
- Zhu, Lina, 9