

XIX Conferencia de la Asociación Española para la Inteligencia Artificial CAEPIA 20/21

22-24 de septiembre de 2021

Málaga, España

ENRIQUE ALBA, FRANCISCO CHICANO, GABRIEL LUQUE, RODRIGO
GIL-MERINO, CARLOS COTTA, DAVID CAMACHO, MANUEL OJEDA-ACIEGO,
SUSANA MONTES, ALICIA TRONCOSO, JOSÉ RIQUELME, EVA ONAINDIA,
MARÍA JOSÉ DEL JESÚS, JOSÉ ANTONIO GÁMEZ, ALBERTO BUGARÍN, MAR
MARCOS, AGAPITO LEDEZMA, JUAN PEDRO LLERENA, JAVIER ECHANOBE,
JAMAL TOUTOUH, SANTIAGO MUIÑOS



Editores

Enrique Alba
Universidad de Málaga
Málaga, España

Gabriel Luque
Universidad de Málaga
Málaga, España

Carlos Cotta
Universidad de Málaga
Málaga, España

Manuel Ojeda-Aciego
Universidad de Málaga
Málaga, España

Alicia Troncoso
Universidad Pablo de Olavide
Sevilla, España

Eva Onaindia
Universidad Politécnica de Valencia
Valencia, España

José Antonio Gámez
Universidad de Castilla La Mancha
Albacete, España

Mar Marcos
Universitat Jaume I
Castellón de la Plana, España

Juan Pedro Llerena
Universidad Carlos III de Madrid
Madrid, España

Jamal Toutouh
Universidad de Málaga
Málaga, España

Francisco Chicano
Universidad de Málaga
Málaga, España

Rodrigo Gil-Merino
Universidad de Málaga
Málaga, España

David Camacho
Universidad Politécnica de Madrid
Madrid, España

Susana Montes
Universidad de Oviedo
Oviedo, España

José Riquelme
Universidad de Sevilla
Sevilla, España

María José del Jesús
Universidad de Jaén
Jaén, España

Alberto Bugarín
Universidad de Santiago de Compostela
Santiago de Compostela, España

Agapito Ledezma
Universidad Carlos III de Madrid
Madrid, España

Javier Echanobe
Universidad del País Vasco
Bilbao, España

Santiago Muiños
Centro Tecnológico AIMEN
Pontevedra, España

ISBN: 978-84-09-30514-8

© Los autores, 2021

Presentación de CAEPIA 20/21

Este volumen contiene un conjunto de artículos seleccionados y revisados por pares enviados a CAEPIA 20/21, la XIX Conferencia de la Asociación Española de Inteligencia Artificial, celebrada en Málaga, España, del 22 al 24 de septiembre de 2021. CAEPIA es un evento bienal español bien establecido sobre Inteligencia Artificial (IA) que comenzó en 1985. Ediciones anteriores tuvieron lugar en Alicante, Málaga, Murcia, Gijón, San Sebastián, Santiago de Compostela, Sevilla, La Laguna, Madrid, Albacete, Salamanca y Granada.

CAEPIA es un foro nacional abierto a investigadores de todo el mundo para presentar y discutir sus últimos avances científicos y tecnológicos en IA. Los autores podían optar por cinco tipos de contribuciones: trabajos inéditos de investigación para un volumen en la serie *Lecture Notes in Artificial Intelligence* de Springer, trabajos inéditos de investigación para estas actas, trabajos destacados ya publicados, proyectos de doctorado, desarrollos de aplicaciones móviles y vídeos divulgativos. La conferencia acogió tanto investigación teórica como metodológica, técnica y aplicada.

Dentro de CAEPIA se organizaron varios talleres y congresos federados relacionados con los temas más relevantes de la IA: XX Congreso Español Sobre Tecnologías y Lógica Fuzzy (ESTYLF); XIV Congreso Español de Metaheurísticas, Algoritmos Evolutivos y Bioinspirados (MAEB); X Simposio de Teoría y Aplicaciones de la Minería de Datos (TAMIDA); y seis talleres. También contamos con un Doctoral Consortium (DC). Este es un foro para que los estudiantes de doctorado interactúen con otros investigadores discutiendo sus planes de trabajo y avances en el doctorado. Como actividad adicional de IA, llevamos a cabo el 4º Concurso de Aplicaciones Móviles con Técnicas de IA, junto con una nueva edición del Concurso de Vídeos de Divulgación de IA.

Todas las actividades anteriores avalan la IA, y nos esforzamos por alcanzar una alta calidad en los artículos científicos, el DC y las competiciones. El programa científico de CAEPIA 20/21 también ofreció una vía para difundir trabajos destacados (Key Works: KW) publicados recientemente en revistas y foros de alto impacto científico. CAEPIA siempre ha tenido como objetivo ser reconocida como una conferencia insignia en IA y, por lo tanto, los artículos fueron revisados por pares. El número total de envíos a CAEPIA 20/21 fue de 186 (en este número no se incluyeron ni DC ni concursos ni presentaciones KW, que suman 83 contribuciones adicionales, y que pasaron por un proceso de evaluación diferente). Los revisores evaluaron la calidad general de los manuscritos presentados, junto con la calidad de la metodología empleada, la solidez de las conclusiones, la importancia del tema, la claridad de la redacción y su organización, entre otros criterios de evaluación. A partir de estas revisiones, los responsables de área, presidentes de congresos y organizadores de talleres y sesiones especiales propusieron un número final de artículos que fueron analizados y aprobados por los editores de este volumen.

CAEPIA 20/21 invitó a dos investigadores de renombre internacional a impartir una charla plenaria. Nuestros dos ponentes plenarios fueron Óscar Cordón (Inteligencia Artificial para Antropología Forense e Identificación Humana) y Yaochu Jin (Optimización Evolutiva Basada en Datos). Nuestra conferencia se celebró como un gran evento dentro de uno aún mayor: la Conferencia Española de Informática (CEDI), que también contó con charlas plenarias muy interesantes.

CAEPIA y los organizadores de CAEPIA 20/21 reconocieron las mejores tesis doctorales y artículos originales en eventos federados escritos tanto por investigadores consolidados como por estudiantes. CAEPIA 20/21 también tuvo como objetivo promover la presencia de mujeres en la investigación de IA. Como en ediciones anteriores, el premio Frances Allen reconoció las dos mejores tesis doctorales defendidas por una mujer durante los dos últimos años.

Los editores de este volumen quieren agradecer a las numerosas personas que contribuyeron al éxito de CAEPIA 20/21: autores, miembros de los comités científicos y los comités de programa, ponentes invitados, organizadores de eventos, gestores de medios electrónicos, etc. También, agradecer el trabajo incansable del comité organizador, nuestros patrocinadores (como VRAIN en Valencia), el equipo de Springer y CAEPIA por su apoyo.

Por último, pero no menos importante, en nombre de los participantes de CAEPIA 20/21, Enrique Alba (presidente) y Francisco Chicano (responsable de este volumen) dan las gracias a la organización de CEDI, la Universidad de Málaga (sede local de la conferencia) y a toda la comunidad española que trabaja en IA (y sus numerosos colaboradores extranjeros) por hacer de este evento un verdadero éxito.

Enrique Alba

Presidente de CAEPIA 20/21

Presentación de la Presidenta de la Asociación Española para la Inteligencia Artificial

Es para mí un inmenso honor presentaros esta nueva edición en Málaga de la serie de conferencias CAEPIA, en la que además se cierra una etapa en la que he tenido el inmenso privilegio de trabajar como Presidenta de la Asociación Española para la Inteligencia Artificial (AEPIA).

Debemos el inicio de estas conferencias bienales en los años 80 a un pequeño, pero muy entusiasta, grupo de investigadores españoles en Inteligencia Artificial (IA). A pesar de que en aquella época la IA estaba en un momento álgido, la situación no era comparable al momento actual, en el que nuestra disciplina se ha convertido en una de las más influyentes en la nueva revolución tecnológica que está teniendo lugar. Además, hay ciertas diferencias con entonces que auguran que quizás, y aunque quede mucho camino aún por recorrer, esta vez la primavera puede ser casi eterna. Hay algunos factores que han confluído para dar este estallido de vida a la IA, como son la disponibilidad de enormes cantidades de datos (uno de los alimentos de los algoritmos actuales de IA), el abaratamiento de la computación en nube o la existencia de plataformas distribuidas y paralelas, que permiten el procesado rápido y económicamente viable de esas grandes cantidades de datos. Confluye también un imparable cambio social, ya que todos interactuamos constantemente con nuestros múltiples dispositivos móviles, en un mundo en el que la interconectividad es imprescindible, y aún lo es más desde la irrupción de la pandemia CoVid-19 en nuestras vidas.

Nuestra disciplina es la cabeza del cambio, y ello implica no sólo cambios tecnológicos y sociales, sino también de poder económico y geopolítico. Existe una cierta dinámica de confrontación entre EEUU y China, en la que esta última ha explicitado que su meta es convertirse en la nación que domine el mundo en IA en el año 2030, mediante una inversión masiva en I+D+I y en políticas de fusión “militar-civil”, aunque por el momento son las empresas americanas las que van en cabeza. Y en este contexto, la UE se ha decidido por una tercera vía, la de intentar fomentar y mantener su talento investigador en IA, además de atraer talento extranjero, y promover asimismo una IA robusta, transparente, auditable, ética y al servicio de la humanidad, con las personas en el centro. Los diferentes países de la UE, entre ellos España, han elaborado Estrategias de Inteligencia Artificial. Según estudios de varias consultoras, y en palabras del Presidente del Gobierno en la presentación de la Estrategia Nacional de IA (ENIA) recientemente este 2 de Diciembre, la IA ha contribuido en 2018 con unos 1700 millones de euros al PIB mundial, y en 2030 se estima que aportará ocho veces más, alrededor de 14 mil millones de euros, una cifra cercana a todo el PIB actual de la UE. Es indudable que la disciplina tiene una enorme capacidad de transformación en la economía.

Nuestra ENIA es más modesta que las estrategias de otros países europeos en inversión, así que es mandatorio para hacer valer nuestras fortalezas como país. Y es en este esfuerzo colectivo donde contribuimos los científicos que trabajamos en los diferentes campos de la IA. En esta conferencia estarán representados muchos de los grupos españoles más relevantes, aportando nuestra investigación y nuestra capacidad de transferencia, parte de la cual se plasma en los trabajos que se presentan tanto en las conferencias plenarias como en las diferentes sesiones, Key notes, Workshops, etc. incluidos en esta XIX CAEPIA, que es co-ocurrente con las conferencias del XX Congreso Español de Tecnologías y Lógica Difusa (ESTYLF), el XIV Congreso Español de Metaheurística y Algoritmos evolutivos y bioinspirados (MAEB), y el X Symposium de Teoría y Aplicaciones de Minería de Datos (TAMIDA), todos ellos co-celebrándose además junto con el Congreso Español De Informática (CEDI). La conferencia también se implica con el talento emergente, incluyendo programas para los estudiantes y los jóvenes investigadores, como es el caso del Doctoral Consortium, los Premios a los Jóvenes investigadores autores de los mejores artículos, o a las mejores Apps y videos divulgativos de IA. En AEPIA también contribuimos con acciones de género positivas mediante los premios Frances Allen a las mejores tesis doctorales realizadas por mujeres, en un intento de reducir la significativa brecha de género e incorporar más talento a la disciplina, en esta ocasión agradecemos su patrocinio en este Premio al Instituto Valenciano de Investigación en Inteligencia Artificial (VRAIN).

Tenemos un tejido investigador del que sentirnos orgullosos, debemos trabajar para mantener, retener y atraer talento, una de las materias primas fundamentales para abrazar el futuro. En esta XIX CAEPIA aportamos una muestra de nuestra contribución a esa apuesta de futuro.

Amparo Alonso Betanzos
Presidenta de AEPIA

Ponentes Plenarios

Óscar Cerdón

Instituto Andaluz Interuniversitario en Ciencia de Datos e Inteligencia Computacional (Instituto DaSCI)

Oscar Cerdón fue el Director Fundador del Centro de Enseñanzas Virtuales (2001-05) y Delegado de la Rectora para la Universidad Digital (2015-19) de la Universidad de Granada (UGR). Fue uno de los investigadores fundadores del European Centre for Soft Computing (2006-11), siendo luego contratado como Investigador Afiliado Distinguido hasta diciembre de 2015. En la actualidad es Catedrático de Universidad en la UGR. Durante más de 25 años, ha impulsado programas de investigación y transferencia en fundamentos y aplicaciones de inteligencia computacional con un gran reconocimiento internacional. Ha publicado más de 380 contribuciones científicas, incluyendo un libro de investigación sobre Genetic Fuzzy Systems (con más de 1400 citas en Google Scholar) y 112 artículos de revista JCR-SCI (68 en Q1 y 38 en D1), ha dirigido 19 tesis doctorales y coordinado 37 proyectos y contratos de investigación (con un presupuesto global de más de 9M€). A fecha de mayo de 2021, sus publicaciones han recibido 5422 citas (H-index=39), estando incluido en el 1 % de investigadores más citados en el mundo (fuente: Web of Science); con 14687 citas y H-index=58 en Google Scholar. También tiene una patente internacional en explotación sobre un sistema inteligente para identificación forense, comercializada en México y Sudáfrica.



Ha recibido el Premio de Jóvenes Investigadores de la UGR (2004), el Premio IEEE CIS Outstanding Early Career Award (en su primera edición, 2011), el Premio Nacional de Informática ARITMEL (2014) por la Sociedad Científica de España, el IEEE Fellow (2018) y el IFSA Fellow (2019). Fue miembro del Grupo de Expertos que desarrolló la Estrategia Española de I+D+I en Inteligencia Artificial para el Ministerio de Ciencia, Innovación y Universidades (2018-19). Es o ha sido Editor Asociado de 19 revistas internacionales, siendo reconocido como Outstanding Associate Editor de IEEE Transactions on Fuzzy Systems (2008) y de IEEE Transactions on Evolutionary Computation (2019). Desde 2004, ha ocupado distintos puestos de representación en EUSFLAT e IEEE Computational Intelligence Society.

En la actualidad investiga en inteligencia artificial para identificación forense (en colaboración con el laboratorio de Antropología Física de la UGR y varios laboratorios forenses y cuerpos y fuerzas de seguridad internacionales) y en modelado basado en agentes y análisis de redes sociales para marketing (en colaboración con R0D Brand Consultants en proyectos para CAPSA, Mercedes, Jaguar-Land Rover, El Corte Inglés, Telefónica, Samsung, Coca Cola Europa, Cola Cao, WiZink,...).

Inteligencia Artificial para Antropología Forense e Identificación Humana

Los métodos de identificación forense basados en el esqueleto empleados por antropólogos, odontólogos y patólogos representan el primer paso en cualquier proceso de identificación humana (ID) y la última oportunidad de identificación de la víctima cuando no puede aplicarse el análisis de ADN o de huellas dactilares. Incluyen métodos como la estimación del perfil biológico (BP), la radiografía comparativa (CR), la superposición craneofacial (CFS) y la comparación de registros dentales. La BP implica el estudio de restos óseos para encontrar rasgos característicos (edad, sexo, estatura y ascendencia) que ayuden a determinar la identidad del individuo. Desempeña un papel crucial en la reducción del rango de coincidencias potenciales durante el proceso de ID, antes de la confirmación mediante una o más técnicas de ID. La CR considera la comparación ante-mortem (AM) y post-mortem (PM) de diferentes huesos y cavidades (senos frontales del cráneo, clavículas, rótulas,...) que han demostrado ser útiles para la identificación positiva por su individualidad y singularidad. La CFS tiene como objetivo superponer un cráneo con algunas imágenes AM de un candidato para determinar si corresponden a la misma persona.

Sin embargo, los profesionales todavía emplean un paradigma de observación utilizando métodos subjetivos introducidos hace muchas décadas basados en la descripción oral y documentación escrita de los hallazgos obtenidos y la comparación manual y visual de los datos AM y PM. El diseño de métodos sistemáticos, automáticos y confiables para apoyar al antropólogo forense en la aplicación de BP, CFS y CR, evitando el uso de procedimientos manuales subjetivos, propensos a errores y que consumen mucho tiempo, es una necesidad para mejorar la identificación forense. El uso de inteligencia artificial, en particular inteligencia computacional (algoritmos evolutivos, conjuntos difusos y aprendizaje profundo), visión por ordenador (registrado de imágenes 3D-2D y procesamiento de imágenes) y aprendizaje automático explicable es una forma muy adecuada para lograr este objetivo. En esta charla presentaremos tres sistemas inteligentes para CFS, CR y estimación de la edad de la muerte a partir del esqueleto, desarrollados en colaboración con el Laboratorio de Antropología Física de la Universidad de Granada en el marco de un proyecto de investigación desarrollado en los últimos quince años. Uno de estos sistemas está protegido por una patente internacional explotada por Panacea Cooperative Research y está comercializado en diferentes países.

Yaochu Jin

Universidad de Surrey, Reino Unido

Yaochu Jin se licenció, estudió y se doctoró en la Universidad de Zhejiang (Hangzhou, China) en 1988, 1991 y 1996, respectivamente, y obtuvo el título de ingeniero de la Universidad del Ruhr (Bochum, Alemania) en 2001.

En la actualidad es profesor distinguido de Inteligencia Computacional en el Departamento de Informática de la Universidad de Surrey, Guildford, Reino Unido, donde dirige el Grupo de Computación e Ingeniería Inspirada en la Naturaleza. Ha sido “Finland Distinguished Professor” de la Universidad de Jyväskylä (Finlandia), “Changjiang Distinguished Visiting Professor” de la Universidad de Northeastern (China) y “Distinguished Visiting Scholar” de la Universidad Tecnológica de Sydney (Australia). Recientemente, el Ministerio Federal de Educación e Investigación de Alemania le ha concedido la “Cátedra Alexander von Humboldt de Inteligencia Artificial”. Sus principales intereses de investigación son la optimización evolutiva asistida por datos, el aprendizaje automático fiable, el aprendizaje evolutivo multiobjetivo, la robótica de enjambre y los sistemas evolutivos de desarrollo.

El Dr. Jin es actualmente editor jefe de IEEE Transactions on Cognitive and Developmental Systems y de Complex & Intelligent Systems. Fue conferenciante distinguido del IEEE y vicepresidente de la Sociedad de Inteligencia Computacional del IEEE. Ha recibido el premio al mejor artículo de 2018 y 2020 de IEEE Transactions on Evolutionary Computation, el premio al mejor artículo de 2014, 2016 y 2019 de IEEE Computational Intelligence Magazine y el premio al mejor artículo del Simposio de IEEE sobre Inteligencia Computacional en Bioinformática y Biología Computacional de 2010. Está reconocido como Investigador Altamente Citado en 2019 y 2020 por el Grupo Web of Science. Es miembro de IEEE.



Optimización Evolutiva Basada en Datos

Muchos problemas de optimización del mundo real no tienen funciones objetivo analíticas y las evaluaciones de los objetivos deben basarse en costosos cálculos o experimentos físicos. Estos problemas de optimización se conocen como problemas de optimización basados en datos. Esta charla ofrece una visión general de la optimización evolutiva asistida por datos de sistemas complejos. Comenzamos con una breve introducción a las ideas básicas de la optimización evolutiva basada en datos, seguida de estrategias avanzadas de gestión de sustitutos que hacen uso de técnicas avanzadas de aprendizaje automático como el aprendizaje semisupervisado, el aprendizaje de transferencia y el aprendizaje de conjunto. Se expondrán ejemplos del mundo real, desde la optimización del diseño de ingeniería hasta la búsqueda de arquitecturas neuronales.

Organización

Presidente

Enrique Alba	Universidad de Málaga, España
--------------	-------------------------------

Presidenta del Comité de Premios

Eva Onaindia	Universidad Politécnica de Valencia, España
--------------	---

Presidentes de Talleres y Trabajos Destacados

Rodrigo Gil-Merino	Universidad de Málaga, España
Araceli Sanchís	Universidad Carlos III de Madrid, España

Presidente de la Sesión General de CAEPIA

Francisco Chicano	Universidad de Málaga, España
-------------------	-------------------------------

Presidentes de Competiciones

Alberto Bugarín	Universidad de Santiago de Compostela, España
José Antonio Gámez	Universidad de Castilla La Mancha, España

Responsable de Actas LNAI

Gabriel Luque	Universidad de Málaga, España
---------------	-------------------------------

Presidentes del Doctoral Consortium

José Riquelme	Universidad de Sevilla, España
María José Del Jesús Díaz	Universidad de Jaén, España

Responsables de Área de CAEPIA

Óscar Corcho	Universidad Politécnica de Madrid, España
Javier del Ser	Tecnalia, España
Juan Manuel Fernández Luna	Universidad de Granada, España
Jesús García Herrero	Universidad Carlos III de Madrid, España
Pedro Larrañaga	Universidad Politécnica de Madrid, España
Patricio Martínez Barco	Universidad de Alicante, España
Miguel Molina Solana	Universidad de Granada, España
Serafín Moral Callejón	Universidad de Granada, España
Eva Onaindia	Universidad Politécnica de Valencia, España
Juan Pavón Mestras	Universidad Complutense de Madrid, España
David Pearce	Universidad Politécnica de Madrid, España
José Luis Pérez de la Cruz Molina	Universidad de Málaga, España
Petia Radeva	Universtat de Barcelona & Centro de Visión por Computador, España

Presidentes de Conferencias Federadas

XX Congreso Español sobre Tecnologías y Lógica Fuzzy (ESTYLF)

Manuel Ojeda-Aciego	Universidad de Málaga, España
Susana Montes	Universidad de Oviedo, España

XIV Congreso Español de Metaheurísticas, Algoritmos Evolutivos y Bioinspirados (MAEB)

Carlos Cotta	Universidad de Málaga, España
David Camacho	Universidad Politécnica de Madrid, España

X Simposio de Teoría y Aplicaciones de Minería de Datos (TAMIDA)

Alicia Troncoso
José Riquelme

Universidad Pablo de Olavide, España
Universidad de Sevilla, España

Organizadores de Talleres**II Taller de Grupos de investigación españoles de IA en Biomedicina (IABiomed)**

Mar Marcos
David Riaño
José M. Juárez

Universitat Jaume I, España
Universitat Rovira i Virgili, España
Universidad de Murcia, España

I Taller de La IA aplicada a los Sistemas Inteligentes de Transporte (IA-SIT)

Paz Sesmero
Pablo Marín
Juan Pedro Llerena

Universidad Carlos III de Madrid, España
Universidad Carlos III de Madrid, España
Universidad Carlos III de Madrid, España

I Taller de Inteligencia Artificial para el Desarrollo (AI4D)

Agapito Ledezma Espino
José Antonio Iglesias Martínez
David Camilo Corrales Muñoz
Elias Said Hung

Universidad Carlos III de Madrid, España
Universidad Carlos III de Madrid, España
Institut national de la recherche agronomique, Francia
Universidad Internacional de La Rioja, España

I Taller de Inteligencia Artificial Cuántica (QUARTI)

Javier Echanobe
Enrique Alba

Universidad del País Vasco, España
Universidad de Málaga, España

I Taller sobre Modelos Generativos Profundos Aplicados (AGM)

Carlos González Val
Juan Manuel Fernández Montenegro
Andrea Fenández Martinez
Santiago Muñios Landin
Omar A. Mures
José A. Iglesias-Gutián

Centro Tecnológico AIMEN, España
Centro Tecnológico AIMEN, España
Centro Tecnológico AIMEN, España
Centro Tecnológico AIMEN, España
Universidad de A Coruña (UDC), España
Universidad de A Coruña (UDC), CITIC, España

I Taller Ciudades inteligentes e Inteligencia Artificial (SCAI)

Jamal Toutouh
Enrique Alba

Universidad de Málaga, España
Universidad de Málaga, España

Organizadores de Sesiones Especiales**ESTYLF. Toma de decisiones con información difusa**

Rosa Rodríguez
Álvaro Labella
Luis Martínez

Universidad de Jaén, España
Universidad de Jaén, España
Universidad de Jaén, España

ESTYLF. Funciones de agregación y preagregación y sus aplicaciones

Humberto Bustince
Graçaliz Dimuro
Javier Fernandez

Universidad Pública de Navarra, España
Universidad Pública de Navarra, España
Universidad Pública de Navarra, España

MAEB. Algoritmos Bioinspirados Profundos

Francisco Chávez
Pablo García

Universidad de Extremadura, España
Universidad de Granada, España

MAEB. Aplicaciones en Medicina y Biotecnología

Ignacio Hidalgo
José Manuel Velasco

Universidad Complutense de Madrid, España
Universidad Complutense de Madrid, España

MAEB. Estrategias de implementación eficiente de heurísticas y metaheurística

José Manuel Colmenar

Universidad Rey Juan Carlos, España

Eduardo García Pardo

Universidad Rey Juan Carlos, España

MAEB. Metaheurísticas Aplicadas al Análisis de Redes Sociales

Jesús Sánchez-Oro

Universidad Rey Juan Carlos, España

Antonio González-Pardo

Universidad Rey Juan Carlos, España

Abraham Duarte

Universidad Rey Juan Carlos, España

Comité de Organización**Publicidad**

Rubén Saborido

Universidad de Málaga, España

Diseño Gráfico

Christian Cintrano

Universidad de Málaga, España

José Ángel Morell

Universidad de Málaga, España

Web

José Ángel Morell

Universidad de Málaga, España

Responsables IT

Francisco Chicano

Universidad de Málaga, España

Andrés Camero

Centro Aeroespacial Alemán, Alemania

Equipo de Soporte IT

Zakaria Abdelmoiz DAHI

Universidad de Málaga, España

Rubén Saborido

Universidad de Málaga, España

Christian Cintrano

Universidad de Málaga, España

Comités de Programa

Sesión General de CAEPIA

Aghaei, Maedeh	Universidad de Barcelona, España
Alonso-Betanzos, Amparo	Universidad de La Coruña, España
Bacardit, Jaume	Newcastle University, Reino Unido
Bahamonde, Antonio	Universidad de Oviedo, España
Bajo, Javier	Universidad Politécnica de Madrid, España
Balaguer-Ballester, Emili	Bournemouth University, Reino Unido
Barbancho, Ana M.	Universidad de Málaga, España
Barreiro, Alvaro	Universidad de La Coruña, España
Barrenechea, Edurne	Universidad Pública de Navarra, España
Barro, Senén	Universidad de Santiago de Compostela, España
Bañares, Jose Angel	Universidad de Zaragoza, España
Bernardos, Ana M.	Universidad Politécnica de Madrid, España
Bielza Lozoya, Concha	Universidad Politécnica de Madrid, España
Borrajo, Daniel	Universidad Carlos III de Madrid, España
Botia, Juan	King's College London, Reino Unido
Botti, Vicent	Universitat Politècnica de València, España
Bugarín, Alberto	Universidad de Santiago de Compostela, España
Bustince, Humberto	Universidad Pública de Navarra, España
Cadenas, José M.	Universidad de Murcia., España
Callejas, Zoraida	Universidad de Granada, España
Cantador, Iván	Universidad Autónoma de Madrid, España
Castelo, Robert	Universitat Pompeu Fabra, España
Corchuelo, Rafael	Universidad de Sevilla, España
Cordon, Oscar	Universidad de Granada, España
Damas, Sergio	Universidad de Granada, España
de Campos, Luis M.	Universidad de Granada, España
de Carvalho, Andre	Universidad de São Paulo, Brasil
De La Ossa, Luis	Universidad de Castilla-La Mancha, España
Del Coz, Juan José	Universidad de Oviedo, España
Dorronsoro, Jose	Universidad Autónoma de Madrid, España
Díez, Francisco Javier	Universidad Nacional de Educación a Distancia, España
Fernández López, Mariano	Universidad San Pablo CEU, España
Fernández-Caballero, Antonio	Universidad de Castilla-La Mancha, España
Ferri, Francesc J.	Universidad de Valencia, España
Galán, José Manuel	Universidad de Burgos, España
Gamez, Jose	Universidad de Castilla-La Mancha, España
García Bringas, Pablo	Universidad de Deusto, España
García-Castro, Raúl	Universidad Politécnica de Madrid, España
Gasca, Rafael M.	Universidad de Sevilla, España
Godo, Lluís	IIIA - CSIC, España
Gomez Romero, Juan	Universidad de Granada, España
Gonzalez, Antonio	Universidad de Granada, España
Graña, Manuel	Universidad del País Vasco, España
Guijarro-Berdiñas, Bertha	Universidad de La Coruña, España
Gómez-Rodríguez, Carlos	Universidade da Coruña, España
Hervas, Cesar	Universidad de Córdoba, España
Inza, Inaki	Universidad del País Vasco, España
Jimenez Linares, Luis	Universidad de Castilla-La Mancha, España
Levy, Jordi	IIIA - CSIC, España
Lopez, Adolfo	Universidad de Valladolid, España
Luaces, Oscar	Universidad de Oviedo, España
López, Victoria	CUNEF Universidad, España
Mandow, Lawrence	Universidad de Málaga, España

Manyà, Felip	IIIA-CSIC, España
Martínez, Luis	Universidad de Jaén, España
Martínez Tomas, Rafael	Universidad Nacional de Educación a Distancia, España
Massa, Jose	Universidad de Buenos Aires, Argentina
Medina-Carnicer, Rafael	Universidad de Córdoba, España
Melian, Belen	Universidad de La Laguna, España
Meseguer, Pedro	IIIA - CSIC, España
Molina, Jose M.	Universidad Carlos III de Madrid, España
Morales-Bueno, Rafael	Universidad de Málaga, España
Muguerza, Javier	Universidad del País Vasco, España
Navas, Ismael	Universidad de Málaga, España
Olivas, Jose Angel	Universidad de Castilla-La Mancha, España
Ossowski, Sascha	Universidad Rey Juan Carlos, España
Palma, Jose	Universidad de Murcia, España
Parraga, C. Alejandro	Centro de Visión por Computador, España
Parreño, Francisco	Universidad de Castilla-La Mancha, España
Patricio, Miguel Angel	Universidad Carlos III de Madrid, España
Peregrin, Antonio	Universidad de Huelva, España
Pla, Filiberto	Universidad Jaume I, España
Pomares, Hector	Universidad de Granada, España
Puerta, Jose M	Universidad de Castilla-La Mancha, España
R. Vela, Camino	Universidad de Oviedo, España
Rizo, Ramon	Universidad de Alicante, España
Rosso, Paolo	Universitat Politècnica de València, España
Salmeron, Antonio	Universidad de Almería, España
Sanchez, Jose Salvador	Universitat Jaume I, España
Sanchez, Luciano	Universidad de Oviedo, España
Saquete, Estela	Universidad de Alicante, España
Segarra, Encarna	Universitat Politècnica de València, España
Soria, Emilio	Universidad de Valencia, España
Taboada, Maria	Universidad de Santiago de Compostela, España
Talavera, Estefania	Universidad de Delft, Países Bajos
Ureña-López, L. Alfonso	Universidad de Jaén, España
Valencia-Garcia, Rafael	Universidad de Murcia, España
Ventura, Sebastián	Universidad de Córdoba, España
Verdegay, José Luis	Universidad de Granada, España

Volumen de LNAI

Aghaei Gavari, Maedeh	Universidad de Barcelona, España
Alonso-Betanzos, Amparo	Universidad de La Coruña, España
Bacardit, Jaume	Newcastle University, Reino Unido
Bahamonde, Antonio	Universidad de Oviedo, España
Bajo, Javier	Universidad Politécnica de Madrid, España
Balaguer-Ballester, Emili	Bournemouth University, Reino Unido
Barbancho, Ana M.	Universidad de Málaga, España
Barreiro, Alvaro	Universidad de La Coruña, España
Barrenechea, Edurne	Universidad Pública de Navarra, España
Barro, Senén	Universidad de Santiago de Compostela, España
Bañares, Jose Angel	Universidad de Zaragoza, España
Bernardos, Ana M.	Universidad Politécnica de Madrid, España
Bielza Lozoya, Concha	Universidad Politécnica de Madrid, España
Borrajo, Daniel	Universidad Carlos III de Madrid, España
Botia, Juan	King's College London, Reino Unido
Botti, Vicent	Universitat Politècnica de València, España
Bustince, Humberto	Universidad Pública de Navarra, España
Cabalar, Pedro	Universidad de La Coruña, España
Cadenas, José M.	Universidad de Murcia, España
Callejas, Zoraida	Universidad de Granada, España
Cantador, Iván	Universidad Autónoma de Madrid, España

Castelo, Robert	Universitat Pompeu Fabra, España
Corchuelo, Rafael	Universidad de Sevilla, España
Damas, Sergio	Universidad de Granada, España
de Campos, Luis M.	Universidad de Granada, España
de Carvalho, Andre	Universidad de São Paulo, Brazil
De La Ossa, Luis	Universidad de Castilla-La Mancha, España
Del Coz, Juan José	Universidad de Oviedo, España
Díaz-Rodríguez, Natalia	ENSTA Paris Institut Polytechnique Paris, France
Dorronsoro, Jose	Universidad Autónoma de Madrid, España
Fandinno, Jorge	Potsdam University, Alemania
Fernández López, Mariano	Universidad San Pablo CEU, España
Fernández-Caballero, Antonio	Universidad de Castilla-La Mancha, España
Ferri, Francesc J.	Universidad de Valencia, España
Galán, José Manuel	Universidad de Burgos, España
Gamez, Jose	Universidad de Castilla-La Mancha, España
García-Castro, Raúl	Universidad Politécnica de Madrid, España
Gasca, Rafael M.	Universidad de Sevilla, España
Godo, Lluís	IIIA - CSIC, España
Gomez Romero, Juan	Universidad de Granada, España
Gonzalez, Antonio	Universidad de Granada, España
Graña, Manuel	Universidad del País Vasco, España
Gómez-Rodríguez, Carlos	Universidade da Coruña, España
Hervas, Cesar	Universidad de Córdoba, España
Inza, Inaki	Universidad del País Vasco, España
Jimenez Linares, Luis	Universidad de Castilla-La Mancha, España
Kushibar, Kaisar	Universidad de Barcelona, España
Levy, Jordi	IIIA - CSIC, España
Lopez, Adolfo	Universidad de Valladolid, España
López, Victoria	CUNEF Universidad, España
Madow, Lawrence	Universidad de Málaga, España
Manyà, Felip	IIIA-CSIC, España
Martínez-Camara, Eugenio	Universidad de Granada, España
Martínez, Luis	Universidad de Jaén, España
Massa, Jose	Universidad de Buenos Aires, Argentina
Medina-Carnicer, Rafael	Universidad de Córdoba, España
Mejia, Marcos	Universitat de Barcelona, España
Melian, Belen	Universidad de La Laguna, España
Meseguer, Pedro	IIIA - CSIC, España
Millan, Eva	Universidad de Málaga, España
Molina, Jose M.	Universidad Carlos III de Madrid, España
Morales-Bueno, Rafael	Universidad de Málaga, España
Nagarajan, Bhalaji	Universitat de Barcelona, España
Núñez, Jose Fernando	Universitat de Barcelona, España
Olivas, Jose Angel	Universidad de Castilla-La Mancha, España
Ossowski, Sascha	Universidad Rey Juan Carlos, España
Parreño, Francisco	Universidad de Castilla-La Mancha, España
Patricio, Miguel Angel	Universidad Carlos III de Madrid, España
Peregrin, Antonio	Universidad de Huelva, España
Pezzano, Giuseppe	Universitat de Barcelona, España
Pla, Filiberto	Universidad Jaume I, España
Pomares, Hector	Universidad de Granada, España
Puerta, Jose M	Universidad de Castilla-La Mancha, España
R. Vela, Camino	Universidad de Oviedo, España
Rizo, Ramon	Universidad de Alicante, España
Rodenas, Javier	Universidad de Barcelona, España
Rodríguez, Rosa	Universidad de Jaén, España
Rosso, Paolo	Universitat Politècnica de València, España
Salmeron, Antonio	Universidad de Almería, España
Sanchez, Jose Salvador	Universitat Jaume I, España

Sanchez, Luciano
 Saquete, Estela
 Segarra, Encarna
 Soria, Emilio
 Talavera, Estefania
 Ureña-López, L. Alfonso
 Valencia-Garcia, Rafael
 Ventura, Sebastián
 Verdegay, José Luis
 Zaided, Mourad

Universidad de Oviedo, España
 Universidad de Alicante, España
 Universitat Politècnica de València, España
 Universidad de Valencia, España
 Universidad de Groningen, Países Bajos
 Universidad de Jaén, España
 Universidad de Murcia, España
 Universidad de Córdoba, España
 Universidad de Granada, España
 Universidad de Gabes, Túnez

ESTYLE

Alcala, Rafael
 Alcala-Fdez, Jesus
 Alcalde, Cristina
 Alonso, Jose M.
 Barrenechea, Edurne
 Barro, Senén
 Benítez, José M.
 Bobillo, Fernando
 Bugarín, Alberto
 Burusco, Ana
 Bustince, Humberto
 Cabrera, Inma P.
 Cabrerizo, Francisco Javier
 Calvo, Tomasa
 Carmona, Pablo
 Castro-Schez, Jose Jesus
 Cordero, Pablo
 Cornejo Piñero, Maria Eugenia
 Couso, Inés
 Cubillo, Susana
 De Andres, Rocio
 De Soto, Adolfo R.
 Del Jesus, Maria Jose
 Díaz, Susana
 Elorza, Jorge
 Escaño, Juan Manuel
 Fuente, Maria
 Galar, Mikel
 García-Lapresta, José Luis
 Godo, Lluís
 Gomez, Daniel
 Gonzalez, Antonio
 Indurain, Esteban
 Jurío, Aránzazu
 Lamata, Maria Teresa
 Liern, Vicente
 Llamazares, Bonifacio
 Lopez-Molina, Carlos
 Magdalena, Luis
 Martínez, Luis
 Marín, Nicolás
 Massanet, Sebastià
 Medina, Jesús
 Molina, Jose M.
 Montero, Javier
 Olivas, Jose Angel
 Paternain, Daniel

Universidad de Granada, España
 Universidad de Granada, España
 Universidad del País Vasco, España
 Universidad de Santiago de Compostela, España
 Universidad Pública de Navarra, España
 Universidad de Santiago de Compostela, España
 Universidad de Granada, España
 Universidad de Zaragoza, España
 Universidade de Santiago de Compostela, España
 Universidad Pública de Navarra, España
 Universidad Pública de Navarra, España
 Universidad de Málaga, España
 Universidad de Granada, España
 Universidad de Alcalá, España
 Universidad de Extremadura, España
 Universidad de Castilla-La Mancha, España
 Universidad de Málaga, España
 Universidad de Cádiz, España
 Universidad de Oviedo, España
 Universidad Politécnica de Madrid, España
 Universidad de Salamanca, España
 Universidad de León, España
 Universidad de Jaén, España
 Universidad de Oviedo, España
 Universidad de Navarra, España
 Universidad de Sevilla, España
 Universidad de Valladolid, España
 Universidad Pública de Navarra, España
 Universidad de Valladolid, España
 IIA - CSIC, España
 Universidad Complutense de Madrid, España
 Universidad de Granada, España
 Universidad Pública de Navarra, España
 Universidad Pública de Navarra, España
 Universidad de Granada, España
 Universidad de Valencia, España
 Universidad de Valladolid, España
 Universidad Pública de Navarra., España
 Universidad Politécnica de Madrid, España
 Universidad de Jaén, España
 Universidad de Granada, España
 Universidad de las Islas Baleares, España
 Universidad de Cádiz, España
 Universidad Carlos III de Madrid, España
 Universidad Complutense de Madrid, España
 Universidad de Castilla-La Mancha, España
 Universidad Pública de Navarra, España

Perez, Raul	Universidad de Granada, España
Pomares, Hector	Universidad de Granada, España
Porcel, Carlos	Universidad de Granada, España
Pradera, Ana	Universidad Rey Juan Carlos, España
Pérez-Fernández, Raúl	Universidad de Oviedo, España
Ramírez Poussa, Eloisa	Universidad de Cádiz, España
Recasens, Jordi	Universidad Politécnica de Cataluña, España
Riera, Juan Vicente	Universidad de las Islas Baleares, España
Romero, Francisco P.	Universidad de Castilla-La Mancha, España
Sainz-Palmero, Gregorio	Universidad de Valladolid, España
Sanchez, Daniel	Universidad de Granada, España
Sanchez, Luciano	Universidad de Oviedo, España
Sanz Delgado, José Antonio	Universidad Pública de Navarra, España
Serrano-Guerrero, Jesus	Universidad de Castilla-La Mancha, España
Sicilia, Miguel-Angel	Universidad de Alcalá, España
Sobrino Cerdeiriña, Alejandro	Universidad de Santiago de Compostela, España
Sánchez Solano, Santiago	IMSE-CNM, España
Torra, Vicenc	Universidad de Skövde, Suecia
Torrens, Joan	Universidad de las Islas Baleares, España
Valls, Aida	Universidad Rovira i Virgili, España
Verdegay, José Luis	Universidad de Granada, España

MAEB

Alvarez, Ada	Universidad Autónoma de Nuevo León, México
Amaya, Jhon Edgar	Universidad Nacional Experimental del Táchira, Venezuela
Araujo, Lourdes	Universidad Nacional de Educación a Distancia, España
Bautista, Joaquín	Universitat Politècnica de Catalunya, España
Bello Orgaz, Gema	Universidad Politécnica de Madrid, España
Blum, Christian	Consejo Superior de Investigaciones Científicas, España
Brito, Julio	Universidad de La Laguna, España
Caballero, Rafael	Universidad de Málaga, España
Camacho, Carlos	Universidad Politécnica de Madrid, España
Casado, Silvia	Universidad de Burgos, España
Castillo, Pedro	Universidad de Granada, España
Chavez, Francisco	Universidad de Extremadura, España
Chica, Manuel	Universidad de Granada, España
Coello Coello, Carlos	CINVESTAV-IPN, México
Colmenar, José Manuel	Universidad Rey Juan Carlos, España
Colomine D., Feijoo	Universidad Nacional experimental del Táchira, Venezuela
Damas, Sergio	Universidad de Granada, España
Del-Jesus, María-José	Universidad de Jaén, España
Dorronsoro, Bernabe	Universidad de Cádiz, España
Duarte, Abraham	Universidad Rey Juan Carlos, España
Duro, Richard	Universidade da Coruna, España
Egea, Jose A.	CEBAS-CSIC, España
Faulin, Javier	Universidad Pública de Navarra, España
Fernandez De Vega, Francisco	Universidad de Extremadura, España
Fernández, Alberto	Universidad de Granada, España
Fernández Leiva, Antonio J.	Universidad de Málaga, España
Gámez, José Antonio	Universidad de Castilla-La Mancha, España
García, Pablo	Universidad de Granada, España
García Pardo, Eduardo	Universidad Politécnica de Madrid, España
García-Martínez, Carlos	Universidad de Córdoba, España
García-Pedrajas, Nicolás	Universidad de Córdoba, España
Gomez Romero, Juan	Universidad de Granada, España
Gonzalez, Antonio	Universidad de Granada, España
Gonzalez, Pedro	Universidad de Jaén, España
Gonzalez-Pardo, Antonio	Universidad Rey Juan Carlos, España
Gutierrez, Pedro Antonio	Universidad de Córdoba, España

Hernández-Díaz, Alfredo G.	Universidad Pablo de Olavide, España
Herrera, Francisco	Universidad de Granada, España
Hervas, Cesar	Universidad de Córdoba, España
Hidalgo, Ignacio	Universidad Complutense de Madrid, España
Huertas, Javier	Universidad Politécnica de Madrid, España
Juan, Angel	Universitat Oberta de Catalunya, España
Laguna, Manuel	Universidad de Colorado Boulder, Estados Unidos
Larranaga, Pedro	Universidad de Madrid, España
Luna, Francisco	Universidad de Málaga, España
Luna, Jose Maria	Universidad de Córdoba, España
Luque, Mariano	Universidad de Málaga, España
Luque-Baena, Rafael	Universidad de Extremadura, España
Magdalena, Luis	Universidad Politécnica de Madrid, España
Martín, Alejandro	Universidad Politécnica de Madrid, España
Melian, Belen	Universidad de La Laguna, España
Mendiburu, Alexander	Universidad del País Vasco, España
Mesejo, Pablo	Universidad de Granada, España
Molina-Solana, Miguel	Universidad de Granada, España
Moreno, José-A.	Universidad de La Laguna, España
Nebro, Antonio J.	Universidad de Málaga, España
Pacheco, Joaquín	Universidad de Burgos, España
Peiró, Juanjo	Universidad de Valencia, España
Pelta, David	Universidad de Granada, España
Peregrin, Antonio	Universidad de Huelva, España
Perez, Raul	Universidad de Granada, España
Ramalhinho, Helena	Universitat Pompeu Fabra, España
Ramírez-Atencia, Cristian	Universidad Politécnica de Madrid, España
Riquelme, José C.	Universidad de Sevilla, España
Risco-Martín, José L.	Universidad Complutense de Madrid, España
Rivas Santos, Victor Manuel	Universidad de Jaén, España
Rodríguez, David	Universidad Nacional Experimental del Táchira, Venezuela
Rodríguez-Fernández, Víctor	Universidad Politécnica de Madrid, España
Romano, Carlos-Andrés	Universidad Politécnica de Valencia, España
Rosete, Alejandro	Instituto Superior Politécnico “José Antonio Echevarría”, Cuba
Rosso, Paolo	Universitat Politècnica de València, España
Ruiz, Ruben	Universitat Politècnica de València, España
Ríos-Mercado, Roger Z.	Universidad Autónoma de Nuevo León, México
Salcedo-Sanz, Sancho	Universidad de Alcalá, España
Sanchez, Luciano	Universidad de Oviedo, España
Santana, Roberto	Universidad del País Vasco, España
Stützle, Thomas	Université Libre de Bruxelles, Bélgica
Sáez, Yago	Universidad Carlos III de Madrid, España
Trujillo, Leonardo	Instituto Tecnológico de Tijuana, México
Udias, Angel	Universidad Rey Juan Carlos, España
Vega-Rodríguez, Miguel A.	Universidad de Extremadura, España
Ventura, Sebastián	Universidad de Córdoba, España
Verdegay, José Luis	Universidad de Granada, España
Villar, Pedro	Universidad de Granada, España
Villegas-Cortez, Juan	Universidad Autónoma Metropolitana, México
Winter, Gabriel	Universidad de Las Palmas de Gran Canaria, España
Zafra Gómez, Amelia	Universidad de Córdoba, España
Álvarez-Valdés, Ramón	Universidad de Valencia, España

TAMIDA

Arbelaitz, Olatz	Universidad del País Vasco, España
Bahamonde, Antonio	Universidad de Oviedo, España
Cano, Alberto	Virginia Commonwealth University, Estados Unidos
Corchado, Emilio	Universidad de Salamanca, España

Cordón, Óscar	Universidad de Granada, España
del Campo, José	Universidad de Málaga, España
Del Coz, Juan José	Universidad de Oviedo, España
del Jesús, M ^a José	Universidad de Jaén, España
Ferri, Francesc J.	Universidad de Valencia, España
García López, Salvador	Universidad de Granada, España
García Pedrajas, Nicolás	Universidad de Córdoba, España
Gibert, Karina	Universidad Politécnica de Catalunya, España
González, Pedro	Universidad de Jaén, España
Hernández Orallo, José	Universidad Politécnica de Valencia, España
Herrera, Francisco	Universidad de Granada, España
Hervás, César	Universidad de Córdoba, España
Larrañaga, Pedro	Universidad Politécnica de Madrid, España
Lozano, José Antonio	Universidad del País Vasco, España
Martínez Ballesteros, María	Universidad de Sevilla, España
Martínez Álvarez, Francisco	Universidad Pablo de Olavide, España
Pérez-Ortiz, María	University College London, Reino Unido
Ramírez, M ^a José	Universitat Politècnica de València, España
Rodríguez, Juan J.	Universidad de Burgos, España
Sánchez i Marrè, Miquel	Universidad Politécnica de Cataluña, España
Ventura, Sebastián	Universidad de Córdoba, España

Taller IABiomed

Araujo, Lourdes	Universidad Nacional de Educación a Distancia, España
Díez, Francisco Javier	Universidad Nacional de Educación a Distancia, España
Fernández-Breis, Jesualdo Tomás	Universidad de Murcia, España
Guijarro-Berdiñas, Bertha	Universidad de La Coruña, España
Larranaga, Pedro	Universidad de Madrid, España
López, Beatriz	Universidad de Girona, España
Martinez Tomas, Rafael	Universidad Nacional de Educación a Distancia, España
Muguerza, Javier	Universidad del País Vasco, España
Palma, Jose	Universidad de Murcia, España
Rodriguez, Alejandro	Universidad Politécnica de Madrid, España
Taboada, Maria	Universidad de Santiago de Compostela, España
Vellido, Alfredo	Universitat Politècnica de Catalunya, España
Villar, José Ramón	Universidad de Oviedo, España

Taller IA-SIT

Armingol Moreno, Jose Maria	Universidad Carlos III de Madrid, España
Molina Lopez, Jose Manuel	Universidad Carlos III de Madrid, España

Taller AI4D

Armuelles, Ivan	Universidad de Panama, Panama
Carbó, Javier	Universidad Carlos III de Madrid, España
Corrales, Juan Carlos	Universidad del Cauca, Colombia
Lopez-Vargas, Ascension	Universidad Carlos III de Madrid, España
Škrjanc, Igor	Universidad de Ljubljana, Eslovenia

Taller QUARTI

DAHI, Zakaria Abdelmoiz	Universidad de Málaga, España
Garcia de Andoin, Mikel	Universidad del País Vasco, España
Torrontegui, Erik	Universidad Carlos III de Madrid, España

Taller AGM

Bazazian, Dena	Universidad de Bristol, Reino Unido
Keane, Mark	University College Dublin, Irlanda
Pintore, Gianni	Center for Advanced Studies, Research and Development in Sardinia, Italia

Ribelles, Jose
Seoane, Antonio

Universitat Jaume I, España
Universidade da Coruña, España

Taller SCAI

Cintrano, Christian
Morell, José Angel
Muñoz, Antonio
Rossit, Diego Gabriel

Universidad de Málaga, España
Universidad de Málaga, España
Universidad de Málaga, España
Universidad Nacional del Sur - CONICET, Argentina

Jurado de Competiciones

Corchado, Juan Manuel
Cordón, Óscar
del Jesús, María José
Guijarro-Berdiñas, Bertha
Luaces, Óscar
Onaindia, Eva
Troncoso, Alicia

Universidad de Salamanca, España
Universidad de Granada, España
Universidad de Jaén, España
Universidad de La Coruña, España
Universidad de Oviedo, España
Universitat Politècnica de València, España
Universidad Pablo de Olavide, España

Jurado de Doctoral Consortium

Duarte, Abraham
Bahamonde, Antonio
Hervás, César
Bustince, Humberto
Gámez, José Antonio
Lozano, José Antonio
Riquelme, José
Rodríguez, Juan José
Gibert, Karina
Magdalena, Luis
del Jesús, María José
Ramírez, María José
Cordón, Óscar
Larrañaga, Pedro
Meseguer, Pedro
Martí, Rafael
Bolón, Verónica

Universidad Rey Juan Carlos, España
Universidad Oviedo, España
Universidad Córdoba, España
Universidad Pública de Navarra, España
Universidad Castilla-La Mancha, España
Universidad del País Vasco, España
Universidad Sevilla, España
Universidad de Burgos, España
Universidad Politécnica de Catalunya, España
Universidad Politécnica de Madrid, España
Universidad Jaén, España
Universidad Politécnica de Valencia, España
Universidad de Granada, España
Universidad Politécnica de Madrid, España
IIIA-CSIC, España
Universidad Valencia, España
Universidad de La Coruña, España

Índice general

Presentación de CAEPIA 20/21	III
Presentación de la Presidenta de AEPIA	V
Ponentes Plenarios	VII
Organización	IX
Comités de Programa	XIII

XIX Conferencia de la Asociación Española para la Inteligencia Artificial

Sesión General

Multi-Objective Bayesian Optimization approach for Environmental Monitoring using an Autonomous Surface Vehicle: Ypacarai Lake Case Study	5
<i>Federico Peralta, Samuel Yanes, Daniel Gutierrez Reina, Sergio Toral</i>	
Genetic Enhanced Model-based Deep Reinforcement Learning for Informative Path Planning in the Lake Ypacarai Scenario	11
<i>Samuel Yanes, Federico Peralta, Daniel Gutiérrez Reina, Sergio Toral</i>	
Bi-objective Evolutionary Approach for WaterResource Patrolling Problem with Autonomous	17
<i>Samuel Yanes, Federico Peralta, Daniel Gutiérrez, Alejandro Tapia, Álvaro Rodríguez, Sergio Toral</i>	
Sistema Inteligente para la Gestión Eficiente de Recogida de Residuos Plásticos	23
<i>Juan José Guerrero, Javier Ferrer, Rubén Saborido, Enrique Alba</i>	
Sistema de Recomendación con Explicaciones Basadas en Texto	29
<i>Pablo Pérez-Núñez, Antonio Bahamonde, Oscar Luaces, Jorge Díez</i>	
Método XAI basado en agrupamiento para la explicación de errores en clasificación	35
<i>Aurora Ramírez, Sebastián Ventura, José Raúl Romero</i>	
Revisión y análisis de conceptos en Inteligencia Artificial Explicable. Una aproximación a la unificación de la terminología	41
<i>Ángela Escobar Mimblera, Pedro González García, María José Del Jesús Díaz</i>	
Challenges of Definition Extraction in Spanish Legal Texts	47
<i>Karen Leticia Vázquez-Flores, Patricia Martín Chozas, Elena Montiel Ponsola</i>	
Bilingual Dataset for Information Retrieval and Question Answering over the Spanish Workers Statute	53
<i>Pablo Calleja, Patricia Martín Chozas, Elena Montiel-Ponsoda, Victor Rodríguez Doncel</i>	
Towards Document Entity Recognition using Close Domain Transfer Learning	58
<i>Adrián Inés Armas, César Domínguez, Jónathan Heras, Eloy Javier Mata, Vico Pascual</i>	
Runners detection based on trajectory analysis	64
<i>Jonay Sarmiento-Ramírez, Javier Lorenzo-Navarro, Modesto Castrillon-Santana, Elena Sánchez-Nielsen</i>	
Evaluación Estética de Fotografías vía Deep Learning y enfoques probabilísticos	70
<i>Luis González Naharro, José Miguel Puerta Callejón, Fernando Rubio Perona, Maria Julia Flores Gallego</i>	

Prediction of Epiretinal Membrane from Retinal Fundus Images Using Deep Learning	76
<i>Ángela Casado-García, Manuel García-Domínguez, Jónathan Heras, Adrián Inés, Didac Royo, Miguel Ángel Zapata</i>	
Evaluation of a CNN+LSTM system for the classification of hand-washing steps	77
<i>Kevin Cikel, Mario Arzamendia, Derlis Gregor, Daniel Gutiérrez, Sergio Toral</i>	
Ultr métricas para la comparación contextualizada de imágenes binarias	83
<i>Sara Iglesias-Rey, Felipe Antunes-Santos, Arnau Mir-Fuentes, Carlos Lopez-Molina</i>	
Towards Combining Object Detection and Text Classification Models for Form Entity Recognition	89
<i>María Villota, Gonzalo Santamaría, César Domínguez, Jónathan Heras, Eloy Mata, Vico Pascual</i>	
AutoML para la composición de workflows diversos mediante programación genética gramatical	90
<i>Rafael Barbudo Lunar, Sebastián Ventura, José Raúl Romero</i>	
Autonomous Learning: Tackling Expansive Mutating Domains	96
<i>Harshal Bharatia</i>	
Learning A Swarm Formation Policy using Deep Reinforcement Learning	102
<i>Jasmina Rais, Fidel Aznar</i>	
Generalized SMOTE: A universal generation oversampling technique for all data types in imbalanced learning	108
<i>Carlos Cernuda, Daniel Reguera-Bakhache, Aitor Aguirre, Mikel Iturbe, Iñaki Garitano, Urko Zurutuza</i>	
Deep Reinforcement and Imitation Learning for Self-Driving Tasks	114
<i>Sergio Hernández-García, Alfredo Cuesta-Infante</i>	
A Proposal to Integrate Deep Q-Learning with Automated Planning to Improve the Performance of a Planning-based Agent	115
<i>Carlos Núñez-Molina, Ignacio Vellido, Vladislav Nikolov-Vasilev, Raúl Pérez, Juan Fdez-Olivares</i>	
Ontologies for run-time self-adaptation of mobile robotic systems	116
<i>Esther Aguado, Ricardo Sanz, Claudio Rossi</i>	
Analysing Ontological Requirements: A Journey from Requirements to Code and Back	122
<i>Alba Fernandez-Izquierdo, María Poveda-Villalón, Raúl García-Castro</i>	
Modelos gráficos probabilísticos basados en mixturas para el problema de partial label ranking	128
<i>Juan C. Alfaro, Juan A. Aledo, Jose A. Gámez</i>	
LabelDetection: Simplifying the Use and Construction of Deep Detection Models	134
<i>Ángela Casado-García, Jónathan Heras</i>	
Algoritmos para el problema del bandido multi-objetivo basados en la escalarización de Chebyshev	135
<i>Lawrence Mandow, Sergio Martín-Albo, Jose-Luis Perez De La Cruz</i>	
Caracterización de conjuntos de puntos mediante cierres no convexos subdivisibles	141
<i>David Novoa-Paradela, Oscar Fontenla-Romero, Bertha Guijarro-Berdiñas</i>	
A memetic algorithm to minimize the total weighted tardiness in the fuzzy flexible job shop	147
<i>Pablo García Gómez, María Camino Rodríguez Vela, Inés González Rodríguez</i>	
Estimación de las tensiones residuales a nivel de grano en una muestra cilíndrica templada de aleación de aluminio AA5083 mediante programación genética	153
<i>Laura Millán, Gabriel Kronberger, Ricardo Fernández Serrano, Oscar Garnica, Patricie Halodova, Gaspar González Doncel, Ignacio Hidalgo</i>	
Towards Automatic Bayesian Optimization: A First Step Involving Acquisition Functions	159
<i>Luis C. Jariego Pérez, Eduardo C. Garrido-Merchán</i>	
Diseñando múltiples genes para una proteína mediante el algoritmo memético SFLA	160
<i>Belen Gonzalez-Sanchez, Miguel A. Vega-Rodríguez, Sergio Santander-Jiménez</i>	
Recognition of Teaching Activities from University Lecture Transcriptions	166
<i>Daniel Diosdado, Alberto Romero, Eva Onaindia</i>	

Human Activity Recognition with Capsule Networks	167
<i>Laura Llopis-Ibor, Alfredo Cuesta-Infante, César Beltrán-Royo, Juan José Pantrigo</i>	
Identification of Contact Failures in Multilayered Media via Inverse Problem and Artificial Neural Networks	168
<i>Lucas Correia da Silva Jardim, Diego Campos Knupp, Roberto Pinheiro Domingos, Luiz Alberto da Silva Abreu, Carlos Cruz Corona, Antônio José da Silva Neto</i>	
Estimación de la biomasa en acuicultura utilizando redes neuronales convolucionales	174
<i>Samuel Lopez, Antonio Peregrin, Juan Carlos Gutierrez, Inmaculada Pulido, Jairo Castro, Fernando Gomez, Alejandro Garrocho</i>	
An unsupervised anomaly detection strategy for industrial processes under data shift scenarios	180
<i>Xabier Barrenetxea Berasategi, Alberto Diez-Olivan, Óscar Córdón</i>	
Train Route Planning as a Multi-Agent Path Finding Problem	186
<i>Mauricio Salerno, Yolanda E-Martín, Raquel Fuentetaja, Alba Gragera, Alberto Pozanco, Daniel Borrajo</i>	
Asynchronous Vector Iteration in Multi-Objective Markov Decision Processes	187
<i>Ekaterina Sedova, Lawrence Mandow, José-Luis Pérez-de-la-Cruz</i>	
A Similarity Measure of Gaussian Process Predictive Distributions	188
<i>Lucía Asencio Martín, Eduardo C. Garrido-Merchán</i>	
Towards Fairness in Classification: Comparison of Methods to Decrease Bias	189
<i>Maitane Martínez-Eguiluz, Oier Irazabal-Urrutia, Olatz Arbelaiz-Gallego</i>	
Modeling Administrative Discretion Using Goal-Directed Answer Set Programming	190
<i>Joaquín Arias, Mar Moreno-Rebato, José A. Rodríguez-García, Sascha Ossowski</i>	
Automatic Generation of Interrelated Organisms on Virtual Environments	191
<i>Santiago Pacheco, Nicolás Ottonello, Sergio Nesmachnow</i>	
An analysis of the indexes measuring the agreement of a profile of rankings	192
<i>Noelia Rico, Camino R. Vela, Irene Diaz</i>	
Trabajos Destacados	
A Hierarchical Architecture for Recognising Intentionality in Mental Tasks on a Brain-Computer Interface ..	201
<i>Asier Salazar-Ramirez, Jose I. Martin, Raquel Martinez Rodriguez, Andoni Arruti, Javier Muguerza, Basilio Sierra</i>	
A Self-Paced Relaxation Response Detection System Based on Galvanic Skin Response Analysis	203
<i>Raquel Martinez Rodriguez, Asier Salazar, Jose Ignacio Martin, Andoni Arruti, Eloy Irigoyen, Javier Muguerza</i>	
Improving Bayesian Inference Efficiency for Sensory Anomaly Detection and Recovery in Mobile Robots ...	205
<i>Manuel Castellano-Quero, Juan-Antonio Fernández-Madrigal, Alfonso García-Cerezo</i>	
Template-free detection and classification of membrane-bound complexes in cryo-electron tomograms	207
<i>Antonio Martinez-Sanchez</i>	
Improving fake faces detection to prevent fraudulent registrations	209
<i>Luis Carabe, Eduardo Cermeño</i>	
Hair Segmentation and Removal in Dermoscopic Images using Deep Learning	211
<i>Lidia Talavera-Martínez, Pedro Bibiloni, Manuel González-Hidalgo</i>	
Navigating at the microscopic scale with Reinforcement Learning	213
<i>Santiago Muñoz-Landín, Alexander Fischer, Viktor Holubek, Frank Cichos</i>	
Online Learning for Probabilistic Load Forecasting	215
<i>Verónica Álvarez, Santiago Mazuelas, José A. Lozano</i>	

CCE: An ensemble architecture based on coupled ANN for solving multiclass problems	217
<i>M. Paz Sesmero, Juan M. Alonso-Weber, Araceli Sanchis</i>	
Reliable and Fast Recurrent Neural Network Architecture Optimization	219
<i>Andrés Camero, Jamal Toutouh, Enrique Alba</i>	
An AI-Powered System for Residential Demand Response	221
<i>Iker Esnaola-Gonzalez, Marko Jelić, Dea Pujić, Francisco Javier Díez, Nikola Tomasevic</i>	
Sobre órdenes admisibles en el conjunto de números borrosos discretos y su aplicación en problemas de toma de decisiones	223
<i>Juan Vicente Riera, Sebastia Massanet, Humberto Bustince, Francisco Javier Fernandez</i>	
Aprendizaje Federado y Privacidad Diferencial: Análisis de herramientas software, el framework Sherpa.ai FL y guías metodológicas para preservar la privacidad de datos	225
<i>Nuria Rodríguez-Barroso, Eugenio Martínez-Cámara, María V. Luzón, Francisco Herrera</i>	
FEPDS: Una propuesta para la extracción de patrones emergentes difusos en flujos continuos de datos	227
<i>Angel M. Garcia-Vico, Huseyin Seker, Cristobal J. Carmona, Pedro Gonzalez, María José Del Jesús</i>	
Miniaturización de aplicaciones JavaScript para el Internet de las Cosas	229
<i>Rubén Saborido</i>	
E2PAMEA: un algoritmo evolutivo para la extracción eficiente de patrones emergentes difusos en entornos big data	231
<i>Angel M. Garcia-Vico, David Elizondo, Francisco Charte, Pedro Gonzalez, Cristobal J. Carmona</i>	
Fostering Diversity in Spatial Evolutionary Generative Adversarial Networks	233
<i>Jamal Toutouh, Erik Hemberg, Una-May O'Reilly</i>	
Revisiting data complexity metrics based on morphology for overlap and imbalance: snapshot, new Overlap Number of Balls metrics and singular problems prospect	235
<i>José Daniel Pascual Triana, David Charte, Marta Andrés Arroyo, Alberto Fernández, Francisco Herrera</i>	

XX Congreso Español sobre Tecnologías y Lógica Fuzzy

Sesión General

Multiple fuzzy Sugeno λ -measures in networks	241
<i>Inmaculada Gutiérrez, Daniel Gomez, Javier Castro, Rosa Espínola</i>	
Explicaciones factuales y contrafactuales en árboles de clasificación difusos	247
<i>Guillermo Fernández, Juan Aledo, José Gámez, José Puerta</i>	
Comparando variabilidades de conjuntos aleatorios	253
<i>Juan Jesús Salamanca, Susana Montes</i>	
A fuzzy probability logic for compound conditionals	256
<i>Tommaso Flaminio, Lluís Godo</i>	
Empleo de contextos L-fuzzy en el diseño de materiales de origen renovable	262
<i>Itsaso Leceta, Cristina Alcalde, Marta Urdanpilleta, Pedro Guerrero, Koro de la Caba, Ana Burusco</i>	
Una versión ternaria del algoritmo de Quine McCluskey para la minimización de base de reglas difusas	268
<i>Leonardo Jara, Antonio González, Raúl Pérez</i>	
Prometheus: Harnessing Fuzzy Logic and Natural Language for Human-centric Explainable Artificial Intelligence	274
<i>Ettore Mariotti, Jose M. Alonso-Moral, Albert Gatt</i>	
Modelando el comportamiento del consumidor con el modelo de 2-tuplas lingüísticas y una heurística basada en el conocimiento de marca	280
<i>Jesús Giráldez-Cru, Manuel Chica, Óscar Cordon, Francisco Herrera</i>	

Generación de descripciones lingüísticas de la contaminación acústica diaria en zonas urbanas	286
<i>Juan Moreno Garcia, Luis Jimenez Linares, Luis Rodriguez-Benitez</i>	
Un sistema para la descripción automática en lenguaje natural de gráficas de sectores: aplicación en datos de calidad del aire	292
<i>Andrea Cascallar-Fuentes, Javier Gallego-Fernández, Alejandro Ramos-Soto, Anthony Saunders, Alberto Bugarín-Diz</i>	
Evaluación empírica de modelos de cuantificación borrosa aplicados a un agente conversacional	298
<i>Mariña Canabal-Juanatey, Jose M. Alonso-Moral, Alejandro Catala, Alberto Bugarín-Diz</i>	
Hypothesis Scoring and Model Refinement Strategies for FM-based RANSAC	304
<i>Alberto Ortiz, Esaú Ortiz, Juan José Miñana, Óscar Valero</i>	
Análisis de distintos tipos de tendencias en sucesiones de contextos L-fuzzy	305
<i>Cristina Alcalde, Ana Burusco</i>	
On the standard fuzzy metric: generalizations and application to model estimation	311
<i>Juan-José Miñana, Alberto Ortiz, Esaú Ortiz, Oscar Valero</i>	
Uso de t-normas para el estudio de la convexidad en conjuntos difusos intervalo-valuados	317
<i>Pedro Huidobro, Pedro Alonso, Humberto Bustince, Vladimír Janis, Susana Montes</i>	
Sesión especial: Toma de decisiones con información difusa	
Preferencias no lineales en Toma de Decisión en Grupo. Amplificación de Valores Extremos	323
<i>Diego García-Zamora, Álvaro Labella Romero, Rosa M. Rodríguez, Luis Martínez</i>	
Estructuras de preferencia para relaciones recíprocas borrosas: caracterización y compatibilidad	329
<i>Fabian Alberto Castiblanco Ruiz, Camilo Franco De Los Ríos, J. Tinguaro Rodríguez, Javier Montero</i>	
Generalización de intervalos en matrices de riesgo imprecisas	331
<i>María Martínez, Juan Baz, Raul Perez-Fernandez, Susana Montes</i>	
Consenso en Entornos Académicos Virtuales: Toma de Decisiones Grupales en el Trabajo de Clase Online ..	337
<i>Aitor Manuel Orellana, Antonio Velez-Estevez, Manuel Jesus Cobo, Ignacio Javier Perez</i>	
Sesión especial: Funciones de agregación y preagregación y sus aplicaciones	
Mapas de características adicionales en redes neuronales profundas a partir de agregaciones basadas en orden	345
<i>Iris Dominguez-Catena, Daniel Paternain, Mikel Galar</i>	
Learning OWA weights by combining fuzzy quantifiers with empirical data	351
<i>Álvaro Cristóbal, Ignacio Huitzil, Fernando Bobillo</i>	
Integrales de Choquet y capacidades 2-aditivas en la construcción de indicadores sintéticos	357
<i>Patrizia Pérez-Asurmendi, Bonifacio Llamazares Rodríguez</i>	
Extensión multidimensional de la integral de Choquet discreta y su aplicación en redes neuronales recurrentes	358
<i>Mikel Ferrero-Jaurrieta, Iosu Rodríguez-Martínez, Graçaliz Pereira Dimuro, Javier Fernández, Humberto Bustince</i>	
Monotonicity in Non-deterministic Computable Aggregations	364
<i>Luis Magdalena, Daniel Gomez, Luis Garmendia, Javier Montero</i>	
Optimizando Desviaciones Moderadas Ponderadas para Interfaces Cerebro Ordenador	370
<i>Javier Fumanal Idocin, Carmen Vidaurre, Marisol Gómez, Asier Urio, Graçaliz Pereira-Dimuro, Humberto Bustince</i>	
Operador de comparación de elementos multivaluados basado en Funciones de Equivalencia Restringsida	376
<i>Aitor Castillo-López, Carlos Lopez-Molina, Javier Fernandez, Humberto Bustince, Mikel Sesma-Sara</i>	

Trabajos Destacados

Computing with Comparative Linguistic Expressions and Symbolic Translation for Decision Making: ELICIT Information	385
<i>Álvaro Labella Romero, Rosa M. Rodríguez, Luis Martínez</i>	
Aggregation of fuzzy quasi-metrics	387
<i>Tatiana Pedraza, Jesús Rodríguez-López, Óscar Valero</i>	
A graded notion of quasi-intents	389
<i>Manuel Ojeda-Hernandez, Inma P. Cabrera, Pablo Cordero</i>	
Community Detection Problem Based on Polarization Measures: An Application to Twitter. The COVID-19 Case in Spain. A Review	391
<i>Inmaculada Gutiérrez, Juan Antonio Guevara, Daniel Gómez, Javier Castro, Rosa Espínola</i>	
Rough-Set-Driven Approach for Attribute Reduction in Fuzzy Formal Concept Analysis	393
<i>María José Benítez-Caballero, Jesús Medina, Eloisa Ramírez-Poussa, Dominik Slezak</i>	
Measuring Consistency of Fuzzy Logic Theories	395
<i>Nicolas Madrid, Manuel Ojeda-Aciego</i>	
Detecting equivalence classes being sublattices from an attribute reduction in FCA	397
<i>Roberto García-Aragón, Jesús Medina, Eloisa Ramírez Poussa</i>	
Non-monotonic multi-adjoint logic programming with constraints	398
<i>M. Eugenia Cornejo, David Lobo, Jesús Medina</i>	

XIV Congreso Español de Metaheurísticas, Algoritmos Evolutivos y Bioinspirados

Sesión General

On the role of the gradient optimizer on evolved GANs	405
<i>Unai Garciarena, Alexander Mendiburu, Roberto Santana</i>	
Aplicación de técnicas de aprendizaje profundo al reconocimiento óptico de partituras SATB	411
<i>Martín Morita Hernández, Francisco Fernández de Vega, Juan Villegas Cortez</i>	
Optimización Evolutiva de parámetros sobre Random Forest en Entornos de Big Data	417
<i>Miguel Angel Rodriguez, Antonio Peregrin</i>	
Un estudio de los diferentes criterios aplicables al problema de resumen extractivo multi-documento	423
<i>Jesus M. Sanchez-Gomez, Miguel A. Vega-Rodríguez, Carlos J. Pérez</i>	
Algoritmos Evolutivos con Memoria Elite para la Selección de Carteras de Inversión	429
<i>Feijoo Colomine, Carlos Cotta, Antonio J. Fernández Leiva</i>	
Un Enfoque Evolutivo Multi-Operador para el Diseño de Rutas Turísticas con Nivel de Interés Dependiente del Tiempo	434
<i>David Pelta, Juan Carlos Ortégón, Marcelino Cabrera</i>	
Cooperación Inter-algorítmica Multiobjetivo para Reconstrucción Filogenética	440
<i>Sergio Santander-Jiménez, Miguel A. Vega-Rodríguez, Leonel Sousa</i>	
El efecto de la multiconectividad en el problema del apagado selectivo de redes 5G ultradensas	446
<i>Alejandro Quiñones, Jesús Galeano, Pablo Zapata, Francisco Luna, Javier Carmona, Juan Francisco Valenzuela</i>	
Influence of the Alternative Objective Functions in the Optimization of the Cyclic Cutwidth Minimization Problem	452
<i>Sergio Cavero, Eduardo G. Pardo, Abraham Duarte</i>	
Bayesian Optimization for Permutations Spaces: A preliminary approach	453
<i>Daniel Jiménez, Josu Ceberio, Jose Antonio Lozano</i>	

Analyzing Elementary Landscapes in Graph Partitioning Problems	459
<i>Ignacio Dorado Llamas, Josu Ceberio Uribe</i>	
A Matheuristic Approach to the Interval Job Shop Scheduling Problem	465
<i>Sezin Afsar, Juan Jose Palacios, Camino R. Vela</i>	
Solving the Permutation Heijunka Flow Shop Scheduling Problem with Non-unit Demands for Jobs	471
<i>Joaquín Bautista-Valhondo</i>	
Combinando GRASP con Iterated Greedy para el problema de la máxima intersección de k-conjuntos	472
<i>Alejandra Casado Ceballos, Sergio Pérez Peló, Jesús Sánchez-Oro Calvo, Abraham Duarte Muñoz</i>	
Optimización multiobjetivo basada en el algoritmo ABC para la codificación de proteínas	478
<i>Belen Gonzalez-Sanchez, Miguel A. Vega-Rodríguez, Sergio Santander-Jiménez</i>	
Impacto Ambiental de la Experimentación con Metaheurísticas: Un Análisis Coste-Beneficio	484
<i>Pavel Novoa-Hernández, David A. Pelta</i>	
Un enfoque multi-objetivo para el problema de la máxima diversidad	490
<i>Pedro Casas-Martínez, Alejandra Casado-Ceballos, Jesús Sánchez-Oro, Eduardo G. Pardo</i>	

Sesión especial: Algoritmos Bioinspirados Profundos

Una Visión Panorámica de la Meta-Cooperación Profunda con Algoritmos Meméticos	499
<i>Jhon Edgar Amaya, Carlos Cotta, Antonio J. Fernández Leiva, Pablo García Sánchez</i>	
En busca de algoritmos genéticos eficientes desde el punto de vista energético	505
<i>Jorge Alvarado Díaz, Francisco Chávez, Pedro Castillo, Josefa Díaz, Francisco Fernández de Vega</i>	
Clusterig Cosmológico: un enfoque del clustering gravitacional clásico inspirado en la estructura y dinámica del cosmos a gran escala	511
<i>Aitor Castillo-López, Javier Fumanal-Idocin, Javier Fernandez, Humberto Bustince</i>	
Cartografía de un Subcampo Científico: Un Estudio sobre la Optimización de Carteras de Inversión	517
<i>Feijoo Colomine, Carlos Cotta, Antonio J. Fernández Leiva</i>	
A la desinformación le gusta la compañía: Representación de bulos de Twitter sobre la COVID-19 mediante embeddings	523
<i>Guillermo Villar-Rodríguez, Javier Huertas-Tato, Alejandro Martín, David Camacho</i>	
Mejorando Redes Profundas para Aprendizaje por Transferencia mediante Capas Dispersas optimizadas con Algoritmos Evolutivos	529
<i>Javier Poyatos Amador, Daniel Molina Cabrera, Aritz Martínez, Siham Tabik, Javier Del Ser Lorente, Francisco Herrera Triguero</i>	

Sesión especial: Aplicaciones en Medicina y Biotecnología

Predicción de la glucosa en sangre mediante un sistema basado en reglas difusas TSK de dos fases	537
<i>Jorge Alvarado Díaz, José Manuel Velasco, Francisco Chávez, Almudena Sánchez, Francisco Fernández de Vega</i>	
Modelo híbrido basado en aprendizaje profundo de creación y etiquetado de células escamosas procedentes de citologías cervicovaginales	543
<i>Andrés Bueno-Crespo, Ana Ortiz-González, Raquel Martínez España, José Martínez-Más, Miguel Ángel Gragera, Manuel Remezal-Solano, Juan-Pedro Martínez-Cendán, Sebastián Ortiz-Reina, Andres Muñoz</i>	
Selección óptima de variables mediante computación evolutiva para algoritmos de clasificación	549
<i>Daniel Parra, Alberto Gutiérrez, Oscar Garnica, Juan Lanchares, José Ignacio Hidalgo</i>	
Calibración y predicción probabilística de un modelo de crecimiento del cáncer de vejiga teniendo en cuenta la incertidumbre de los datos	555
<i>Juan Carlos Cortés, Elena López-Navarro, Ana Moscardó-García, Rafael-J. Villanueva</i>	

Heurística para el desarrollo de un sistema experto epidemiológico aplicado al COVID-19	560
<i>Beatriz González-Pérez, Concepción Nuñez, José L. Sánchez, Gabriel Valverde, Jose Manuel Velasco</i>	
Calibración de un modelo basado en agentes que describe la evolución de la resistencia del <i>Acinetobacter baumannii</i> a la colistina en la Ciudad de Valencia	566
<i>Juan A. Aledo, Carlos Andreu-Villaró, Juan Carlos Cortés, Juan C. Orengo, Rafael-J. Villanueva</i>	
A Blood Glucose classification and prediction system by Neural Networks	572
<i>Alvaro Delgado, Alejandro Varela, Félix Tena, Jose Manuel Velasco, J. Ignacio Hidalgo</i>	

Sesión especial: Estrategias de implementación eficiente de heurísticas y metaheurística

Heurísticas para la mejora de la mantenibilidad de proyectos software	581
<i>Javier Yuste, Eduardo G. Pardo, Abraham Duarte</i>	
Un algoritmo eficiente para el problema de disposición de instalaciones en dos filas	587
<i>Raúl Martín-Santamaría, Alberto Herrán, Abraham Duarte, José Manuel Colmenar</i>	
Algoritmos de estimación del tiempo de ventana en la recogida de pedidos en almacenes logísticos	593
<i>Sergio Gil Borrás, Eduardo García Pardo, Antonio Alonso Ayuso, Abraham Duarte, Ernesto Jiménez Merino</i>	

Sesión especial: Metaheurísticas Aplicadas al Análisis de Redes Sociales

Nuevos algoritmos metaheurísticos para el análisis de la influencia de los usuarios en las redes sociales	601
<i>Isaac Lozano-Osorio, Jesús Sánchez-Oro, Abraham Duarte, Óscar Cordón</i>	
Detección de comunidades con un enfoque multi-objetivo utilizando VNS	607
<i>Sergio Pérez-Peló, Jesús Sánchez-Oro, Antonio Gonzalez-Pardo, Abraham Duarte</i>	
Análisis de Redes Sociales basado en las Conquistas de César Borgia	613
<i>Javier Idocin, Oscar Cordon, Amparo Alonso-Betanzos, Humberto Bustince, Javier Fernandez</i>	
Evaluación de modelos multilingües pre-entrenados en similitud semántica para la lucha contra la desinformación	619
<i>Álvaro Huertas-García, Alejandro Martín, Javier Huertas-Tato, David Camacho</i>	

Trabajos Destacados

Hybrid GRASP-SO-PR algorithm. The closest vs. the most balanced	627
<i>Ana Dolores López-Sánchez, Jesús Sánchez-Oro, Anna Martínez-Gavara, Alfredo G. Hernández-Díaz, Abraham Duarte</i>	
Complete Shipment constraints in Container Loading Problem	629
<i>Iván Giménez-Palacios, María Teresa Alonso, Ramon Alvarez-Valdes, Francisco Parreño</i>	
Scatter search for the capacitated dispersion problem	631
<i>Anna Martínez-Gavara, Jesús Sánchez-Oro, Rafael Martí</i>	
A review and empirical analysis of diversity and dispersion problems	633
<i>Rafael Martí, Anna Martínez-Gavara, Sergio Pérez-Peló, Jesús Sánchez-Oro</i>	

X Simposio de Teoría y Aplicaciones de Minería de Datos

Sesión General

Nearest Neighbors-based Forecasting for Electricity Demand Time Series in Streaming	639
<i>L. Melgar-García, D. Gutiérrez-Avilés, C. Rubio-Escudero, A. Troncoso</i>	
STree: a Single Multi-class Oblique Decision Tree Based on Support Vector Machines	640
<i>Ricardo Montañana, José A. Gámez, José M. Puerta</i>	

Studying the Effect of Different L_p Norms in the Context of Time Series Ordinal Classification	641
<i>David Guijo-Rubio, Víctor Manuel Vargas, Pedro Antonio Gutiérrez, César Hervás-Martínez</i>	
Universal Adversarial Examples in Speech Command Classification	642
<i>Jon Vadillo, Roberto Santana</i>	
Automatic Labelling using QA	648
<i>Eugenia Taillefer, Manuel Baena-García, Rafael Morales-Bueno</i>	
Reemplazo de la función de pooling de Redes Neuronales Convolucionales por combinaciones lineales de funciones crecientes	652
<i>Iosu Rodríguez, Julio Lafuente, Mikel Sesma, Francisco Herrera, Pablo Ursúa, Humberto Bustince</i>	
Evaluation of the Transformer Architecture for Univariate Time Series Forecasting	658
<i>Pedro Lara-Benítez, Luis Gallego-Ledesma, Manuel Carranza-García, José M. Luna-Romera</i>	
Electricity Consumption Time Series Forecasting Using Temporal Convolutional Networks	659
<i>J. F. Torres, M. J. Jiménez-Navarro, F. Martínez-Álvarez, A. Troncoso</i>	
ReLU-based Activations: Analysis and Experimental Study for Deep Learning	660
<i>Víctor Manuel Vargas, David Guijo-Rubio, Pedro Antonio Gutiérrez, César Hervás-Martínez</i>	
Clustering: Un paquete R para facilitar el análisis de algoritmos de agrupamiento	661
<i>Luis Alfonso Pérez Martos, Pedro González García, Cristóbal José Carmona del Jesús</i>	
Computer Science doctoral research in Spain. A relational perspective	667
<i>Adrián Arnaiz-Rodríguez, José Miguel Ramírez-Sanz, José Luis Garrido-Labrador, Alicia Olivares-Gil</i>	
Estudio comparativo de medidas de disimilitud para Clustering Multi-Instancia	673
<i>Aurora Esteban Toscano, Amelia Zafra Gómez, Sebastián Ventura Soto</i>	
A cluster merge algorithm to classify on/offline records of researchers in programming training	679
<i>Sonia Estévez, Victoria Lopez</i>	
Valoración de Anomalías Multi-paso Basada en Histogramas para Detección de Anomalías no Supervisada ..	685
<i>Ignacio Aguilera-Martos, Marta García-Barzana, Diego García-Gil, Jacinto Carrasco, David López, Manuel Arias-Rodil, Julián Luengo, Francisco Herrera</i>	
Estudio experimental sobre XGBoost, Rotation Forest y su combinación	691
<i>Francisco Javier González-Moya, Álar Arnaiz-González, Juan J. Rodríguez</i>	
Estudio de estrategias basadas en vecinos para clasificación multi-instancia multi-etiqueta	697
<i>Eva Gibaja, Amelia Zafra</i>	
Comparing BERT against traditional machine learning text classification	703
<i>Santiago González-Carvajal, Eduardo C. Garrido-Merchán</i>	
Aprendizaje de clasificadores binarios justos mediante Optimización Multiobjetivo	706
<i>Jorge Casillas, David Villar</i>	
Análisis predictivo efectivo en centros de datos: una solución basada en inteligencia artificial	712
<i>David Garcia-Retuerta, Roberto Casado-Vara, Juan M. Corchado</i>	
Uso de técnicas de Inteligencia Artificial en el ámbito de la Atención Temprana	717
<i>Ignacio Sierra, Norberto Díaz-Díaz, Carlos D. Barranco, Rocío Carrasco-Villalón</i>	
Un sistema de apoyo al experto en actividad física basado en Metaheurísticas	723
<i>Joaquín Roiz Pagador, Roberto Ruiz, África Calvo-Lluch, Adán González-Reina</i>	
StreetQR: Informative assistance device for street name plates and places of interest	729
<i>Jesús Benito-Picazo, Jorge García-González, Gonzalo Ramos-Jiménez, Ezequiel López-Rubio</i>	
Personalización de Recomendaciones en TripAdvisor Usando Deep Learning	735
<i>Iñigo Luis López-Riobóo Botana, Amparo Alonso Betanzos, Verónica Bolón Canedo, Antonio Bahamonde Rionda</i>	

Sistema de Recomendación Transparente y Escrutable para Recomendaciones Personalizadas	741
<i>Antonio Rodríguez-Revuelta, Pablo Pérez-Núñez, Jorge Díez</i>	
Automatic Head Appearance Recognition for Individuals Affected by PIMD	747
<i>Carmen Campomanes-Álvarez, B. Rosario Campomanes-Álvarez, Pelayo Quirós</i>	
Selección de características en una base de datos sobre el cáncer colorrectal	753
<i>José Antonio Delgado-Osuna, Daniel Ranchal-Parrado, Carlos García-Martínez, Sebastián Ventura</i>	
Identification of outliers through classification based on evolution indices and their application to COVID-19	757
<i>Yury Andrea Jiménez Agudelo, Victoria Lopez, Walter Chanava</i>	
Open Data and Experimental data source combination's impact on underground water analysis in Isla Margarita, Venezuela	763
<i>Karina Gibert, Dante Conti, Iñigo Arregui Pérez de Loza</i>	
Análisis descriptivo de cáncer de mama usando minería de datos	769
<i>Antonio Manuel Trasierras, Jose Maria Luna, Sebastián Ventura</i>	
Aplicación de técnicas de Aprendizaje Automático para el análisis de datos de acciones de juego en voleibol	775
<i>Francisco Aragón Royón, Elia Mercado Palomino, Aurelio Ureña Espá, José M. Benítez</i>	

Trabajos Destacados

Rotation Forest for Big Data	783
<i>Mario Juez-Gil, Álgar Arnaiz-González, Juan J. Rodríguez, Carlos López-Nozal, César García-Osorio</i>	
Rotation Forest for Multi-target Regression	785
<i>Juan J. Rodríguez, Mario Juez-Gil, Carlos López-Nozal, Álgar Arnaiz-González</i>	
Experimental evaluation of ensemble classifiers for imbalance in Big Data	787
<i>Mario Juez-Gil, Álgar Arnaiz-González, Juan J. Rodríguez, César García-Osorio</i>	

Talleres, Competiciones y Doctoral Consortium

II Taller de Grupos de investigación españoles de IA en Biomedicina (IABiomed)

Methods and measures to quantify ICU patient heterogeneity	793
<i>David Cuadrado, David Riaño, Josep Gomez, Alejandro Rodríguez, Maria Bodi</i>	
Neural Style Transfer to deal with the Domain Shift Problem on Glioblastoma Segmentation	795
<i>Manuel García, Cesar Dominguez, Jónathan Heras, Eloy Javier Mata, Vico Pascual</i>	
Fall detection for healthy and autonomous elderly people	800
<i>Mirko Fáñez Kertelj, José Ramón Villar, Enrique de La Cal, Víctor Gonzalez, Javier Sedano, Samad Barri Khojasteh</i>	
A platform for the design and execution of clinical data transformation and reasoning workflows	802
<i>José Alberto Maldonado, Mar Marcos, Jesualdo Tomás Fernández-Breis, Vicente Miguel Giménez-Solano, María del Carmen Legaz-García, Begoña Martínez-Salvador</i>	
DISNET: extracting public phenotypical knowledge of diseases	804
<i>Lucía Prieto Santamaría, Ernestina Menasalvas Ruiz, Alejandro Rodríguez-González</i>	
Using Machine Learning for the prediction of mutations in the initiation codon	806
<i>Javier Castell Díaz, Francisco Abad-Navarro, Jesualdo T. Fernández-Breis, M. Eugenia de la Morena Barrio, Javier Corral de la Calle</i>	
A Helping Hand to CDSS Rule Technologies in Fighting Antibiotic Resistance for Hospital Settings	812
<i>Bernardo Canovas-Segura, Natalia Iglesias, Antonio Morales, Jose M. Juarez, Manuel Campos</i>	

XX Congreso Español sobre Tecnologías y Lógica Fuzzy



XX Congreso Español sobre Tecnologías y Lógica Fuzzy

SESIÓN GENERAL



Multiple fuzzy Sugeno λ -measures in networks.

1st Inmaculada Gutiérrez

*Facultad de Estudios Estadísticos
Universidad Complutense de Madrid
Madrid, España
inmaguti@ucm.es
ORCID: 0000-0002-7103-393X*

2nd Daniel Gómez

*Facultad de Estudios Estadísticos
Instituto de Evaluación Sanitaria
Universidad Complutense de Madrid
Madrid, España
dagomezc@estad.ucm.es
ORCID: 0000-0001-9548-5781*

3rd Javier Castro

*Facultad de Estudios Estadísticos
Instituto de Evaluación Sanitaria
Universidad Complutense de Madrid
Madrid, España
jcastroc@estad.ucm.es
ORCID: 0000-0002-5345-8870*

4th Rosa Espínola

*Facultad de Estudios Estadísticos
Instituto de Evaluación Sanitaria
Universidad Complutense de Madrid
Madrid, España
rosaev@estad.ucm.es
ORCID: 0000-0003-2336-3659*

Abstract—In this paper we take up the community detection problem based on fuzzy measures. We focus on the existence of a family of vectors which define some additional information about the individuals of a network, from which we obtain multiple fuzzy Sugeno λ -measures. We introduce a new knowledge representation model, which combines the information of those fuzzy measures with a crisp graph: the multi-dimensional extended vector fuzzy graph. We suggest a particular application of it, devoted to the community detection problem. To solve it, we define a method, based on the consideration of the multi-dimensional weighted graph associated with multiple fuzzy measures.

Index Terms—Fuzzy measure, Sugeno λ -measure, Community detection problem, Extended vector fuzzy graph, Weighted graph associated with a fuzzy measure

I. INTRODUCTION

The community detection problem is one of the most important topics in the field of Social Networks Analytics (SNA). Classical methods have their basis on the structure of the graph, an assumption which has provided good results. Nevertheless, in the last years, some authors have agreed on the importance of adding some additional information apart from the graph when dealing with SNA problems, particularly, when finding communities in a network. Different approaches can be found. Particularly, Gutiérrez et al. have been working in the incorporation of some additional information modeled by fuzzy measures to enrich the process [1]–[3].

We take up the idea introduced in [4] of finding communities by considering the additional information modeled by a vector, for example, the ratings of a film given by a group of people connected among them. Several synergies

among individuals may appear from the ratings, modeled by a fuzzy Sugeno λ -measure [5], [6], so the community structure may be affected. Draw from this initial assumption, now we suggest the management of multiple vectors. For example, imagine each vector represents the interest of a group of people regarding an activity, so that one vector is about sports, another about music or literature, and so on. From these interests, there may appear different synergies among the people, which will be represented by a multi-dimensional family of fuzzy Sugeno λ -measures, obtained from mentioned vectors. To handle these fuzzy Sugeno λ -measures, it is defined the multi-dimensional weighted graph associated with them (vector MAGW). It represents the synergies among individuals regarding the information given by the vectors. This tool has a many application in any SNA problem. Particularly, we will use it to address a community detection problem with additional information.

On the other hand, as a generalization of the extended vector fuzzy graph [4], we define the multi-dimensional extended vector fuzzy graph (MEVFG), which combines the knowledge of a crisp graph with that of multiple fuzzy Sugeno λ -measures obtained from vectors. On its basis, we approach a community detection problem based on fuzzy measures. We define an algorithm inspired by the Louvain method. It is based on modularity optimization [7] and local moving [8]. We also propose a specification of it to deal with 1-additive measures, situation in which the proposed algorithm is polynomial-time complexity.

The remainder of the paper is organized as follows. In Section II we provide several useful definitions which will be used later. Then, in Section III we address an scenario in which there are multiple fuzzy measures defining some information about a graph. The starting point is the existence of a family of vectors, each of which defines an evidence about the

This research has been partially supported by the Government of Spain, Grant Plan Nacional de I+D+i, MTM2015-70550-P, PGC2018096509-B-I00, BDNS Identif.: 498817 EB25/20 and CT17/17 - CT18/17.

individuals of the graph, i.e. the nodes. After that, we propose a particular application of the multi-dimensional extended vector fuzzy graph devoted to the community detection in Section IV. We conclude the work in Section V with some conclusions and future research lines.

II. PRELIMINARIES

In this section we show several definitions which set the basis of this paper. We start with the definition of a graph or network, a tool which set the basis of this work.

Definition 1. Graph / Network [9].

A graph or network is a pair $G = (V, E)$ where $V = \{1, 2, \dots, n\}$ is a set of individuals named nodes or vertices, and $E = \{\{i, j\} \mid i, j \in V\}$ is a non ordered set of pairs of nodes, named edges or arcs.

Another way to characterize a graph is by means of its adjacency matrix. This matrix, usually denoted by A , represents the direct connections between the nodes, in the sense that $A_{ij} = 1$ if $\{i, j\} \in E$, and $A_{ij} = 0$, otherwise.

On the other hand, we will consider non-weighted graphs, (and edge exists or not, so $A_{ij} = 1$ or $A_{ij} = 0$), and weighted graphs. In this case, there is a weight function defined on the set of edges, $w : E \rightarrow \mathbb{R}$, so there is a weight or value assigned to each edge. In this scenario, $A_{ij} = w_{ij}$, where w_{ij} is the weight of the edge $\{i, j\} \in E$.

Besides the graphs, the cornerstone of this paper is the use of fuzzy Sugeno λ -measures [4], functions to which we force to be fuzzy measures [6] and Sugeno λ -measures [5]. Particularly, we focus on this type of functions when they are obtained from a vector.

Definition 2. Fuzzy Sugeno λ -measure obtained from a vector $(\mu_{x,p})$ [4].

Let $x = (x_1, \dots, x_n)$ denote a vector defining any evidence about the elements of the n -set V , where $x_i \geq 0 \forall i$. Let $p \in (0, 1]$ denote a parameter. The function $\mu_{x,p}$ is fuzzy Sugeno λ -measure obtained from a vector, and it is characterized as follows:

$$\begin{aligned} \mu_{x,p}(i) &= \frac{px_i}{\sum_{k=1}^n x_k}, \forall i \in V \\ \text{and} \\ \mu_{x,p}(A \cup B) &= \mu_{x,p}(A) + \mu_{x,p}(B) + \lambda \mu_{x,p}(A) \mu_{x,p}(B), \\ &\forall A, B \subseteq V, \text{ with } A \cap B = \emptyset \text{ and} \\ &\lambda + 1 = \prod_{i=1}^n (1 + \lambda \mu_{x,p}(i)) \end{aligned}$$

To simplify the visualization and understanding of the relations defined by $\mu_{x,p}$, we will suggest the definition of the weighted graph associated with it. To characterize this graph, we consider the Shapley value. It is an essential tool in the

frame of Game Theory which was also adapted to the fuzzy measures background.

Definition 3. Shapley value [10].

Let $\mu : 2^V \rightarrow [0, 1]$ denote a fuzzy measure, where $|V| = n$. For every $i \in V$ its Shapley index is calculated as:

$$Sh_i(\mu) = \sum_{K \subseteq V \setminus \{i\}} \frac{(n - |K| - 1)! |K|!}{n!} (\mu(K \cup \{i\}) - \mu(K))$$

The Shapley value of the fuzzy measure $\mu : \mathcal{P}(V) \rightarrow [0, 1]$ is defined by the vector $Sh(\mu) = (Sh_1(\mu), \dots, Sh_n(\mu))$.

Hence, on the basis of the Shapely value and considering an aggregation operator, we recall the concept of weighted graph associated with a fuzzy measure [11], adapted for the scenario in which the fuzzy measure is obtained from a vector [4]. This graph is a tool which represents the synergies and relations between every pair of elements of the set V , according to the knowledge modeled by $\mu_{x,p}$.

Definition 4. Weighted graph associated with $\mu_{x,p}$, $G_{\mu_{x,p}}$ (AWG of vector) [4].

Let V denote a n -set, and let x denote a n -vector defining any evidence about the elements of V . Let $\mu_{x,p}$ denote the fuzzy Sugeno λ -measure obtained from x . The weighted graph associated with $\mu_{x,p}$, $G_{\mu_{x,p}}$, (AWG of vector) is that whose adjacency matrix is:

$$X_{ij} = \phi \left(Sh_i(\mu_{x,p}) - Sh_i^j(\mu_{x,p}), Sh_j(\mu_{x,p}) - Sh_j^i(\mu_{x,p}) \right) \quad (1)$$

being $\phi : [-1, 1]^2 \rightarrow [0, 1]$ a bi-variate aggregation operator [12]; $Sh_i(\mu_{x,p})$ and $Sh_i^j(\mu_{x,p})$ the Shapley values of i on μ_x when it is in a coalition with all the elements of V or $V \setminus \{j\}$, respectively [10].

For every pair of elements of V , the AWG represents how each individual is affected by the absence of the other in a coalition regarding $\mu_{x,p}$.

The calculation of the Shapely value may be hard for general fuzzy measures; nevertheless it is much easier for additive fuzzy measures. In [4] it was demonstrated that $\mu_{x,p}$ is a fuzzy Sugeno λ -measure. Particularly, when $p = 1$, it is 1-additive [13] (denoted by μ_x^a). To facilitate the calculation of the Shapley value when defining the weighted graph associated with $\mu_{x,p}$, we focus on this scenario.

To end this section, we present another tool which plays an essential role to address the community detection problem based on fuzzy measures is the extended fuzzy graph, firstly introduced in [11] and then adapted to the particular case of fuzzy Sugeno λ -measures in [4].

Definition 5. Extended vector fuzzy graph (EVFG) [4].

Let $G = (V, E)$ denote a crisp graph and let x denote a vector, so that $\mu_{x,p}$ is the fuzzy Sugeno λ -measure obtained from x . The triplet $\tilde{G}_x = (V, E, \mu_{x,p})$ is an extended vector fuzzy graph, (EVFG).

III. MULTIPLE FUZZY SUGENO λ - MEASURES IN A GRAPH

In the framework of networks analysis, we assume the existence of several information sources about the nodes of the graphs, given by a family of vectors, $(x^1, \dots, x^r)$, so that each of them defines (independently) an evidence about the individuals of a set. From these family of vectors, we characterize the fuzzy Sugeno λ -measures [4] $\mu_{x^1,p^1}, \dots, \mu_{x^r,p^r}$.

For a proper understanding and visualization of these fuzzy measures, we suggest the definition of the corresponding multi-dimensional weighted graph associated with them. It is a generalization to a multi-dimensional scale of the weighted graph associated with a fuzzy Sugeno λ -measure obtained from a vector (AWG of vector).

Definition 6. Multi-dimensional weighted graph associated with a family of fuzzy Sugeno λ -measures (vector MAWG).

Let $(x^1, \dots, x^r)$ denote a family of vectors; each one defines an evidence about the elements of a set V with $|V| = n$. $(\mu_{x^1,p^1}, \dots, \mu_{x^r,p^r})$ denote the fuzzy Sugeno λ -measures obtained from these vectors. The multi-dimensional weighted graph associated with $(\mu_{x^1,p^1}, \dots, \mu_{x^r,p^r})$ is that whose adjacency matrices are $(X^1, \dots, X^r)$, being X^ℓ the adjacency matrix of the AWG of μ_{x^ℓ,p^ℓ} , where, $\forall \ell = 1, \dots, r, \forall i, j \in V$,

$$X_{ij}^\ell = \phi^\ell \left(Sh_i(\mu_{x^\ell,p^\ell}) - Sh_i^j(\mu_{x^\ell,p^\ell}), Sh_j(\mu_{x^\ell,p^\ell}) - Sh_j^i(\mu_{x^\ell,p^\ell}) \right) \quad (2)$$

being $\phi : [-1, 1]^2 \rightarrow [0, 1]$ a bi-variate aggregation operator [12]; $Sh_i(\mu_{x,p})$ and $Sh_i^j(\mu_{x,p})$ the Shapley indices of i on μ_x when it is in a coalition with all the elements of V or $V \setminus \{j\}$, respectively [10].

Then, we define a knowledge representation model which combines the information provided by a crisp graph with some additional information independent of its structure which is modeled by some vectors.

Definition 7. Multi-dimensional extended vector fuzzy graph (MEVFG).

Let $G = (V, E)$ denote a crisp graph, and let $(x^1, \dots, x^r)$ denote a family of vectors so that each of them defines any evidence about the elements of V . Let $(\mu_{x^1,p^1}, \dots, \mu_{x^r,p^r})$ denote the family of fuzzy Sugeno λ -measures obtained from vectors $(x^1, \dots, x^r)$. Then, $\tilde{G} = (V, E, (\mu_{x^1,p^1}, \dots, \mu_{x^r,p^r}))$ is a multi-dimensional extended vector fuzzy graph, (MEVFG).

Let us note that the MEVFG goes further than other existent tool: it allows the characterization of several synergies among the individuals, regardless their connections in the crisp graph. Furthermore, due to the properties of the fuzzy Sugeno λ -measures, some relations between elements can be inferred from the knowledge about some individual evidence.

IV. COMMUNITY DETECTION PROBLEM IN A MULTI-DIMENSIONAL EXTENDED VECTOR FUZZY GRAPH

The MEVFG has numerous applications. In this paper we take up the idea introduced in [4] related to community detection in graphs according to the existence of a fuzzy Sugeno λ -measure. Now we work with the multi-dimensional extended vector fuzzy graph $\tilde{G} = (V, E, (\mu_{x^1,p^1}, \dots, \mu_{x^r,p^r}))$ obtained from the combination of a crisp graph $G = (V, E)$ and a family of vectors $(x^1, \dots, x^r)$ defining additional information about the elements of V , from which we define the fuzzy Sugeno λ -measures $(\mu_{x^1,p^1}, \dots, \mu_{x^r,p^r})$. We agree that, the more information is analyzed, the more cohesive and realistic are the groups detected.

The proposed methodology, named *Multi-dimensional Sugeno Louvain*, is based on the Louvain Algorithm [14]. The main point of our algorithm is to summarize all the knowledge of the MEVFG into two matrices: the adjacency matrix of the crisp graph, A , represents the direct connections between the nodes (edges), and $\mathbb{X}$ summarizes the additional information given by the family of vectors $(x^1, \dots, x^r)$.

- **Step 1: definition of the vector MAWG.** Given the fuzzy Sugeno λ -measures $(\mu_{x^1,p^1}, \dots, \mu_{x^r,p^r})$ obtained from $(x^1, \dots, x^r)$, matrices $(X^1, \dots, X^r)$ have to be defined according to equation (2).
- **Step 2: information aggregation.** Matrices $X^1, \dots, X^r$ are aggregated to obtain the matrix $\mathbb{X}$. The aggregation function $\Phi : \Pi^r \rightarrow \Pi$ is used, being Π the set of quadratic n -matrices. Particularly, we suggest the use of a matrix aggregator based on the classical aggregation operators with “element to element” transformation: $\mathbb{X} = \Phi(X^1, \dots, X^r)$.

After this aggregation process, the method *Duo Louvain*, summarized by its pseudo-code in the Algorithm 1, has to be applied (see [15], [16] for more details), considering the matrix $M = \theta(A, \mathbb{X})$, being $\theta : \Pi^2 \rightarrow \Pi$ an aggregation function. That method can consider the information of two matrices when finding communities in a graph (A is used to find “feasible” communities, and any other matrix to calculate the maximum of modularity). Note that the notion of group and its size depend on the operator Φ considered [16]. The new community detection method defined to find groups in a multi-dimensional extended vector fuzzy graph is summarized in the Algorithm 2 and it is named *Multi-dimensional Sugeno-Louvain*.

Algorithm 1 *Duo Louvain*

```

1: Input:  $(A, M)$ ;
2: Output:  $P$ ;
3: Preliminary
4:  $C_i \leftarrow \{i\}, \forall i \in V$  (each node  $i$  is an isolated community);
5:  $P \leftarrow (1, 2, \dots, n)$  (initial partition);
6: end Preliminary
7: Phase 1
8: Take  $o = (o^1, \dots, o^i, \dots, o^n) \in \pi(V)$ ;
9:  $stop \leftarrow 0$ ;
10: while ( $stop == 0$ ) do
11:    $stop \leftarrow 1$ 
12:   for ( $i = 1$ ) to ( $n$ ) do
13:      $(e_1, \dots, e_h) \leftarrow H(o^i)$  (find the neighbours of  $o^i$  in  $A$ );
14:     for ( $j = 1$ ) to ( $h$ ) do
15:       Calculate  $\Delta Q_{o^i}(e_j)$  in  $M$ ;
16:     end for
17:      $j^* \leftarrow \left\{ e_\ell \mid \Delta Q_{o^i}(j^*) = \max_{\ell \in \{1, \dots, h\}} \{\Delta Q_{o^i}(e_\ell)\} \right\}$ ;
18:     if ( $\Delta Q_{o^i}(j^*) > 0$ ) then
19:        $C_{P(o^i)} \leftarrow C_{P(o^i)} \setminus \{o^i\}$ ;
20:        $C_{P(j^*)} \leftarrow C_{P(j^*)} \cup \{o^i\}$ ;
21:        $P(o^i) \leftarrow P(j^*)$ ;
22:      $stop \leftarrow 0$ ;
23:   end if
24: end for
25: end while
26: end Phase 1
27: Phase 2
28: Calculate  $A^*$  from  $A$  (nodes of  $A^*$  are the communities previously found in  $A$ );
29: Calculate  $M^*$  from  $M$  (nodes of  $M^*$  are the communities previously found in  $M$ );
30: if ( $A^* \neq A$ ) then
31:    $A \leftarrow A^*$ ;
32:    $M \leftarrow M^*$ ;
33:   Apply Phase 1 and Phase 2;
34: end if
35: end Phase 2
36: return( $P$ );

```

Algorithm 2 *Multi-dimensional Sugeno-Louvain*

```

1: Input:  $(A, (x^1, \dots, x^r), (p^1, \dots, p^r))$ ,  $A$  represents  $G = (V, E)$ ;
2: Output:  $P$ ;
3: Preliminary
4: for ( $\ell = 1$ ) to ( $r$ ) do
5:   Calculate  $\mu_{x^\ell, p^\ell}$  (fuzzy Sugeno  $\lambda$ -measure from  $x^\ell$ );
6:    $X_{ij}^\ell \leftarrow \phi^\ell (Sh_i(\mu_{x^\ell, p^\ell}) - Sh_i^j(\mu_{x^\ell, p^\ell}) Sh_j(\mu_{x^\ell, p^\ell}) - Sh_j^i(\mu_{x^\ell, p^\ell})),$ 
      $\forall i, j \in V$ ;
7: end for
8:  $\mathbb{X} \leftarrow \Phi(X^1, \dots, X^r)$ ;
9:  $M \leftarrow \theta(A, \mathbb{X})$ ;
10: end Preliminary
11:  $P \leftarrow DuoLouvain(A, M)$ ;
12: return( $P$ );

```

The exponential complexity concerning the Shapley value may be avoided by considering additive fuzzy measure. Then, we suggest a specific application of the *Multi-dimensional Sugeno-Louvain* Algorithm, named *1-additive Multi-dimensional Sugeno Louvain*, which involves a particular 1-additive characterization of μ_{x^ℓ, p^ℓ} , denoted by $\mu_{x^\ell}^a$. On this basis, the calculation of the Shapley index is immediate from vector x^ℓ as follows: $Sh_i(\mu_{x^\ell}^a) = \frac{x_i^\ell}{\sum_{k=1}^n x_k^\ell}$ and $Sh_i^j(\mu_{x^\ell}^a) = \frac{x_i^\ell}{\sum_{k=1, k \neq j}^n x_k^\ell}$. Then, the complexity of the method *1-additive Multi-dimensional Sugeno-Louvain* (its pseudo-code is showed in Algorithm 3), whose only difference with respect to the *Multidimensional-dimensional Sugeno-Louvain* is the calculation of X^ℓ (line 6) in *pseudo-code* is equal to the Louvain Algorithm.

Algorithm 3 *1-additive Multi-dimensional Sugeno Louvain*

```

1: Input:  $(A, (x^1, \dots, x^r))$ ,  $A$  is a representation of  $G = (V, E)$ ;
2: Output:  $P$ ;
3: Preliminary
4: for ( $\ell = 1$ ) to ( $r$ ) do
5:    $X_{ij}^\ell \leftarrow \phi\{|\frac{x_i^\ell}{\sum_{k=1}^n x_k^\ell} - \frac{x_i^\ell}{\sum_{k=1, k \neq j}^n x_k^\ell}|, |\frac{x_j^\ell}{\sum_{k=1}^n x_k^\ell} - \frac{x_j^\ell}{\sum_{k=1, k \neq i}^n x_k^\ell}|\}$ ;
6: end for
7:  $\mathbb{X} \leftarrow \Phi(X^1, \dots, X^r)$ ;
8:  $M \leftarrow \theta(A, \mathbb{X})$ ;
9: end Preliminary
10:  $P \leftarrow DuoLouvain(A, M)$ ;
11: return( $P$ );

```

Example 1.

We consider graph $G = (V, E)$ whose nodes are connected as a chain of size 12, as it can be seen in matrix A .

$$A = \begin{pmatrix} 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \end{pmatrix}$$

Fig. 1. Adjacency matrix of graph $G = (V, E)$

There exist 4 vectors of additional information, defining some evidences about the elements of V ,

- $x^1 = (9, 9.5, 10, 1, 0.5, 1, 9.5, 8, 10, 1, 1.5, 1)$
- $x^2 = (10, 9.5, 9, 1, 0.5, 1, 9, 9, 9.5, 1.5, 2, 0.5)$
- $x^3 = (9.5, 8.5, 10, 1.5, 1, 1, 10, 9.5, 9.5, 0.9, 1, 1)$
- $x^4 = (9, 9.5, 10, 1, 1, 1, 10, 9.5, 9, 0.5, 1, 1)$

Each vector represents the opinion (ratings) of 12 people about a films.

From vectors (x^1, x^2, x^3, x^4) , we define the family of 1-additive fuzzy Sugeno λ -measures $(\mu_{x^1}^a, \mu_{x^2}^a, \mu_{x^3}^a, \mu_{x^4}^a)$. The adjacency matrices of the corresponding vector MAWG are (X^1, X^2, X^3, X^4) , showed in the Figure 2.

We accept that there are more synergies between those people who have similar preferences. Any classical community detection algorithm based on modularity optimization provides the partition $P = \{\{1, 2, 3, 4\}, \{5, 6, 7, 8\}, \{9, 10, 11, 12\}\}$. Nevertheless, if the additional information is considered, the partition provided by the Multi-dimensional Sugeno-Louvain 1-additive is $P^x = \{\{1, 2, 3\}, \{4, 5, 6\}, \{7, 8, 9\}, \{10, 11, 12\}\}$, which respect the synergies between the elements (showed in matrices X^1, X^2, X^3, X^4 , whose highest values are bold), as well as the structure of the crisp graph established by the edges.

$$X^1 = \frac{1}{10000} \begin{pmatrix} 0 & \mathbf{260} & \mathbf{274} & 24 & 12 & 24 & \mathbf{260} & \mathbf{215} & \mathbf{274} & 24 & 36 & 24 \\ \mathbf{260} & 0 & \mathbf{292} & 25 & 12 & 25 & \mathbf{277} & \mathbf{227} & \mathbf{292} & 25 & 38 & 25 \\ \mathbf{274} & \mathbf{292} & 0 & 26 & 13 & 26 & \mathbf{292} & \mathbf{239} & \mathbf{310} & 26 & 40 & 26 \\ 24 & 25 & 26 & 0 & 1 & 3 & 25 & 21 & 26 & 3 & 4 & 3 \\ 12 & 12 & 13 & 1 & 0 & 1 & 12 & 10 & 13 & 1 & 2 & 1 \\ 24 & 25 & 26 & 3 & 1 & 0 & 25 & 21 & 26 & 3 & 4 & 3 \\ \mathbf{260} & \mathbf{277} & \mathbf{292} & 25 & 12 & 25 & 0 & \mathbf{227} & \mathbf{292} & 25 & 38 & 25 \\ \mathbf{215} & \mathbf{227} & \mathbf{239} & 21 & 10 & 21 & \mathbf{227} & 0 & \mathbf{239} & 21 & 32 & 21 \\ \mathbf{274} & \mathbf{292} & \mathbf{310} & 26 & 13 & 26 & \mathbf{292} & \mathbf{239} & 0 & 26 & 40 & 26 \\ 24 & 25 & 26 & 3 & 1 & 3 & 25 & 21 & 26 & 3 & 4 & 3 \\ 36 & 38 & 40 & 4 & 2 & 4 & 38 & 32 & 40 & 4 & 0 & 4 \\ 24 & 25 & 26 & 3 & 1 & 3 & 25 & 21 & 26 & 3 & 4 & 0 \end{pmatrix}$$

$$X^2 = \frac{1}{10000} \begin{pmatrix} 0 & \mathbf{287} & \mathbf{269} & 26 & 13 & 26 & \mathbf{269} & \mathbf{269} & \mathbf{287} & 39 & 53 & 13 \\ \mathbf{287} & 0 & \mathbf{256} & 25 & 12 & 25 & \mathbf{256} & \mathbf{256} & \mathbf{272} & 37 & 50 & 12 \\ \mathbf{269} & \mathbf{256} & 0 & 23 & 12 & 23 & \mathbf{242} & \mathbf{242} & \mathbf{256} & 35 & 48 & 12 \\ 26 & 25 & 23 & 0 & 1 & 3 & 23 & 23 & 25 & 4 & 5 & 1 \\ 13 & 12 & 12 & 1 & 0 & 1 & 12 & 12 & 12 & 2 & 3 & 1 \\ 26 & 25 & 23 & 3 & 1 & 0 & 23 & 23 & 25 & 4 & 5 & 1 \\ \mathbf{269} & \mathbf{256} & \mathbf{242} & 23 & 12 & 23 & 0 & \mathbf{242} & \mathbf{256} & 35 & 48 & 12 \\ \mathbf{269} & \mathbf{256} & \mathbf{242} & 23 & 12 & 23 & \mathbf{242} & 0 & \mathbf{256} & 35 & 48 & 12 \\ \mathbf{287} & \mathbf{272} & \mathbf{256} & 25 & 12 & 25 & \mathbf{256} & \mathbf{256} & 0 & 37 & 50 & 12 \\ 39 & 37 & 35 & 4 & 2 & 4 & 35 & 35 & 37 & 0 & 8 & 2 \\ 53 & 50 & 48 & 5 & 3 & 5 & 48 & 48 & 50 & 8 & 0 & 3 \\ 13 & 12 & 12 & 1 & 1 & 1 & 12 & 12 & 12 & 2 & 3 & 0 \end{pmatrix}$$

$$X^3 = \frac{1}{10000} \begin{pmatrix} 0 & \mathbf{232} & \mathbf{278} & 36 & 24 & 24 & \mathbf{278} & \mathbf{264} & \mathbf{264} & 22 & 24 & 24 \\ \mathbf{232} & 0 & \mathbf{244} & 32 & 21 & 21 & \mathbf{244} & \mathbf{232} & \mathbf{232} & 19 & 21 & 21 \\ \mathbf{278} & \mathbf{244} & 0 & 38 & 25 & 25 & \mathbf{295} & \mathbf{278} & \mathbf{278} & 23 & 25 & 25 \\ 36 & 32 & 38 & 0 & 4 & 4 & 38 & 36 & 36 & 3 & 4 & 4 \\ 24 & 21 & 25 & 4 & 0 & 3 & 25 & 24 & 24 & 2 & 3 & 3 \\ 24 & 21 & 25 & 4 & 3 & 0 &quad 25 & 24 & 24 & 2 & 3 & 3 \\ \mathbf{278} & \mathbf{244} & \mathbf{295} & 38 & 25 & 25 & 0 & \mathbf{278} & \mathbf{278} & 23 & 25 & 25 \\ \mathbf{264} & \mathbf{232} & \mathbf{278} & 36 & 24 & 24 & \mathbf{278} & 0 & \mathbf{264} & 22 & 24 & 24 \\ \mathbf{264} & \mathbf{232} & \mathbf{278} & 36 & 24 & 24 & \mathbf{278} & \mathbf{264} & 0 & 22 & 24 & 24 \\ 22 & 19 & 23 & 3 & 2 & 2 & 23 & 22 & 22 & 0 & 2 & 2 \\ 24 & 21 & 25 & 4 & 3 & 3 & 25 & 24 & 24 & 2 & 0 & 3 \\ 24 & 21 & 25 & 4 & 3 & 3 & 25 & 24 & 24 & 2 & 3 & 0 \end{pmatrix}$$

$$X^4 = \frac{1}{10000} \begin{pmatrix} 0 & \mathbf{256} & \mathbf{269} & 23 & 23 & 23 & \mathbf{269} & \mathbf{256} & \mathbf{242} & 12 & 23 & 23 \\ \mathbf{256} & 0 & \mathbf{287} & 25 & 25 & 25 & \mathbf{287} & \mathbf{272} & \mathbf{256} & 12 & 25 & 25 \\ \mathbf{269} & \mathbf{287} & 0 & 26 & 26 & 26 & \mathbf{305} & \mathbf{287} & \mathbf{269} & 13 & 26 & 26 \\ 23 & 25 & 26 & 0 & 3 & 3 & 26 & 25 & 23 & 1 & 3 & 3 \\ 23 & 25 & 26 & 3 & 0 & 3 & 26 & 25 & 23 & 1 & 3 & 3 \\ 23 & 25 & 26 & 3 & 3 & 0 & 26 & 25 & 23 & 1 & 3 & 3 \\ \mathbf{269} & \mathbf{287} & \mathbf{305} & 26 & 26 & 26 & 0 & \mathbf{287} & \mathbf{269} & 13 & 26 & 26 \\ \mathbf{256} & \mathbf{272} & \mathbf{287} & 25 & 25 & 25 & \mathbf{287} & 0 & \mathbf{256} & 12 & 25 & 25 \\ \mathbf{242} & \mathbf{256} & \mathbf{269} & 23 & 23 & 23 & \mathbf{269} & \mathbf{256} & 0 & 12 & 23 & 23 \\ 12 & 12 & 13 & 1 & 1 & 1 & 13 & 12 & 12 & 0 & 1 & 1 \\ 23 & 25 & 26 & 3 & 3 & 3 & 26 & 25 & 23 & 1 & 0 & 3 \\ 23 & 25 & 26 & 3 & 3 & 3 & 26 & 25 & 23 & 1 & 3 & 0 \end{pmatrix}$$

Fig. 2. Adjacency of the vector MAGW of $(\mu_{x^1}^a, \mu_{x^2}^a, \mu_{x^3}^a, \mu_{x^4}^a)$

V. CONCLUSIONS

In this paper we address a new perception of the community detection problem in networks. In a multi-dimensional scale, we suggest the inclusion of some additional information defined by a family of vectors to the process of finding groups in a crisp graph. We assert that, the more information is analyzed, the more realistic the results obtained.

In the particular context of fuzzy Sugeno λ -measures, we introduce two new tools which facilitates the handling of a family of them. The first one is the multi-dimensional weighted graph associated with multiple fuzzy Sugeno λ -measures, a representation tool which shows the synergies between every

pair of elements regarding the knowledge modeled by the measures considered.

On the other hand, for an scenario modeled by a crisp graph $G = (V, E)$, about whose nodes there is some additional information defined by vectors $(x^1, \dots, x^r)$, we define the multi-dimensional extended vector fuzzy graph $\tilde{G} = (V, E, (\mu_{x^1, p^1}, \dots, \mu_{x^r, p^r}))$, being $(\mu_{x^1, p^1}, \dots, \mu_{x^r, p^r})$ the fuzzy Sugeno λ -measures obtained from those vectors [4].

On the basis of this tool, we approach the community detection problem based on multiple fuzzy measures. We define a new algorithm, *Multi-dimensional Sugeno Louvain*, which has two main points: (1) an aggregation process to summarize all the knowledge modeled by $\tilde{G}$ into two matrices; (2) the application of the *Duo Louvain* Algorithm. We suggest a particular application of that algorithm to use 1-additive fuzzy measures, so that the calculation of the Shapley value is immediate. So, we can affirm that the complexity of the *1-additive Multi-dimensional Sugeno Louvain* Algorithm is equal to the Louvain Algorithm.

At the moment, we are currently working on the evaluation and testing of the proposed methodology. To carry on with it, we will differentiate two different steps. The first one, will be devoted to the consideration of synthetic networks. We will apply our algorithm in several benchmark models [17], then we will consider the Normalized Mutual Information (NMI) [18] to evaluate the results obtained. After that, we will work with some real cases. We are particularly interested in analyzing the data obtained from Social Networks, such as Twitter or Facebook.

The testing process is still in early stage, but the preliminary results we have obtained are very promising on the goodness of the proposed method.

REFERENCES

- [1] M. Barroso, I. Gutiérrez, D. Gómez, J. Castro, and R. Espínola, "Group definition based on flow in community detection," in *Information Processing and Management of Uncertainty in Knowledge-Based Systems*. Springer, 2020, pp. 524–538.
- [2] I. Gutiérrez, J. A. Guevara, D. Gómez, J. Castro, and R. Espínola, "Community detection problem based on polarization measures. an application to Twitter: the COVID-19 case in Spain," vol. 9, no. 4, 2021.
- [3] I. Gutiérrez, M. Barroso, D. Gómez, J. Castro, and R. Espínola, "Pattern-based clustering problem based on fuzzy measures," *Developments of Artificial Intelligence Technologies in Computation and Robotics*, vol. 12, pp. 412–420, 2020.
- [4] I. Gutiérrez, D. Gómez, J. Castro, and R. Espínola, "Fuzzy sugeno λ -measures and theirs applications to community detection problems." Glasgow, UK: IEEE International Conference on Fuzzy Systems, 2020, pp. 1–6.
- [5] M. Sugeno, "Theory of fuzzy integrals and its applications," Ph.D. dissertation, Tokyo Institute of Technology, 1974.
- [6] —, "Fuzzy measures and fuzzy integrals: A survey," *Fuzzy Automata and Decision Processes*, vol. 78, pp. 89–102, 1977.
- [7] M. Newman and M. Girvan, "Finding and evaluating community structure in networks," *Physical Review E*, vol. 69, no. 026113, 2004.
- [8] L. Waltman and N. J. Van Eck, "A smart local moving algorithm for large-scale modularity-based community detection," *Europ. Phys. Journal B*, vol. 86, no. 11, 2013.
- [9] J. Sylvester, "Chemistry and algebra," *Nature* 17, vol. 3, no. 284, pp. 103–104, 1878.
- [10] L. S. Shapley, "A value for n -person games," *Contribute. Theory Games*, vol. 2, no. 28, pp. 307–317, 1953.
- [11] I. Gutiérrez, D. Gómez, J. Castro, and R. Espínola, "A new community detection algorithm based on fuzzy measures," in *Advances in Intelligent Systems and Computing series*, vol. 1029, Intelligent and Fuzzy Techniques in Big Data Analytics and Decision Making. INFUS 2019. Springer, Cham, 2020, pp. 133–140.
- [12] R. R. Yager and J. Kacprzyk, Eds., *The Ordered Weighted Averaging Operators. Theory and Applications*. Boston, MA, USA: Springer Science & Business Media, 1997.
- [13] M. Grabisch, " k -order additive discrete fuzzy measures and their representation," *Fuzzy Sets Systems*, vol. 92, no. 2, pp. 167–189, 1997.
- [14] V. Blondel, J. Guillaume, R. Lambiotte, and E. Lefevre, "Fast unfolding of communities in large networks," *Journal Of Statistical Mechanics-Theory And Experiment*, vol. 10, no. P10008, 2008.
- [15] I. Gutiérrez, D. Gómez, J. Castro, and R. Espínola, "Fuzzy measures: A solution to deal with community detection problems for networks with additional information," *JIFS*, vol. 39, no. 5, pp. 6217–6230, 2020.
- [16] —, "Multiple bipolar fuzzy measures: an application to community detection problems for networks with additional information," *IJCIS*, vol. 13, no. 1, pp. 1636–1649, 2020.
- [17] D. Bader, A. Kappes, H. Meyerhenke, P. Sanders, C. Schulz, and D. Wagner.
- [18] T. O. Kvalseth, "On normalized mutual information: measure derivations and properties," *Entropy*, vol. 19, no. 12, p. 631, 2017.

Explicaciones factuales y contrafactuales en árboles de clasificación difusos

Guillermo Tomás Fernández, Juan A. Aledo, Jose A. Gámez y Jose M. Puerta

Laboratorio de Sistemas Inteligentes y Minería de Datos (SIMD), I3A

Universidad de Castilla-La Mancha. Albacete, España

{Guillermo.Fernandez, JuanAngel.Aledo, Jose.Gamez, Jose.Puerta}@uclm.es

Resumen—Los algoritmos de clasificación actuales han alcanzado una gran popularidad debido a su eficiencia para generar modelos capaces de resolver problemas de alta complejidad. En particular, son los algoritmos denominados de caja negra los que mejores resultados ofrecen, ya que se benefician de esta enorme cantidad de datos para aprender modelos cada vez más precisos. Sin embargo, su principal desventaja frente a otros algoritmos más simples, p.e. un árbol de decisión, es la pérdida de interpretación tanto del modelo como de las clasificaciones individuales, lo que supone un grave inconveniente de cara a muchas aplicaciones en las que proveer una explicación es hoy día recomendable, e incluso obligatorio. Una práctica habitual es construir un modelo *explicable* que mimetice el comportamiento del clasificador más complejo en la zona circundante a la instancia a explicar.

Sin embargo, la generación de explicaciones en estos modelos de caja blanca tampoco es trivial, lo que ha generado una intensa investigación en torno a ellos. Es habitual generar dos tipos de explicaciones, factuales y contrafactuales, que se complementan para informar al decisor por qué se ha clasificado la instancia en una determinada clase o categoría. En este trabajo proponemos la definición de explicaciones factuales y contrafactuales en el marco de los árboles de clasificación difusos, en los que al contrario de su contraparte *crisp* una instancia puede disparar más de una rama. Nuestra propuesta se centra en definir explicaciones factuales que contienen más de una regla, en contraposición al estándar habitual que se limita a incluir una única regla en la explicación factual. Además, introducimos la idea de explicación factual *robusta* y la generación de explicaciones contrafactuales a partir de la clasificación realizada y la explicación factual generada, que puede tener más de una regla.

Index Terms—Inteligencia artificial explicable (XAI); Lógica difusa; Árboles de decisión difusos; Explicaciones factuales; Explicaciones contrafactuales; Robustez.

I. INTRODUCCIÓN

La gran capacidad de decisión de los algoritmos modernos de clasificación ha originado un enorme incremento en la variedad de campos en los que se están aplicando dichos algoritmos. Problemas que hasta recientemente necesitaban obligatoriamente la intervención de un experto (o un sistema experto, con la complejidad y especificidad que este requiere) se están resolviendo mediante la aplicación de técnicas que se aprovechan de la gran cantidad de datos que se pueden recoger para aprender de manera automática un clasificador capaz de resolver el problema.

Sin embargo, a pesar de que estos algoritmos están alcanzando cotas de precisión cada vez más altas, lo hacen a costa de la interpretabilidad final para el usuario. El razonamiento

que realizan estos sistemas es cada vez más complejo, lo que crea una necesidad de tener una creencia ciega en estos sistemas. En ciertos ámbitos críticos, esta creencia ciega no es suficiente para motivar el uso de este tipo de clasificadores. Este hecho se ve también refrendado por la legislación, p.e. el *derecho a la explicación*, incluido en la Regulación General de Protección de Datos aprobada por la Unión Europea [1], que no solo afecta a humanos sino también a técnicas de inteligencia artificial y sistemas informáticos.

Para atender las necesidades comentadas, surge la Inteligencia Artificial Explicable (XAI), una línea de investigación en auge que se centra en explicar aquellos modelos y sistemas que por sus características no resultan interpretables por un usuario. Dado que normalmente este tipo de algoritmos, denominados de *caja negra*, son los que mejores resultados obtienen en los distintos problemas, es especialmente interesante desarrollar sistemas capaces de explicar sus decisiones. En este sentido, los métodos *agnósticos* [2]–[4] explican la decisión realizada por un modelo complejo (habitualmente de caja negra) mediante la construcción de un modelo más sencillo de explicar que mimetiza al modelo complejo en la vecindad de la instancia cuya clasificación debe ser explicada.

Sin embargo, incluso cuando se construyen modelos de caja blanca, la generación de la explicación no es trivial, si no que es habitual generar explicaciones de distinto tipo: que indican por qué se ha clasificado en una determinada categoría (explicaciones factuales) y que indican por qué no ha sido clasificada en otra categoría (explicaciones contrafactuales) [5]. Existen distintas definiciones para este tipo de explicaciones en función del formalismo utilizado: árboles de decisión [4], [6], regresión [2], clasificadores probabilísticos [7], [8], etc.

En este trabajo nos centramos en el uso de clasificadores basados en árboles de clasificación difusos [9], en los que al contrario de sus homónimos *crisp*, una instancia puede disparar más de una regla. Parece razonable que si p.e. dos reglas difusas (extraídas del árbol) se disparan con similar grado de importancia (activación), ambas puedan ser usadas para explicar la clasificación, en lugar de seleccionar únicamente una de ellas. Partiendo de esta premisa nuestra contribución en este artículo se centra en la definición de explicaciones factuales que puedan contener, en caso de necesidad, más de una regla difusa. A partir de estas definiciones planteamos el concepto de *robustez* de una explicación factual y la generación de explicaciones contrafactuales a partir de la instancia a clasificar y la explicación factual (posiblemente no unitaria)

generada.

II. DEFINICIÓN DEL PROBLEMA

Consideremos un problema de clasificación supervisada, que consiste en asignar una clase c_i perteneciente a un conjunto predefinido $C = \{c_1, \dots, c_m\}$ a una instancia x . Sea esta instancia $x = (x_1, \dots, x_n)$ una configuración de valores sobre n variables predictoras, $X_1, \dots, X_n$.

Asumamos que asociada a cada variable predictora X_i hay una variable difusa (lingüística) $F_i = \{v_{i,1}, \dots, v_{i,k_i}\}$ definida mediante una partición de Ruspini de k_i conjuntos difusos ordenados. Utilizaremos $v_{i,j}$ para referirnos indistintamente tanto al conjunto difuso como a la etiqueta lingüística asociada al mismo. Dado un valor $\delta \in \text{dom}(X_i)$, sea

$$\mu_i(\delta) = (\mu_{i,1}(\delta), \dots, \mu_{i,k_i}(\delta))$$

el vector de grados de pertenencia de δ a los k_i conjuntos difusos de F_i . Nótese que $\sum_{j=1}^{k_i} \mu_{i,j}(\delta) = 1$ (Figura 1). Definimos

$$f(\delta) = \arg \max_{1, \dots, k_i} \mu_i(\delta),$$

como el índice correspondiente al conjunto difuso al que δ tiene mayor grado de pertenencia. En concreto, $\mu_{i,z_i}(\delta)$ es el grado de pertenencia del valor δ en F_i para el conjunto v_{i,z_i} .

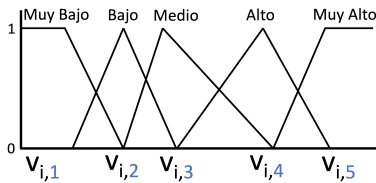


Figura 1. Ejemplo de partición de Ruspini con 5 conjuntos difusos para la variable i -ésima. Se muestran los conjuntos difusos $v_{i,j}$ y las etiquetas lingüísticas.

Por último, sea T un árbol de decisión difuso aprendido a partir de un conjunto de datos de entrenamiento $TR = \{(x_1^j, \dots, x_n^j, c^j)\}_{j=1}^N$, donde N es el número de instancias y $c^j \in C$ es la clasificación (categoría o clase) asociada a cada instancia.

Nuestro objetivo es obtener una explicación $e = \langle R_f, R_{cf} \rangle$ para la instancia x a partir de las reglas¹ que se extraen del árbol T , donde R_f es una explicación factual y R_{cf} es una explicación contrafactual.

III. ESTADO DEL ARTE

III-A. Explicaciones factuales y contrafactuales para clasificadores difusos

En [6] se propone un método para generar explicaciones factuales y contrafactuales a partir de árboles de decisión tanto crisp como difusos. En ambos casos la explicación factual contiene una única regla y se genera una explicación contrafactual por cada categoría distinta a la predicha como clasificación.

¹De forma más general se podría extender a un conjunto de reglas difusas.

A la hora de extraer los factuales, los autores consideran el árbol como una serie de nodos unidos por aristas pesadas. En un árbol crisp, la arista tendrá un peso binario, mientras que en un árbol difuso tendrá un valor real entre 0 y 1 (según el grado de activación de la instancia con el nodo). Estos nodos y aristas generan caminos desde la raíz hasta las hojas, cada uno con un peso asociado según el valor de las aristas que lo compongan. Finalmente, para obtener el factual los autores utilizan un α -corte en los caminos generados, de tal manera que todos los caminos con un peso menor que α se eliminen. Así, garantizan que sólo existe un único factual en un árbol difuso, de igual manera que en el caso crisp existe un único camino desde la raíz hasta la hoja que clasifica un ejemplo.

En cuanto a las explicaciones contrafactuales, existen $m - 1$ explicaciones donde m es el número de clases. En particular, una explicación contrafactual es un camino desde la raíz del árbol hasta una hoja que tenga una clase distinta a la de la instancia a explicar.

III-B. Árboles de Decisión Difusos

En [9] se describe un método que primero genera de manera automática los conjuntos difusos que se utilizarán y posteriormente el proceso de aprendizaje del árbol usando dichos conjuntos. En este trabajo consideramos los árboles producidos por este algoritmo como clasificadores cuya inferencia se pretende explicar.

Dado un árbol T ya entrenado, cada rama tiene la forma

$$b : l_{i_1}[b] \wedge l_{i_2}[b] \wedge \dots \wedge l_{i_b}[b] \wedge h[b]$$

donde

- i_b es el número de literales en el antecedente de la regla que se deriva de b .
- $l_i[b] : (F_i, v_{i,z_i})$ es un literal compuesto por la variable difusa F_i y la etiqueta v_{i,z_i} , con $1 \leq z_i \leq k_i$.
- $h[b] : \{c_1 = w_1[b], c_2 = w_2[b], \dots, c_m = w_m[b]\}$ es un nodo hoja, donde se registra la importancia $w_i[b]$ para cada clase en esta hoja, con $\sum_{i=1}^m w_i[b] = 1$.

De cada rama b y cada clase c_i con $w_i[b] > 0$ obtenemos una regla

$$r_i[b] : l_{i_1}[b] \wedge l_{i_2}[b] \wedge \dots \wedge l_{i_b}[b] \rightarrow c_i,$$

con peso asociado $w(r_i[b]) = w_i[b]$.

Denotaremos como $R(c_i)$ el conjunto de todas las reglas que se pueden extraer del árbol con consecuente c_i .

Llamamos **regla de peso máximo** para la rama b a

$$r[b] = \arg \max_{r_i[b]} w_i[b].$$

En caso de empate se elige una única regla de manera arbitraria o aleatoria.

Al realizar la inferencia, una instancia x dispara $p(x) \geq 1$ reglas (posiblemente de distintas ramas). Denotemos por $R(x) = \{r_1(x), \dots, r_{p(x)}(x)\}$ el conjunto de reglas disparadas.

Se define el grado de emparejamiento de una regla $r(x) \in R(x)$ compuesta por i_b literales $l_i : (F_i, v_{i,z_i})$ con la instancia $x = (x_1, \dots, x_n)$ como

$$md(r(x)) = \min_{i \in \{i_1, \dots, i_b\}} (\mu_{i, z_i}(x_i))$$

Por otra parte, se define el grado de activación $AD(r(x))$ de una instancia x para la regla $r(x) \in R(x)$ como

$$AD(r(x)) = md(r(x)) \cdot w(r(x))$$

Dado un conjunto de reglas $\Gamma(x) \subseteq R(x)$, obtenemos la clase asociada a $\Gamma(x)$ mediante el método *weighted vote* como

$$c(\Gamma(x)) = \arg \max_{c_i} \sum_{\substack{r(x) \in \Gamma(x) \\ c(r(x)) = c_i}} AD(r(x)), \quad (1)$$

con $c(r(x))$ la clase asociada a la regla $r(x)$. En caso de empate, este se rompe aleatoria o arbitrariamente. En particular, la clase $c(x)$ predicha por el conjunto de reglas $R(x)$ se obtiene como $c(x) = c(R(x))$.

IV. METODOLOGÍA

En este trabajo, proponemos una definición de explicación factual que expande las encontradas hasta ahora en la literatura. Si bien las explicaciones factuales contienen una única regla y esto es razonable en entornos *crisp*, en el caso difuso es posible disparar más de una regla con grados de activación *similares*, por lo que elegir únicamente una de ellas puede suponer una pérdida de información importante en la explicación facilitada. Proponemos una serie de definiciones que contemplan la posible inclusión de más de una regla en la explicación factual.

IV-A. Explicaciones factuales

Las reglas factuales son aquellas que nos ayudan a explicar la clasificación de una instancia en una determinada categoría. Dada una instancia x , y $c(x)$ su clasificación, sea $R'(c(x))$ el subconjunto de reglas de peso máximo en $R(x) \cap R(c(x))$, es decir, las reglas de peso máximo activadas por x con consecuente $c(x)$. Entonces, se define un factual $R_f(x)$ de x como un subconjunto de $R'(c(x))$ que explica la clasificación para la instancia x . Como caso particular, cuando $R'(c(x))$ sólo tiene una regla, esta será el factual.

IV-A1. Definiciones de factual: Sea x una instancia, y $|R'(c(x))| = p'(x)$. Sean $r_1(x), r_2(x), \dots, r_{p'(x)}(x)$ los elementos de $R'(c(x))$ que suponemos ordenados de mayor a menor grado de activación, es decir, $AD(r_1(x)) \geq AD(r_2(x)) \geq \dots \geq AD(r_{p'(x)}(x))$. A continuación proponemos tres formas de obtener un factual de x :

- Definimos el **factual asociado a la media**, $m\text{-}fact(x)$, como el subconjunto $R_f(x) = \{r_1(x), \dots, r_q(x)\}$ formado por las q reglas de $R'(c(x))$ para las que se cumple

$$AD(r_j(x)) \geq \frac{\sum_{j=1}^{p'(x)} AD(r_j(x))}{p'(x)}, \forall j \in \{1, \dots, q\} \quad (2)$$

- Dado $\lambda \in (0, 1]$, definimos el **factual asociado al λ -cociente**, $c(\lambda)\text{-}fact(x)$, como el subconjunto $R_f(x) = \{r_1(x), \dots, r_q(x)\}$ formado por las primeras q reglas de $R'(c(x))$ para las que se cumple

$$\frac{AD(r_q(x))}{AD(r_{q+1}(x))} > 1 + \lambda \wedge \frac{AD(r_j(x))}{AD(r_{j+1}(x))} \leq 1 + \lambda \quad (3)$$

$$\forall j \in \{1, \dots, q-1\}$$

- Dados $\lambda \in (0, 1]$ y $\beta \in (0, 1]$, definimos el **λ -cociente de masa mínima β** , $c(\lambda, \beta)\text{-}fact(x)$, como una variante de $c(\lambda)\text{-}fact(x)$ en la que exigimos que el sumatorio de todos los grados de activación de las reglas en el factual sea mayor que un umbral β . Es decir, el subconjunto $R_f(x) = \{r_1(x), \dots, r_q(x)\}$ formado por las primeras q reglas de $R'(c(x))$, para las que se cumple:

$$\sum_{j=1}^q AD(r_j(x)) \geq \beta \wedge \frac{AD(r_q(x))}{AD(r_{q+1}(x))} > 1 + \lambda \quad (4)$$

IV-B. Robustez de una explicación factual

Sea

$$R_x(c_i) = \{r \in R(x) : c(r) = c_i\} = R(x) \cap R(c_i),$$

el conjunto de reglas activadas por la instancia x con consecuente c_i . Definimos

$$R_f^*(x) = R_f(x) \cup \bigcup_{\substack{i=1, \dots, m \\ c_i \neq c(x)}} R_x(c_i)$$

como el conjunto de reglas del factual y las disparadas por la instancia para el resto de las clases (independientemente de que sean o no de peso máximo).

Consideramos que un factual $R_f(x)$ es robusto cuando:

$$c(R_f^*(x)) = c(R_f(x))$$

En la literatura, la mayoría de métodos consideran la explicación factual para $c(x)$ como una única regla. Sin embargo, existen casos en los que esto podría no resultar robusto si el proceso de inferencia utiliza varias reglas para determinar la clasificación de la instancia. Este es otro motivo para la definición de explicaciones factuales que admitan múltiples reglas, aumentando así la robustez de la explicación.

IV-C. Explicaciones contrafactuales

Una explicación contrafactual para una instancia x es aquella que nos muestra los cambios mínimos que habría que realizar a x para que cambie de clase. Se generará una explicación contrafactual por cada clase alternativa. Proporcionamos distintos métodos.

1. Contrafactual con respecto a la instancia x y una clase $\bar{c} \neq c(x)$, $R_{cf}(x, \bar{c})$. Dada una clase $\bar{c} \in C$, denotaremos por $R'(\bar{c})$ el conjunto de todas las reglas del árbol de peso máximo y consecuente $\bar{c}$.

Definimos la distancia de r a la instancia x como

$$d(r, x) = \sum_{F_i \in V(r)} \left[(1 - \mu_{i, z_i}(x_i)) \cdot \frac{|z_i - f(x_i)|}{k_i} \right] \quad (5)$$

donde $V(r)$ es el conjunto de variables difusas que aparecen en la regla r y v_{i,z_i} el conjunto difuso asociado al literal de la variable F_i en r .

El factor $(1 - \mu_{i,z_i}(x_i))$ considera el grado de pertenencia de la instancia a cada uno de esos conjuntos difusos, de tal manera que no solo se tiene en cuenta la distancia al conjunto sino también si x pertenece al mismo.

En cuanto al cociente $\frac{|z_i - f(x_i)|}{k_i}$, tiene en cuenta la distancia entre el índice del conjunto v_{i,z_i} del literal y el índice del conjunto al que la instancia x tiene mayor grado de pertenencia para la variable F_i .

Definimos la explicación contrafactual $R_{cf}(x, \bar{c})$ como

$$\arg \min_{r \in R'(\bar{c})} d(r, x)$$

2. Contrafactual con respecto al factual $R_f(x)$ y una clase $\bar{c} \neq c(x)$, $R_{cf}(x, R_f(x), \bar{c})$. Nuestro objetivo ahora es generar una explicación contrafactual que se diferencie lo mínimo posible de la explicación factual, pero clasifique la instancia x en $\bar{c}$.

Comencemos definiendo la distancia entre una regla contrafactual r y una regla ρ del factual $R_f(x)$ con $V(r) = V(\rho)$ como

$$d_{rule}(r, \rho, x) = \sum_{F_i \in V(r)} \left[(1 - \mu_{i,z_{r_i}}(x_i)) \cdot \frac{|z_{r_i} - z_{\rho_i}|}{k_i} \right] \quad (6)$$

con $v_{i,z_{r_i}}$ (resp. $v_{i,z_{\rho_i}}$) el conjunto difuso asociado al literal de la variable F_i en r (resp. ρ).

Por otra parte, dado $S \subseteq V(r)$, denotamos por $r \downarrow^S$ la simplificación de la regla r resultante de eliminar en el antecedente los literales asociados a las variables que no aparecen en el subconjunto S .

Definimos entonces la distancia entre r y $R_f(x) = \{\rho_1, \dots, \rho_q\}$ como

$$d(r, R_f(x)) = \sum_{\rho_j \in R_f(x)} [(1 - md(\rho_j(x))) \cdot cf_dist(r, \rho_j, x)] \quad (7)$$

con

$$cf_dist(r, \rho_j, x) = |V(r) \cup V(\rho_j)| - |V(r) \cap V(\rho_j)| + d_{rule}(r \downarrow^{V(r) \cap V(\rho_j)}, r \downarrow^{V(r) \cap V(\rho_j)}, x) \quad (8)$$

donde $|V(r) \cup V(\rho_j)| - |V(r) \cap V(\rho_j)|$ es el número de variables difusas diferentes entre las reglas r y ρ_j

Es decir, $d(r, R_f(x))$ es la suma de la distancia contrafactual cf_dist (Ecuación 8) de r a cada $\rho_j \in R_f(x)$, ponderadas por el grado de emparejamiento md de estas reglas ρ_j a la instancia x . De este modo tendrá más influencia una regla contrafactual que esté muy cerca de la “mejor” regla del factual que una que esté muy cerca de la “peor”.

Al calcular este tipo de distancias, se puede interpretar de manera distinta qué significa “realizar el mínimo número de cambios”. Hay dos tipos de cambios diferenciados entre dos reglas: una modificación de una variable que ya existe; y una adición o substracción de una variable que no existe. En este trabajo, se le ha dado más importancia a añadir o eliminar variables, haciendo que su contribución a la distancia sea mayor. El razonamiento detrás de esta decisión es que dada una regla, la modificación de una variable que ya existe en la

regla es una menor perturbación que considerar una variable nueva, haciendo así que la longitud de la regla se modifique.

Finalmente, definimos la explicación contrafactual $R_{cf}(x, R_f(x), \bar{c})$ como

$$\arg \min_{r \in R'(\bar{c})} d(r, R_f(x)).$$

V. EJEMPLO ILUSTRATIVO

V-A. Instancia a través del árbol difuso

Para ilustrar los conceptos introducidos se ha entrenado un árbol de decisión difuso usando el método propuesto en [9] para el conjunto de datos *wine* [10], definido mediante 13 atributos numéricos y una clase con tres categorías (*type 0, type 1, type 2*). Para el ejemplo hemos seleccionado una instancia que mostramos a continuación indicando el grado de pertenencia a las variables difusas aprendidas para la construcción del árbol (se le han asignado etiquetas del conjunto (Muy Bajo, Bajo, Medio, Alto, Muy Alto)). Por claridad, solo se muestran las variables que aparecen en las reglas disparadas y explicaciones generadas. Se usa el formato: $F_i : \{“v_{i,1}”: \mu_{v_{i,1}}(x_i), \dots, “v_{i,k_i}”: \mu_{v_{i,k_i}}(x_i)\}$ para los k_i conjuntos difusos de cada variable, y se muestran sólo las variables que aparecen en las reglas por claridad:

flavanoids : {“muy bajo” : 0.14, “bajo” : 0.86,
“alto” : 0, “muy alto” : 0},
color_intensity : {“bajo” : 0.357, “medio” : 0.643,
“alto” : 0},
alcohol : {“muy bajo” : 0, “bajo” : 0.616,
“alto” : 0.384, “muy alto” : 0},
hue : {“bajo” : 0.45, “medio” : 0.55,
“alto” : 0},
proanthocyanins : {“bajo” : 0.64, “medio” : 0.36,
“alto” : 0}
od280/od315 : {“bajo” : 0.274, “medio” : 0.726,
“alto” : 0}
class : *type 2*

V-B. Factual

Utilizando la Ecuación 2 para obtener el factual $R_f(x)$, resultan cuatro reglas (consecuente = *type 2*):

- $r_1 : (flavanoids \text{ bajo}) \wedge (color_intensity \text{ medio}) \wedge (alcohol \text{ bajo}) \wedge (hue \text{ bajo}), w(r_1) = 1$
 $AD(r_1(x)) = 0.45$
- $r_2 : (flavanoids \text{ bajo}) \wedge (color_intensity \text{ medio}) \wedge (alcohol \text{ alto}) \wedge (proanthocyanins \text{ bajo}), w(r_2) = 1$
 $AD(r_2(x)) = 0.38$
- $r_3 : (flavanoids \text{ bajo}) \wedge (color_intensity \text{ bajo}) \wedge (alcohol \text{ alto}) \wedge (od280/od315 \text{ medio}), w(r_3) = 1$
 $AD(r_3(x)) = 0.36$
- $r_4 : (flavanoids \text{ bajo}) \wedge (color_intensity \text{ medio}) \wedge (alcohol \text{ alto}) \wedge (proanthocyanins \text{ medio}), w(r_4) = 0.92$
 $AD(r_4(x)) = 0.33$

Se muestra además el grado de activación (Sección III-B). A modo de ejemplo, para r_4 sería

$$AD(r_4(x)) = md(r_4(x)) \cdot w(r_4(x)) = \min(0.86, 0.643, 0.384, 0.36) \cdot 0.92 = 0.33$$

Podemos ver que las cuatro reglas son ciertamente parecidas y con un grado de activación similar. En función de una ligera variación de *alcohol* y *color_intensity*, la variable que explicaría la clasificación de la instancia cambia (*hue*, *proanthocyanins* u *od280/od315*).

Probablemente para el usuario sea más interesante tener esta información más completa que únicamente la regla r_1 .

V-C. Robustez

Para comprobar la robustez del factual, el primer paso es comprobar si existen reglas del árbol que se disparan para el resto de clases (*type* 0 y/o *type* 1). En este caso, se disparan suficientes reglas de la clase *type* 1 para alcanzar un grado de activación acumulado igual a 1.22. Este hecho provocaría que si únicamente incluyéramos r_1 como explicación factual, dicho factual no sería robusto. Sin embargo, el factual formado por $\{r_1, r_2, r_3, r_4\}$ sí lo es puesto que $\sum_{i=1}^4 AD(r_i(x)) = 1.52$.

V-D. Contrafactual

Dado que el clasificador obtiene el valor *type* 2 para la clase, pueden existir explicaciones contrafactuales para las clases *type* 0 y *type* 1. Usaremos la segunda definición propuesta en la Sección IV-C.

V-D1. Clase *type* 1: Consideramos todas las reglas de peso máximo con consecuente *type* 1. Las dos más cercanas al factual en función de la distancia definida son:

$$\begin{aligned} r_{c1} &: (\text{flavanoids bajo}) \wedge (\text{color_intensity medio}) \wedge \\ & (\text{alcohol alto}) \wedge (\text{proanthocyanins alto}) \\ d(r_{c1}, R_f(x)) &= 12.98 \\ r_{c2} &: (\text{flavanoids bajo}) \wedge (\text{color_intensity medio}) \wedge \\ & (\text{alcohol muy alto}) \wedge (\text{proanthocyanins alto}) \\ d(r_{c2}, R_f(x)) &= 14.68 \end{aligned}$$

De estas dos reglas, sería r_{c1} el contrafactual para la clase *type* 1.

V-D2. Clase *type* 0: Razonando igual, las dos reglas que minimizan la distancia a la explicación factual son

$$\begin{aligned} r_{c1} &: (\text{flavanoids alto}) \wedge (\text{alcohol muy alto}) \\ d(r_{c1}, R_f(x)) &= 24.18 \\ r_{c2} &: (\text{flavanoids muy alto}) \wedge (\text{alcohol alto}) \\ d(r_{c2}, R_f(x)) &= 24.72 \end{aligned}$$

Igual que antes, r_{c1} será el contrafactual para la clase *type* 0.

VI. EVALUACIÓN

VI-A. Conjuntos de datos

Los experimentos se han realizado con los conjuntos de datos *iris* [10], [11], *wine* [10], [12] y *beer* [13]. Los dos primeros conjuntos de datos se utilizan comúnmente en problemas de clasificación y el tercero se ha usado en XAI [?].

VI-B. Metodología experimental

Para cada conjunto de datos se ha entrenado un árbol de decisión difuso usando *holdout* como metodología de validación (70 % entrenamiento y 30 % test). Así, las instancias cuya clasificación debe ser explicada son las del conjunto de *test*.

VI-C. Recursos computacionales

Todos los algoritmos han sido programados en Python 3.8, debido a la potencia de librerías como *scikit-learn* [14] que ayudan con el tratamiento de datos y las estructuras necesarias. Para el algoritmo de aprendizaje del árbol de clasificación y la definición de las variables difusas se ha seguido [9], en particular, la implementación del árbol FMDT (*Fuzzy Multiway Decision Tree*). Posteriormente, se han añadido los métodos necesarios para calcular las explicaciones factuales y contrafactuales propuestas en este artículo. Para garantizar la reproducibilidad de las pruebas realizadas, tanto el código como (el acceso a) los datos se publicará en un repositorio abierto en Github.

VI-D. Explicaciones factuales y robustez.

En la Tabla I se muestra información sobre las explicaciones factuales obtenidas y su robustez. En el caso de los criterios que dependen de umbrales (λ , β) se han probado varios valores. Las columnas de la tabla representan:

- **Config:** Método para obtener la explicación factual (Sección IV-A1) y parámetros utilizados.
- $q > 1$: Proporción de instancias (del conjunto de test) para las que la explicación factual incluye más de una regla.
- **NR-Fact:** Proporción de instancias (del conjunto de test) cuya explicación no es robusta.
- **Len:** Longitud media de las explicaciones factuales (medida en número de reglas).

Además, se calculan las siguientes constantes, que son independientes del método considerado para generar la explicación factual:

- **ExistCF:** Proporción de instancias (del conjunto de test) para las que $R(x)$ contiene reglas con consecuente distinto a la clase predicha ($c(x)$).
- **NR-Rule:** Proporción de instancias (del conjunto de test) para las que la explicación factual únicamente por la regla con mayor grado de activación no es robusto.

Tabla I
EXPERIMENTOS SOBRE LA ROBUSTEZ

Configuración	Iris $q > 1$	ExistCF: 0.6 NR-Fact	NR-Rule: 0.13 Len	Wine $q > 1$	ExistCF: 0.677 NR-Fact	NR-Rule: 0.1 Len	Beer $q > 1$	ExistCF: 0.625 NR-Fact	NR-Rule: 0.1 Len
$m\text{-fact}(x)$	0.2	0.06	1.22	0.559	0.025	2.05	0.475	0.1	1.58
$c(\lambda)\text{-fact}(x)$ $\lambda: 0.1$	0.06	0.06	1.1	0.0508	0.075	1.033	0.175	0.1	1.2
$c(\lambda)\text{-fact}(x)$ $\lambda: 0.25$	0.12	0.06	1.22	0.237	0.025	1.644	0.225	0.08	1.325
$c(\lambda, \beta)\text{-fact}(x)$ $\lambda: 0.1 \beta: 0.5$	0.06	0.06	1.1	0.0508	0.075	1.033	0.188	0.1	1.213
$c(\lambda, \beta)\text{-fact}(x)$ $\lambda: 0.1 \beta: 0.7$	0.28	0.03	1.34	0.13	0	1.305	0.3	0.04	1.4
$c(\lambda, \beta)\text{-fact}(x)$ $\lambda: 0.25 \beta: 0.5$	0.12	0.06	1.22	0.237	0.025	1.644	0.225	0.085	1.325
$c(\lambda, \beta)\text{-fact}(x)$ $\lambda: 0.25 \beta: 0.7$	0.34	0.03	1.64	0.254	0	1.711	0.325	0.04	1.475

En la Tabla I se puede observar que la media suele ser el método que obtiene reglas más largas, sin que ello necesariamente implique que sean más robustas. Los otros dos métodos, que tienen en cuenta la similaridad entre las reglas del factual, obtienen unas reglas algo más cortas y además suelen ser más robustas. Además, en todas las bases de datos se ve que existen más reglas que no son robustas a factuales, independientemente de cómo estén contruidos. Así, se demuestra una necesidad de buscar esta robustez con múltiples reglas en el factual.

VI-E. Explicaciones contrafactuales

Evaluremos los siguientes parámetros:

- **Número de contrafactuales (NumCF)**. El número de posibles contrafactuales por clase. Un mayor número de posibles contrafactuales representa un mayor factor de ramificación del árbol.
- **Mejor distancia mínima (BestMinDist)**. La distancia del mejor contrafactual a la explicación factual.
- **Longitud del Contrafactual (CFLen)**. Calculada como el número medio de literales que contiene el contrafactual.

En la Tabla II se muestran estos datos para los distintos métodos propuestos para calcular la explicación contrafactual (**Configuración (Config)**). Se empleará una única configuración de parámetros.

Tabla II
EXPERIMENTOS SOBRE LOS CONTRAFACUALES.

Configuración	Iris CFLen	NumCF: 3.73 BestMinDist	Wine CFLen	NumCF: 3.73 BestMinDist	Beer CFLen	NumCF: 9.37 BestMinDist
CF de x	1.69	2.42	3.86	8.97	2.62	0.95
CF de $m\text{-}fact(x)$	1.53	0.391	2.61	1.93	2.203	0.721
CF de $c(\lambda)\text{-}fact(x)$ $\lambda : 0.1$	1.53	0.339	2.79	0.767	2.191	0.477
CF de $c(\lambda, \beta)\text{-}fact(x)$ $\lambda : 0.1 \beta : 0.7$	1.53	0.72	2.56	2.11	2.2	0.93

De los resultados podemos observar que en general los contrafactuales que se extraen con respecto a la instancia tienen una longitud ligeramente superior que lo generados con respecto a la explicación factual. Esto se debe a que al estar la instancia definida sobre todos los atributos, no penaliza necesariamente que la regla contrafactual tenga más literales, sino que tenga menor distancia a sus valores (conjuntos difusos). Por otro lado, la gran diferencia de distancia entre los contrafactuales $m\text{-}fact$ y los contrafactuales $c\text{-}fact$ se debe a que las explicaciones factuales $c\text{-}fact$ contienen más reglas. Al calcularse la distancia como la suma de la distancia a todas las reglas del factual, cuanto mayor sea el factual (como se observa en la Tabla I) mayor será la distancia.

VII. CONCLUSIONES Y TRABAJO FUTURO

Partiendo de la hipótesis de que en un entorno de razonamiento difuso generar explicaciones factuales con una única regla puede suponer una pérdida de información, se han propuesto definiciones alternativas que permiten incluir múltiples reglas. Como se ha comprobado en los experimentos, en la mayoría de los casos esto no es necesario, pero sí en un porcentaje nada despreciable. Se ha introducido además el concepto de robustez para la explicación factual, comprobándose en la evaluación realizada la mejoría que en este sentido supone introducir más de una regla (en los casos más inciertos). Finalmente se han propuesto definiciones de explicación contrafactual ligadas a las proporcionadas para las explicaciones factuales.

Por tratarse de un primer trabajo en esta línea de investigación, es obvio que existen distintas vías de ampliación, como es su experimentación/integración en métodos agnósticos de explicación. Concretamente, pretendemos generalizar el estudio presentado en [4] al caso de los árboles de clasificación

difusos. Consideramos además de interés estudiar métodos de simplificación de las explicaciones factuales generadas, de forma que estas sean lo más compactas posible, mejorando así su comprensión por parte de los usuarios.

AGRADECIMIENTOS

Este trabajo ha sido parcialmente financiado por el Gobierno de España, fondos FEDER y la Junta de Comunidades de Castilla-La Mancha mediante los proyectos PID2019-106758GB-C33 /AEI/ 10.13039/501100011033 y SBPLY/17/180501/000493. Guillermo Tomás Fernández también está financiado por FPU19/02930 del MCIU.

REFERENCIAS

- [1] G. D. P. Regulation, "General data protection regulation (gdpr)," Intersoft Consulting, Accessed in October 24, vol. 1, 2018.
- [2] M. T. Ribeiro, S. Singh, and C. Guestrin, "why should i trust you?" explaining the predictions of any classifier," in *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining*, 2016, pp. 1135–1144.
- [3] —, "Anchors: High-precision model-agnostic explanations," in *Thirty-Second AAAI Conference on Artificial Intelligence*, 2018.
- [4] R. Guidotti, A. Monreale, F. Giannotti, D. Pedreschi, S. Ruggieri, and F. Turini, "Factual and counterfactual explanations for black box decision making," *IEEE Intell. Syst.*, vol. 34, no. 6, pp. 14–23, 2019.
- [5] I. Stepin, J. M. Alonso, A. Catalá, and M. Pereira-Fariña, "A survey of contrastive and counterfactual explanation generation methods for explainable artificial intelligence," *IEEE Access*, vol. 9, pp. 11974–12001, 2021.
- [6] I. Stepin, J. M. Alonso, A. Gatala, and M. Pereira-Fariña, "Generation and evaluation of factual and counterfactual explanations for decision trees and fuzzy rule-based classifiers," in *2020 IEEE International Conference on Fuzzy Systems (FUZZ-IEEE)*. IEEE, 2020, pp. 1–8.
- [7] A. Shih, A. Choi, and A. Darwiche, "A symbolic approach to explaining bayesian network classifiers," in *Proceedings of the Twenty-Seventh International Joint Conference on Artificial Intelligence, IJCAI 2018, July 13-19, 2018, Stockholm, Sweden*, J. Lang, Ed., 2018, pp. 5103–5111.
- [8] E. Albini, A. Rago, P. Baroni, and F. Toni, "Relation-based counterfactual explanations for bayesian network classifiers," in *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence, IJCAI 2020, C. Bessiere, Ed., 2020*, pp. 451–457.
- [9] A. Segatori, F. Marcelloni, and W. Pedrycz, "On distributed fuzzy decision trees for big data," *IEEE Transactions on Fuzzy Systems*, vol. 26, no. 1, pp. 174–192, 2017.
- [10] D. Dua and C. Graff, "UCI machine learning repository," 2017. [Online]. Available: <http://archive.ics.uci.edu/ml>
- [11] R. A. Fisher, "The use of multiple measurements in taxonomic problems," *Annals of eugenics*, vol. 7, no. 2, pp. 179–188, 1936.
- [12] M. Forina et al., "An extendible package for data exploration, classification and correlation. institute of pharmaceutical and food analysis and technologies," 1998.
- [13] G. Castellano, C. Castiello, and A. M. Fanelli, "The fisdet software: Application to beer style classification," in *2017 IEEE International Conference on Fuzzy Systems (FUZZ-IEEE)*. IEEE, 2017, pp. 1–6.
- [14] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot, and E. Duchesnay, "Scikit-learn: Machine learning in Python," *Journal of Machine Learning Research*, vol. 12, pp. 2825–2830, 2011.

Comparando variabilidades de conjuntos aleatorios

Juan Jesús Salamanca

Departamento de Estadística e I.O. y D.M.

Universidad de Oviedo

Gijón, España

salamancajuan@uniovi.es

Susana Montes

Departamento de Estadística e I.O. y D.M.

Universidad de Oviedo

Oviedo, España

montes@uniovi.es

Resumen—Los conjuntos aleatorios son modelos de probabilidad destacados en Probabilidades Imprecisas. Un problema que nos ocupa es el de establecer alguna decisión en un ambiente de incertidumbre. Aquí, los ordenes estocásticos juegan un papel fundamental, ya que suponen técnicas de decisión en este contexto. En esta contribución introducimos un nuevo orden dispersivo para conjuntos aleatorios. Es decir, un método de preferencia entre dos conjuntos aleatorios que está basado en la variabilidad de los elementos aleatorios. Además, estudiaremos brevemente las propiedades más destacadas que presenta este orden.

Index Terms—conjunto aleatorio, orden dispersivo, decisión bajo imprecisión

I. INTRODUCCIÓN

Los ordenes estocásticos son herramientas muy útiles en Estadística (véanse las referencias [4] y [6], por ejemplo). Dentro de estos ordenes nos encontramos los ordenes dispersivos, que son aquellas herramientas que permiten comparar dos elementos aleatorios en términos de sus variabilidades.

Así, encontramos en primer lugar al orden dispersivo usual.

Definición 1. [6, Sec. 3.B] Dadas dos variables univariantes X e Y , el orden dispersivo usual elige X sobre Y si

$$F_X^{-1}(\alpha) - F_X^{-1}(\beta) \leq F_Y^{-1}(\alpha) - F_Y^{-1}(\beta)$$

donde F_X^{-1} y F_Y^{-1} denotan las funciones inversas continuas por la derecha de las funciones de distribución F_X y F_Y de X e Y , respectivamente.

El principal escollo que presenta esta técnica es su dificultad para ser extendida a un punto aleatorio de un espacio métrico arbitrario.

Por otro lado encontramos el orden dispersivo débil.

Definición 2. [2] Dados dos puntos aleatorios X e Y de un espacio métrico $(\mathcal{M}, d)$, diremos que X es preferido a Y en el orden dispersivo débil si se satisface

$$\mathbb{P}\{d(X, X') \leq r\} \leq \mathbb{P}\{d(Y, Y') \leq r\}$$

Este trabajo ha sido parcialmente financiado por el proyecto PGC2018-098623-B-I00 del Ministerio de Ciencia, Innovación y Universidades. Los autores quieren agradecer dicha financiación.

para todo $r \geq 0$, donde X' e Y' son puntos aleatorios idénticamente distribuidos a X e Y , respectivamente y $\mathbb{P}$ denota la medida de probabilidad subyacente.

Una dificultad que presenta esta otra técnica es que no produce comparación en numerosas ocasiones.

Otros ordenes dispersivos se basan en contracciones. Sin embargo, tales aplicaciones pueden no tener sentido para algunas clases de espacios topológicos (véase [5]).

Finalmente nos encontramos con el orden dispersivo por bolas.

Definición 3. [5] Dados dos puntos aleatorios X e Y de un espacio métrico $(\mathcal{M}, d)$, se dice que X es preferido a Y en el orden dispersivo por bolas si para todo punto p existe otro punto q tal que

$$\mathbb{P}\{X \in \mathbb{B}_q(r)\} \geq \mathbb{P}\{Y \in \mathbb{B}_p(r)\}$$

se satisface para todo $r \geq 0$, donde $\mathbb{B}_x(r)$ denota la bola (cerrada) de centro x y radio r .

Podemos interpretar este orden diciendo que, para todo punto p , la distribución de $d(Y, p)$ está más dispersa que la distribución de $d(X, q)$ para algún punto q . Es decir, controlamos la variabilidad de los puntos aleatorios a través de bolas geométricas. Esta última técnica tiene la gran ventaja de ser lo suficientemente maleable como para poder extenderse a otros esquemas de trabajo.

Por otro lado, conviene notar que en numerosas ocasiones la información disponible es difusa, vaga o imprecisa en general. Este es el punto de partida de lo que se ha convenido en llamar teoría de Probabilidades Imprecisas (véase [7]). En líneas generales, esta rama de la Estadística estudia modelos de probabilidad bajo imprecisión e incertidumbre. En este esquema de trabajo, los conjuntos aleatorios, en su interpretación disyuntiva u óptica, son modelos destacados de probabilidades imprecisas (véase [1]).

En este contexto, necesitamos alguna herramienta que nos permita elegir como preferido un conjunto aleatorio frente a otro en términos de sus variabilidades. En esta contribución

presentaremos un orden dispersivo nuevo para conjuntos aleatorios. Este orden dispersivo está basado en una comparación de las funciones de plausibilidad de los conjuntos aleatorios, con lo que se puede entender como una extensión del orden dispersivo por bolas anteriormente comentado.

Recuérdese que la función de plausibilidad caracteriza al conjunto aleatorio, tal y como se recoge en el celebrado teorema de Choquet (véase [3] y las referencias allí citadas). En este sentido contemplamos la información estadística que poseen los conjuntos aleatorios a comparar.

II. UN NUEVO ORDEN DISPERSIVO PARA CONJUNTOS ALEATORIOS

Recordamos en primer lugar la noción de conjunto aleatorio cerrado. Para ello, comencemos fijando un espacio topológico Hausdorff, segundo-contable y localmente compacto $\mathcal{M}$, al que le dotaremos de una distancia d . La distancia nos permite considerar las bolas (geométricas); así, denotaremos por $\mathbb{B}_p(r)$ a la bola de centro p y radio r . Al mismo tiempo denotemos a la clase de conjuntos cerrados de $\mathcal{M}$ por $\mathcal{F}$.

Definición 4. [3] Un conjunto aleatorio (cerrado) $\hat{X}$ sobre $\mathcal{M}$ es una aplicación $\hat{X} : (\Omega, \sigma, \mathbb{P}) \rightarrow \mathcal{F}$ tal que

$$\left\{ w \in \Omega : \hat{X}(w) \cap K \neq \emptyset \right\} \in \sigma$$

donde K es un conjunto compacto de $\mathcal{M}$ y $(\Omega, \sigma, \mathbb{P})$ es un espacio de probabilidad completo.

De la definición anterior se deduce inmediatamente que está bien definida la aplicación $Pl_{\hat{X}}$ que asocia a cada compacto K de $\mathcal{M}$ la probabilidad $Pl_{\hat{X}}(K) = \mathbb{P} \left\{ w \in \Omega : \hat{X}(w) \cap K \neq \emptyset \right\}$. Esta aplicación se llama *función de plausibilidad* (del conjunto aleatorio).

El celebrado resultado de Choquet establece que todo conjunto aleatorio queda caracterizado por su función de plausibilidad (véase [3]).

En este trabajo estamos interesados en obtener una técnica de comparación entre dos conjuntos aleatorios en términos dispersivos. Para tal fin, extenderemos en este sentido el orden dispersivo que se da en [5] para puntos aleatorios.

Definición 5. Sean $\hat{X}$ e $\hat{Y}$ dos conjuntos aleatorios de $\mathcal{M}$. Diremos que $\hat{X}$ es preferido a $\hat{Y}$ en el orden dispersivo de plausibilidad si para todo punto p existe otro punto q tal que

$$Pl_{\hat{X}}(\mathbb{B}_q(r)) \geq Pl_{\hat{Y}}(\mathbb{B}_p(r)) \quad (1)$$

se satisface para todo $r \geq 0$, donde $Pl_{\hat{X}}$ y $Pl_{\hat{Y}}$ denotan las funciones de plausibilidad de $\hat{X}$ e $\hat{Y}$, respectivamente.

Es claro observar que el orden introducido solo tiene en cuenta el grado de dispersión de los conjuntos aleatorios. De hecho, recoge la información acerca de cómo se aleja el

conjunto aleatorio respecto de puntos. Si un conjunto aleatorio es preferido a otro, entonces podemos decir que, en algún sentido, se acerca más a un punto el primero que el segundo.

Las propiedades más destacables de esta nueva técnica se recogen a continuación. Antes de ello, recordemos el concepto de isometría, que generaliza a las traslaciones y rotaciones de un espacio euclídeo. Dado un espacio métrico $(\mathcal{M}, d)$, una *isometría* ϕ es una aplicación $\phi : \mathcal{M} \rightarrow \mathcal{M}$ biyectiva que conserva distancias, es decir, que para cualesquiera puntos p y q , se tiene: $d(p, q) = d(\phi(p), \phi(q))$. Dado un conjunto aleatorio $\hat{X}$ definimos el conjunto aleatorio $\phi \circ \hat{X}$ como: $(\phi \circ \hat{X})(w) = \phi(\hat{X}(w))$, para todo w del espacio de probabilidad subyacente.

Proposición 1. El orden dispersivo de plausibilidad presenta las siguientes propiedades:

- (a) Es transitivo. De forma más precisa, sean $\hat{X}$, $\hat{Y}$ y $\hat{Z}$ tres conjuntos aleatorios de $\mathcal{M}$. Si $\hat{X}$ es preferido a $\hat{Y}$ e $\hat{Y}$ es preferido a $\hat{Z}$, entonces $\hat{X}$ es preferido a $\hat{Z}$.
- (b) Es reflexivo y, por tanto, un preorden, pero no es un orden.
- (c) Es compatible con isometrías. De forma más precisa, sea ϕ una isometría de $(\mathcal{M}, d)$. Entonces, el conjunto aleatorio $\hat{X}$ es preferido al conjunto aleatorio $\hat{Y}$ si y solo si $\phi \circ \hat{X}$ es preferido a $\hat{Y}$ si y solo si $\hat{X}$ es preferido a $\phi \circ \hat{Y}$.

Idea de la demostración. - Para demostrar este resultado hay que tener en cuenta:

- (a) La propiedad de transitividad se obtiene a partir de la transitividad que se hereda de la comparación de las correspondientes funciones de plausibilidad, y haciendo uso de las definiciones.
- (b) Es trivial que es reflexivo, sin más que considerar como q el valor de p . Sin embargo se pueden encontrar dos conjuntos aleatorios distintos tales que cualquiera de los dos sea preferido al otro.
- (c) Se puede deducir a partir del siguiente hecho:

$$\begin{aligned} Pl_{\phi \circ \hat{X}}(\mathbb{B}_p(r)) &= \mathbb{P} \left\{ w : \phi \circ \hat{X}(w) \cap \mathbb{B}_p(r) \neq \emptyset \right\} = \\ &= \mathbb{P} \left\{ w : \phi^{-1} \left(\phi \circ \hat{X}(w) \cap \mathbb{B}_p(r) \right) \neq \emptyset \right\} = \\ &= \mathbb{P} \left\{ w : \hat{X}(w) \cap \phi^{-1} \circ \mathbb{B}_p(r) \neq \emptyset \right\} = \\ &= \mathbb{P} \left\{ w : \hat{X}(w) \cap \mathbb{B}_{\phi^{-1}(p)}(r) \neq \emptyset \right\} = \\ &= Pl_{\hat{X}}(\mathbb{B}_{\phi^{-1}(p)}(r)). \end{aligned}$$

En este sentido, la isometría ϕ conserva la función de plausibilidad.

Por un lado, la propiedad de transitividad es una propiedad muy importante que un buen orden debe presentar. Por otro lado, que el orden sea compatible con las isometrías está asociado a su carácter de ser un orden dispersivo, puesto

que no tiene en cuenta “posiciones” del conjunto aleatorio, sino cómo puede alejarse o acercarse a un punto dado. De hecho, recuérdese que el concepto de isometría generaliza propiamente al grupo de traslaciones y rotaciones de un espacio euclídeo. En particular, esta propiedad indica que el orden es invariante frente a estas transformaciones de un espacio euclídeo.

El orden que estamos introduciendo ha de proveer argumentos lógicos y de peso para el/la decisor/a. A continuación daremos un argumento para tal fin.

Dado que estamos trabajando en términos de probabilidades imprecisas, el conjunto aleatorio se puede interpretar como un modelo que engloba un conjunto de modelos de probabilidad exacta. En concreto, una selección no es sino un modelo de probabilidad exacto que es compatible con la información dada por el conjunto aleatorio, tal como se deduce de la siguiente definición.

Definición 6. [3] Una selección X de un conjunto aleatorio $\hat{X}$ sobre $(\mathcal{M}, d)$ es un punto aleatorio de $\mathcal{M}$ tal que $X(w) \in \hat{X}(w)$ casi seguro.

Podemos dar el siguiente resultado, que se particulariza para conjuntos aleatorios discretos (en esta contribución, aquellos cuyo espacio de probabilidad es discreto, por definición).

Teorema 1. Sean $\hat{X}$ e $\hat{Y}$ dos conjuntos aleatorios discretos sobre $(\mathcal{M}, d)$. Si $\hat{X}$ es preferido a $\hat{Y}$ en el orden dispersivo de plausibilidad, entonces para todo punto p y para toda selección Y de $\hat{Y}$ podemos encontrar un punto q y una selección X de $\hat{X}$ tal que

$$\mathbb{E}[d^2(Y, p)] \geq \mathbb{E}[d^2(X, q)].$$

En particular, para toda selección de $\hat{Y}$ se puede encontrar una selección de $\hat{X}$ con menor o igual varianza.

Idea de la demostración. Por reducción al absurdo. Asumamos que existe un punto p' y una selección Y' de $\hat{Y}$ tales que $\mathbb{E}[d^2(Y', p')] \geq \mathbb{E}[d^2(X, q)]$ se satisface para todo punto q y selección X de $\hat{X}$. Esto implica que existe un real no-negativo r para el cual $\mathbb{P}\{d(Y', p') \leq r\} > \mathbb{P}\{d(X, q)\}$ se satisface para todo punto q y selección X de $\hat{X}$. Nótese que la función de plausibilidad satisface $Pl_{\hat{Y}}(\mathbb{B}_{p'}(r)) \geq \mathbb{P}\{d(Y', p') \leq r\}$. El siguiente paso es usar que $\hat{X}$ es preferido a $\hat{Y}$ para deducir que debe existir un punto q' tal que $Pl_{\hat{X}}(\mathbb{B}_{q'}(r)) > \mathbb{P}\{d(X, q') \leq r\}$ se satisface para toda selección X de $\hat{X}$. La contradicción se encuentra con el hecho de que se puede obtener una selección X de $\hat{X}$ satisfaciendo $Pl_{\hat{X}}(\mathbb{B}_p(r)) = \mathbb{P}\{X \in \mathbb{B}_p(r)\}$, y para ello basta definir X como sigue: $X(w) \in \arg \min d_H(\hat{X}(w), p)$, donde d_H denota la distancia Hausdorff (que X es punto aleatorio se deduce inmediatamente a partir de la propiedad de que $\hat{X}$ es discreto).

La última afirmación de este teorema se deduce inmediatamente de la anterior.

El argumento entonces es claro: el/la decisor/a está eligiendo el conjunto aleatorio con el que es posible extraer un modelo exacto con menor varianza. Este es un argumento de peso que soporta la técnica que hemos introducido.

Ejemplo 1. Consideremos $\mathbb{R}$ dotado de su distancia usual d . Consideremos los conjuntos aleatorios $\hat{X} = [X - r, X + r]$ e $\hat{Y} = [Y - s, Y + s]$, donde X está normalmente distribuida con media μ_x y desviación típica σ_x , Y también está normalmente distribuida con media μ_y y desviación típica σ_y y r es una constante estrictamente mayor que cero. Podemos interpretar estos conjuntos aleatorios como modelos exactos de probabilidad a los que se les añade un error o incertidumbre de valor r . Es fácil probar que $\hat{X}$ es preferido a $\hat{Y}$ si $\sigma_x \leq \sigma_y$. Obsérvese que esta decisión está en consonancia con el teorema anterior, dado que en este caso seremos capaces de extraer una selección, un modelo exacto, de $\hat{X}$ con menor varianza que cualquier selección de $\hat{Y}$.

III. CONCLUSIONES

En general, un problema de decisión en un ambiente impreciso acarrea una dificultad extra respecto a los problemas clásicos de teoría de la probabilidad. En esta contribución hemos introducido un nuevo orden dispersivo para conjuntos aleatorios. Este orden está basado en comparar las funciones de plausibilidad sobre la familia de bolas centradas en un punto.

Sería interesante seguir estudiando las propiedades que presenta este nuevo orden, y cómo simplificar el proceso de decisión. En este sentido, cabe estudiar qué particularidades presenta el orden para los casos en los que el espacio métrico sea bien finito o bien tiene la estructura de un espacio vectorial o afín euclídeo.

Por otro lado, también cabe extender este orden para otros modelos destacados en Probabilidades Imprecisas, como las variables aleatorias fuzzy.

REFERENCIAS

- [1] I. Couso and D. Dubois. Statistical reasoning with set-valued information: Ontic vs. epistemic views. *International Journal of Approximate Reasoning*, pages 1502–1518, 2014.
- [2] A. Giovagnoli and H.P. Wynn. Multivariate dispersion orderings. *Statistics & probability letters*, pages 325–332, 1995.
- [3] I. Molchanov. Springer, 2005.
- [4] A. Müller and D. Stoyan. *Comparison Methods for Stochastic Models and Risks*. Wiley, 2002.
- [5] J.J. Salamanca. A universal, canonical dispersive ordering in metric spaces. *Journal of Statistical Planning and Inference*, 212:1–13, 2021.
- [6] M. Shaked and J.G. Shanthikumar. *Stochastic orders*. Springer Science & Business Media, 2007.
- [7] P. Walley. *Statistical reasoning with imprecise probabilities*. Chapman and Hall, 1991.

A fuzzy probability logic for compound conditionals

Tommaso Flaminio

IIIA – CSIC

08193 Bellaterra, Spain

ORCID 0000-0002-9180-7808

tommaso@iia.csic.es

Lluís Godo

IIIA – CSIC

08193 Bellaterra, Spain

ORCID 0000-0002-6929-3126

godo@iia.csic.es

Abstract—In this paper we propose a fuzzy modal logic for conditional probability that allows to represent and reason about the probability of not only basic conditional expressions of the form “ φ given ψ ”, written $(\varphi \mid \psi)$, but also compound conditional sentences such as “ φ given ψ and γ given χ ”, written $(\varphi \mid \psi) \wedge (\gamma \mid \chi)$, and more in general, any Boolean combination of basic ones. In order to formalize compound conditional formulas we will adopt the recently defined *Logic for Boolean Conditionals* (LBC) and hence formalize conditional probability as a simple (unconditional) probability of conditional sentences. In addition to such basic fuzzy modal logic for the probability of compound conditionals, we will also present some extensions and prove that each of them is sound and complete w.r.t. to a suitable class of probabilistic models. Furthermore, we will prove how to recover the usual interpretation of conditional probability, showing that, under minimal requirements, in these logics the probability of a basic conditional $(\varphi \mid \psi)$ can be safely taken as the conditional probability of φ given ψ , i.e. as the ratio $P(\varphi \wedge \psi)/P(\psi)$.

Index Terms—Conditional probability; Compound conditional; Fuzzy logic; Fuzzy modal logic.

I. INTRODUCTION

Fuzzy sets-based models and numerical uncertainty models, although sharing the feature of evaluating sentences in a totally ordered scale, usually the real unit interval $[0, 1]$, account for radically different notions of gradualness. From a formal point of view, these differences can be easily grasped if we consider their corresponding logics: fuzzy logics and uncertainty logics (in particular probability logics), respectively. In fact, while the former are truth-functional, i.e. the truth-value of a compound formula like $\varphi \vee \psi$ only depends on the truth-values of its components φ and ψ , the latter are not, since, for instance, the probability of $\varphi \vee \psi$ cannot be computed only from the probability of φ and the probability of ψ (it is also needed to know what is the probability of $\varphi \wedge \psi$).

Despite these differences, however, probability logics can be properly handled in a fuzzy logical setting by expanding the language of a fuzzy logic with a unary modality $P(\cdot)$ and interpreting, for every classical formula φ , the modal formula $P(\varphi)$ as “ φ is probable”. Clearly, $P(\varphi)$ is a fuzzy proposition, whose truth-degree can be taken as the probability of φ . More precisely, the fuzzy modal logic $\text{FP}(\mathbb{L})$, as firstly introduced in [11] and improved in [10], extends the language of Łukasiewicz logic $\mathbb{L}$ by the modal operator $P(\cdot)$ and uses the ground logic $\mathbb{L}$ to express the basic properties of a probability function.

In particular, it is worth to recall that the finite additivity of P can be expressed in $\text{FP}(\mathbb{L})$ by using the Łukasiewicz connective $\oplus$ whose standard interpretation is the truncated sum: for all $x, y \in [0, 1]$, $x \oplus y = \min\{1, x + y\}$. Very recently, in [1] the authors have studied in depth the relationship of this fuzzy logic-based approach to more traditional probability logics after Halpern et al. see e.g. [13].

In addition to simple probability, the paper [8] presents the logic $\text{FP}(\mathbb{L}\Pi)$ to deal with conditional probability by considering, instead of $\mathbb{L}$, the stronger logic $\mathbb{L}\Pi$. Such formalism can be roughly regarded as the expansion of Łukasiewicz logic by the connectives of product conjunction $\odot$ and product implication $\rightarrow_{\Pi}$. The standard semantics of $\odot$ and $\rightarrow_{\Pi}$ interprets them, respectively, by the usual product $\cdot$ and the function $x \rightarrow_{\Pi} y = 1$ if $x \leq y$ and $x \rightarrow_{\Pi} y = y/x$ otherwise. Thus, if $P(\psi)$ is not zero, the conditional probability $P(\varphi \mid \psi)$ can be written in $\text{FP}(\mathbb{L}\Pi)$ as $P(\psi) \rightarrow_{\Pi} P(\varphi \wedge \psi)$ and hence interpreted in its semantics as $P(\varphi \wedge \psi)/P(\psi)$. A related approach can be found in [9], where Popper conditional probabilities are formalised in a similar setting.

In this paper we propose a fuzzy modal logic $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ for conditional probability that extends $\text{FP}(\mathbb{L}\Pi)$ in the expressive power. In particular, $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ formalizes conditional events by the recently defined logic LBC (Logic of Boolean Conditionals) for conditional events. The latter allows to represent not only basic conditional expressions “ φ given ψ ”, written $(\varphi \mid \psi)$, but also compound conditional sentences such as “ φ given ψ and γ given χ ”, written in LBC as $(\varphi \mid \psi) \wedge (\gamma \mid \chi)$, or more in general, any Boolean combination of basic ones [5]. For each of such (basic and compound) conditional sentences, $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ permits to represent and reason about their probability. Thus, the conditional probability of “ φ given ψ ” is treated in $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ as the *unconditional* probability of the basic conditional formula $(\varphi \mid \psi)$.

In addition, we will present extensions of $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ that capture a more refined notion of probability functions. For $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ and each of its extensions, we prove soundness and completeness results w.r.t. suitable classes of probability models.

This paper is organized as follows. Section II gathers extensive preliminaries: on the Logic for Boolean Conditionals (LBC) in Subsection II-A; on the ground propositional logic $\mathbb{L}\Pi$ in Subsection II-B; and on the fuzzy modal logic for conditional probability $\text{FP}(\mathbb{L}\Pi)$ in Subsection II-C. In Section

III we will define the probability logic for compound conditionals FP(LBC, $\mathbb{L}\Pi$). In the same section we will consider the class of *separable* models and prove completeness. Moreover, we will show that, for every basic conditional $(\varphi \mid \psi)$ such that ψ has positive probability, the logic FP(LBC, $\mathbb{L}\Pi$) proves that the modal formula $P(\varphi \mid \psi)$ is equivalent to $P(\psi) \rightarrow_{\Pi} P(\varphi \wedge \psi)$ and hence, in every separable model, $P(\varphi \mid \psi)$ is interpreted as the ratio $\mu(\varphi \wedge \psi)/\mu(\psi)$. A first extension of FP(LBC, $\mathbb{L}\Pi$), namely the logic FP(LBC, $\mathbb{L}\Pi$)⁺, will be defined in Section IV and there we will show it to be complete w.r.t. to a subclass of separable models called *positive separable* models. Section V deals with the logic FP(LBC, $\mathbb{L}\Pi$)_c⁺, a further extension meant to capture the behavior of the so-called *canonical* extensions to $\mathcal{C}(\mathbf{A})$ of positive (unconditional) probabilities on $\mathbf{A}$ in [5]. For FP(LBC, $\mathbb{L}\Pi$)_c⁺ we prove soundness and completeness w.r.t. to the proper subclass of positive separable models that will be called *canonical* in that section. Conclusions and future work on this subject will be discussed in the final Section VI.

II. PRELIMINARIES

A. The logic LBC

In this section we recall from [5] the Logic of Boolean Conditionals (LBC). The idea is to consider basic conditional formulas of the form $(\varphi \mid \psi)$ as primitive objects that can be freely combined with Boolean connectives. A difference with the so-called *measure-free conditionals* is that the combination of two basic conditionals need not be another basic conditional, only in some special cases specified in the axioms of the logic. Indeed, formulas of LBC correspond to Boolean combinations of basic conditional formulas $(\varphi \mid \psi)$, where φ, ψ are classical propositions. In more detail, let $\mathbf{L}$ be a propositional language built from a finite set of propositional variables $p_1, p_2, \dots, p_k$ and classical logic connectives $\wedge, \vee, \neg, \rightarrow, \leftrightarrow$. We will denote by $\vdash_{CPL}$ derivability in Classical Propositional Logic. Based on $\mathbf{L}$, we define the language LBC of conditionals by the following stipulations:

- Basic (or atomic) conditional formulas, expressions of the form $(\varphi \mid \psi)$ where $\varphi, \psi \in \mathbf{L}$ and such that $\not\vdash_{CPL} \neg\psi$, are in LBC.
- Further, if $\Phi, \Psi \in \text{LBC}$, then $\neg\Phi, \Phi \wedge \Psi \in \text{LBC}$.¹ Other connectives like $\vee, \rightarrow$ and $\leftrightarrow$ are defined as usual.

Note that we do not allow the nesting of conditionals, as usually done in the vast literature on the modal approaches to Conditional Logics. Actually, purely propositional formulas from $\mathbf{L}$ can also be considered to be part of LBC since, as a matter of fact, any proposition φ can be identified with the conditional $(\varphi \mid \top)$, where $\top$ is an abbreviation for $\psi \vee \neg\psi$.

Definition 2.1: The *Logic of Boolean conditionals* (LBC for short) has the following axioms:

¹We use the same symbols for connectives in $\mathbf{L}$ and in LBC without danger of confusion.

For any tautology of CPL, the formula resulting from a uniform replacement of the variables by basic conditionals.

- (A1) $(\psi \mid \psi)$
- (A2) $\neg(\varphi \mid \psi) \leftrightarrow (\neg\varphi \mid \psi)$
- (A3) $(\varphi \mid \psi) \wedge (\delta \mid \psi) \leftrightarrow (\varphi \wedge \delta \mid \psi)$
- (A4) $(\varphi \mid \psi) \leftrightarrow (\varphi \wedge \psi \mid \psi)$
- (A5) $(\varphi \mid \psi) \leftrightarrow (\varphi \mid \chi) \wedge (\chi \mid \psi)$, if $\vdash_{CPL} \varphi \rightarrow \chi$ and $\vdash_{CPL} \chi \rightarrow \psi$
- (R1) from $\vdash_{CPL} \varphi \rightarrow \psi$ derive $(\varphi \mid \chi) \rightarrow (\psi \mid \chi)$
- (R2) from $\vdash_{CPL} \chi \leftrightarrow \psi$ derive $(\varphi \mid \chi) \leftrightarrow (\varphi \mid \psi)$
- (MP) Modus Ponens: from Φ and $\Phi \rightarrow \Psi$ derive Ψ

The notion of proof in LBC, $\vdash_{LBC}$, is defined as usual from the above axioms and rules.

In [5] it is shown that the Lindenbaum algebra corresponding to LBC, that is, the algebra of LBC-formulas modulo logical equivalence, is a certain type of Boolean algebra, called Boolean algebra of conditionals, which is finite if the set propositional variables is so, that is our case. Then, the algebra is atomic and its set of atoms are conjunctions of basic conditionals of length $n-1$, where $n = 2^m$ with m being the number of propositional variables, of the following form:

$$(\alpha_1 \mid \top) \wedge (\alpha_2 \mid \neg\alpha_1) \wedge \dots \wedge (\alpha_{n-1} \mid \bigwedge_{i=1, n-2} \alpha_i),$$

where $\alpha_1, \dots, \alpha_{n-1}$ are propositional atoms of the Lindenbaum algebra of the underlying propositional language $\mathbf{L}$.

The semantics of the LBC logic is based on sequences $\bar{w} = \langle w_1, \dots, w_n \rangle$ of pairwise different propositional interpretations for the underlying language $\mathbf{L}$, in such a way that $\bar{w}$ satisfies a conditional $(\varphi \mid \psi)$, written $\bar{w} \models_{LBC} (\varphi \mid \psi)$, if $w_i \models_{CPL} \varphi$ for the lowest index i such that $w_i \models \psi$.

B. The logic $\mathbb{L}\Pi$

The $\mathbb{L}\Pi$ is a powerful fuzzy logic system that suitably combines the connectives from Łukasiewicz logic with the connectives of Product fuzzy logic [4]. The language of the $\mathbb{L}\Pi$ logic is built in the usual way from a countable set of propositional variables, three binary connectives $\rightarrow_L$ (Łukasiewicz implication), $\odot$ (Product conjunction) and $\rightarrow_{\Pi}$ (Product implication), and the truth constant $\bar{0}$. A truth-evaluation is a mapping e that assigns to every propositional variable a real number from the unit interval $[0, 1]$ and extends to all formulas as follows:

$$\begin{aligned} e(\bar{0}) &= 0, & e(\varphi \odot \psi) &= e(\varphi) \cdot e(\psi) \\ e(\varphi \rightarrow_L \psi) &= \min(1 - e(\varphi) + e(\psi), 1), \\ e(\varphi \rightarrow_{\Pi} \psi) &= \begin{cases} 1, & \text{if } e(\varphi) \leq e(\psi) \\ e(\psi)/e(\varphi), & \text{otherwise} \end{cases} \end{aligned}$$

The truth constant $\bar{1}$ is defined as $\varphi \rightarrow_L \varphi$. In this way we have $e(\bar{1}) = 1$ for any truth-evaluation e . Moreover, many other connectives can be defined from those introduced above:

$$\begin{aligned} \neg_L \varphi &\text{ is } \varphi \rightarrow_L \bar{0} \\ \neg_{\Pi} \varphi &\text{ is } \varphi \rightarrow_{\Pi} \bar{0}, \\ \varphi \wedge \psi &\text{ is } \varphi \&(\varphi \rightarrow_L \psi), \\ \varphi \vee \psi &\text{ is } \neg_L(\neg_L \varphi \wedge \neg_L \psi), \end{aligned}$$

$\varphi \oplus \psi$	is	$\neg_L \varphi \rightarrow_L \psi$,
$\varphi \otimes \psi$	is	$\neg_L (\neg_L \varphi \oplus \neg_L \psi)$,
$\varphi \ominus \psi$	is	$\varphi \& \neg_L \psi$,
$\varphi \leftrightarrow_L \psi$	is	$(\varphi \rightarrow_L \psi) \& (\psi \rightarrow_L \varphi)$,
$\Delta \varphi$	is	$\neg_{\Pi} \neg_L \varphi$,
$\nabla \varphi$	is	$\neg_{\Pi} \neg_{\Pi} \varphi$,

with the following interpretations:

$$\begin{aligned}
e(\neg_L \varphi) &= 1 - e(\varphi), \\
e(\neg_{\Pi} \varphi) &= \begin{cases} 1, & \text{if } e(\varphi) = 0 \\ 0, & \text{otherwise} \end{cases}, \\
e(\varphi \wedge \psi) &= \min(e(\varphi), e(\psi)), \\
e(\varphi \vee \psi) &= \max(e(\varphi), e(\psi)), \\
e(\varphi \oplus \psi) &= \min(1, e(\varphi) + e(\psi)), \\
e(\varphi \otimes \psi) &= \max(0, e(\varphi) + e(\psi) - 1), \\
e(\varphi \ominus \psi) &= \max(0, e(\varphi) - e(\psi)), \\
e(\varphi \leftrightarrow_L \psi) &= 1 - |e(\varphi) - e(\psi)|, \\
e(\Delta \varphi) &= \begin{cases} 1, & \text{if } e(\varphi) = 1 \\ 0, & \text{otherwise} \end{cases}, \\
e(\nabla \varphi) &= \begin{cases} 1, & \text{if } e(\varphi) > 0 \\ 0, & \text{otherwise} \end{cases}.
\end{aligned}$$

The logic $\mathbb{L}\Pi$ is defined Hilbert-style as the logical system whose axioms and rules are the following:²

- (i) Axioms of Łukasiewicz Logic:
- (ii) Axioms of Product Logic
- (iii) The following additional axioms relating Łukasiewicz and Product logic connectives:
 - ($\neg$) $\neg_{\Pi} \varphi \rightarrow_L \neg_L \varphi$
 - (Δ) $\Delta(\varphi \rightarrow_L \psi) \equiv \Delta(\varphi \rightarrow_{\Pi} \psi)$
 - ($L\Pi$) $\varphi \odot (\psi \odot \chi) \equiv (\varphi \odot \psi) \odot (\varphi \odot \chi)$
- (iv) Deduction rules of $\mathbb{L}\Pi$ are modus ponens for $\rightarrow_L$ (modus ponens for $\rightarrow_{\Pi}$ is derivable), and necessitation for Δ : from φ derive $\Delta \varphi$.

C. The probability logic $FP(\mathbb{L}\Pi)$

This final section on preliminaries is on the probability logic $FP(\mathbb{L}\Pi)$ (FP for Fuzzy Probability), introduced in [8] and defined as a sort of modal extension with a unary operator $P(\cdot)$ over the fuzzy logic $\mathbb{L}\Pi$ described in the previous section. This logic allows for reasoning about the probability of classical propositions.

Actually, the propositional language L is extended by a fuzzy unary modal operator P . If φ is a proposition of L , then $P\varphi$ is a modal proposition whose intended reading is that “ φ is *probable*”, and whose truth-degree will be taken as the probability of φ .

The language of $FP(\mathbb{L}\Pi)$ is defined as follows. *Formulas* of $FP(\mathbb{L}\Pi)$ are of two types:

- Non-modal: they are exactly the (classical) formulas of L , i.e. those built from a set Var of propositional variables $\{p_1, p_2, \dots, p_n, \dots\}$ using the classical binary connectives $\wedge$ and $\neg$. Other connectives like $\vee$ and $\rightarrow$ are defined from $\wedge$ and $\neg$ in the usual way. We shall denote them by lower case Greek letters φ, ψ , etc.

- Modal: they are built from elementary modal formulas of the form $P\varphi$, where φ is a non-modal formula, using the connectives of $\mathbb{L}\Pi$ ($\rightarrow_L, \odot, \rightarrow_{\Pi}$). We shall denote them by upper case Greek letters Φ, Ψ , etc.

These are all the formulas of $FP(\mathbb{L}\Pi)$. Notice that nested modalities, among other things, are not allowed.

Axioms and rules of $FP(LBC, \mathbb{L}\Pi)$ are as follows:

- (CPL) All axioms and rules of classical propositional logic restricted to classical, non-modal, formulas;
- ($\mathbb{L}\Pi$) All axioms and rules of $\mathbb{L}\Pi$ for modal formulas;
- (P) The following axioms and rules for the modality P : for all propositions $\varphi, \psi \in L$,
 - (P1) $P(\varphi \rightarrow \psi) \rightarrow_L (P(\varphi) \rightarrow_L P(\psi))$
 - (P2) $P(\neg \varphi) \leftrightarrow_L \neg P(\varphi)$
 - (P3) $P(\varphi \vee \psi) \leftrightarrow_L [P(\varphi) \oplus (P(\psi) \ominus P(\varphi \wedge \psi))]$
 - (Nec) if $\vdash_{CPL} \varphi$, derive $P(\varphi)$.

Models of $FP(\mathbb{L}\Pi)$ are probability Kripke structures $K = (W, e, \mu)$, where:

- W is a non-empty set of possible worlds;
- $e : V \times W \rightarrow \{0, 1\}$ provides for each world a *Boolean* (two-valued) evaluation of the proposition variables, that is, $e(p, w) \in \{0, 1\}$ for each propositional variable $p \in Var$ and each world $w \in W$; and
- $\mu : 2^W \rightarrow [0, 1]$ is a finitely additive probability measure on a Boolean algebra of subsets of W such that for each p , the set $\{w \mid e(p, w) = 1\}$ is measurable (cf. [10] 8.4.1).

A truth evaluation e is extended to non-modal formulas in the classical way, to elementary modal formulas as follows:

$$e(P\varphi, w) = \mu(\{w \in W \mid e(\varphi, w) = 1\}),$$

and to compound modal formulas by using the truth-functions of the $\mathbb{L}\Pi$ logic.

Soundness and completeness of the logic $FP(\mathbb{L}\Pi)$ w.r.t. to the class of probability Kripke models is proved in [8]: if $T \cup \{\Phi\}$ is a finite set of $FP(\mathbb{L}\Pi)$ -formulas, then T proves Φ in $FP(\mathbb{L}\Pi)$ iff for any probability Kripke model $K = (W, e, \mu)$ and any world $w \in W$, $e(\Phi, w) = 1$ whenever $e(\Psi, w) = 1$ for all $\Psi \in T$.

III. PROBABILITY LOGIC OVER CONDITIONALS

In this section we define a logic to reason about the probability of basic and compound conditionals over the fuzzy logic $\mathbb{L}\Pi$. In the same line as with the logic $FP(\mathbb{L}\Pi)$ described in Section II-C, we extend the language of LBC with a fuzzy (unary) modal operator P , so that, for every basic conditional $(\varphi \mid \psi)$, the intended meaning of a formula $P(\varphi \mid \psi)$ is that the conditional “ φ given ψ ” is *probable*, and that the truth-degree of $P(\varphi \mid \psi)$ is the probability of the conditional “ $(\varphi \mid \psi)$ ”. The relation of this probability to the usual notion of conditional probability of φ given ψ will become clear later.

The logic $FP(LBC, \mathbb{L}\Pi)$ is obtained by replacing, in the definition of $FP(\mathbb{L}\Pi)$, classical logic for events by the conditional logic LBC defined as in Section II-A. Formulas of $FP(LBC, \mathbb{L}\Pi)$ are of two types:

²This definition, proposed in [3], is actually a simplified version of the original definition of $L\Pi$ given in [4].

- **Conditional formulas** are formulas of the logic LBC, that is, basic conditionals of the form $(\varphi \mid \psi)$ for all classical formulas φ and ψ such that ψ is not a classical logic contradiction (in other words, $\not\vdash_{CPL} \neg\psi$) and compound conditional formulas obtained as Boolean combinations of basic ones. Compound conditional formulas will be denoted as $\Phi, \Psi, \dots$;

- **Modal formulas**: for every (basic or compound) conditional formula Φ , $P(\Phi)$ is an atomic modal formula. Compound modal formulas are combinations of atomic ones by means of the $\mathbb{L}\Pi$ connectives.

Thus, for instance, $P(\varphi \mid \psi)$, $P((\varphi \mid \psi) \wedge (\gamma \mid \delta))$ and $P((\varphi \mid \psi) \wedge (\gamma \mid \delta)) \rightarrow_L P(\chi \mid \tau)$ are compound modal formulas for all classical formulas $\varphi, \psi, \gamma, \delta, \chi, \tau$ such that $\not\vdash \neg\psi, \not\vdash \neg\delta, \not\vdash \neg\tau$. However, neither $(\varphi \mid \psi) \rightarrow_L P(\gamma \mid \delta)$ nor $P((\varphi \mid \psi) \oplus P(\chi \mid \tau))$ are well formed formulas in this language.

Axioms and rules of $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ are as follows:

- (LBC) All axioms and rules of LBC restricted to conditional formulas;
- ($\mathbb{L}\Pi$) All axioms and rules of $\mathbb{L}\Pi$ for modal formulas;
- (P) The axioms and rules for the modality P are those for $\text{FP}(\mathbb{L}\Pi)$, but now for all conditional formulas $\Phi, \Psi \in \text{LBC}$, plus a new rule (Sep):
 - (P1) $P(\Phi \rightarrow \Psi) \rightarrow_L (P(\Phi) \rightarrow_L P(\Psi))$;
 - (P2) $P(\neg\Phi) \leftrightarrow_L \neg P(\Phi)$;
 - (P3) $P(\Phi \vee \Psi) \leftrightarrow_L [P(\Phi) \oplus (P(\Psi) \ominus P(\Phi \wedge \Psi))]$;
 - (Nec) if $\vdash_{\text{LBC}} \Phi$, derive $P(\Phi)$;
 - (Sep) if $\vdash_{CPL} (\varphi \rightarrow \chi) \wedge (\chi \rightarrow \psi)$, derive $P((\varphi \mid \chi) \wedge (\chi \mid \psi)) \leftrightarrow_L P(\varphi \mid \chi) \odot P(\chi \mid \psi)$.

The notion of proof according to these axioms and rules will be denoted $\vdash_{\text{FP}}$. The axioms and rules of $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ are meant to capture the behavior of an unconditional, *separable* probability measure on the Lindenbaum algebra of the logic LBC. In fact, let us recall from [5] that, in particular, a probability μ on a Boolean algebra of conditionals $\mathcal{C}(\mathbf{A})$ is *separable* if for all $a, b, c \in A \setminus \{\perp\}$ such that $a \leq b \leq c$, then $\mu((a \mid b) \wedge (b \mid c)) = \mu(a \mid b) \cdot \mu(b \mid c)$. As we will show later on, separability is captured by the rule (Sep) above.

In what follows Ω will denote the set of Boolean interpretations for the variables Var , and $\text{Seq}_n(\Omega)$ will denote the set of sequences of n pairwise different interpretations from Ω .

Definition 3.1: A probability LBC-Kripke model is a structure $\mathcal{C} = \langle W, e, \mu \rangle$ where

- W is a set of worlds;
- $e : W \rightarrow \text{Seq}_n(\Omega)$ maps every world $w \in W$ to a LBC-evaluation $e(w) = \bar{w} \in \text{Seq}_n(\Omega)$;
- μ is a probability on $2^{e[W]}$, where $e[W] = \{e(w) \mid w \in W\} \subseteq 2^{\text{Seq}_n(\Omega)}$.

Each LBC-Kripke model $\mathcal{C} = \langle W, e, \mu \rangle$ induces a probability on LBC formulas in the natural way, in particular for each basic conditional $(\varphi \mid \psi)$ we define:

$$\mu((\varphi \mid \psi)) = \mu(\{\bar{w} \in \text{Seq}_n(\Omega) \mid w \in W, \bar{w} \models (\varphi \mid \psi)\}),^3$$

³Without danger of confusion, we will write $\mu(\varphi \mid \psi)$ for $\mu((\varphi \mid \psi))$.

and similarly for every compound conditional. Notice that, when $e[W] = \text{Seq}_n(\Omega)$, the Boolean algebra $2^{e[W]}$ on which μ is defined, actually is the Lindenbaum algebra of LBC. It was proved in [5, Theorem 7.3] that such Lindenbaum algebra is isomorphic to $\mathcal{C}(\mathbf{L})$, that is, the conditional algebra generated by $\mathbf{L}$, the Lindenbaum algebra of classical propositional logic. Thus, every LBC-Kripke model determines a probability measure μ on $\mathcal{C}(\mathbf{L})$.

Definition 3.2: A probability LBC-Kripke model $\mathcal{C} = \langle W, e, \mu \rangle$ is called *separable* when

$$\mu((\varphi \mid \psi) \wedge (\psi \mid \chi)) = \mu(\varphi \mid \psi) \cdot \mu(\psi \mid \chi) \quad (1)$$

for every φ, ψ, χ such that $\vdash_{CPL} (\varphi \rightarrow \psi) \wedge (\psi \rightarrow \chi)$ and $\not\vdash_{CPL} \neg\psi$.

An immediate consequence of the definition above is that every μ of a separable model, satisfies $\mu(\varphi \mid \top) = \mu((\varphi \mid \psi) \wedge (\psi \mid \top)) = \mu(\varphi \mid \psi) \cdot \mu(\psi \mid \top)$ for all φ, ψ such that $\vdash_{CPL} \varphi \rightarrow \psi$ and $\not\vdash \neg\psi$. Note that the first equality is due to Axiom (A5) of LBC.

Given a formula F of $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ and a separable model $\mathcal{S} = \langle W, e, \mu \rangle$, the evaluation of F in $\mathcal{S}$ at $w \in W$ is inductively defined by the following stipulations:

- If F is a conditional formula Φ , then $\|F\|_{\mathcal{S}, w} = w_1(\Phi) \in \{0, 1\}$, where $e(w) = \langle w_1, w_2, \dots \rangle$;
- If $F = P(\Phi)$ is atomic modal, then $\|P(\Phi)\|_{\mathcal{S}, w} = \mu(\Phi)$;
- If F is compound modal, then $\|F\|_{\mathcal{S}, w}$ is computed by evaluating its atomic components and then by using the truth-functionality of the connectives of $\mathbb{L}\Pi$ in the standard algebra $[0, 1]$.

Notice that if F is modal, then $\|F\|_{\mathcal{S}, w}$ does not depend on the world w . Finally, the truth-degree of F in $\mathcal{C}$ is defined as $\|F\|_{\mathcal{S}} = \inf_{w \in W} \|F\|_{\mathcal{S}, w}$.

Definition 3.3: If $T \cup \Phi$ is a set of modal formulas, Φ logically follows from T , written $T \models_{\text{SFP}} \Phi$, when for all separable probability LBC-Kripke model $\mathcal{S}$, if $\|F\|_{\mathcal{S}} = 1$ for every $F \in T$, then $\|\Phi\|_{\mathcal{S}} = 1$ as well.

Next, we prove that $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ is sound and complete with respect to the class of separable models. Its proof, a main part of which will follow from a general result we will recall below, is based on the fact that the logic for events, LBC is locally finite. This means that the Lindenbaum algebra $\mathbf{CL}$ over a finite set of variables is finite. Indeed, as proved in [5, Theorem 7.3], $\mathbf{CL}$ is isomorphic to $\mathcal{C}(\mathbf{L})$ to the Boolean algebra of conditionals of the Lindenbaum algebra of classical logic, on the same set of variables.

Theorem 3.1: The logic $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ is sound and complete for deductions from probabilistic modal formulas w.r.t. to the class of separable models.

Proof: Soundness of (P1), (P2), (P3) and (Nec) follows directly from [10, Lemma 8.4.5.]. Thus, let us show that (Sep) holds in every separable model. If $(\varphi \rightarrow \chi) \wedge (\chi \rightarrow \psi)$ is a theorem of classical logic and ψ is not a contradiction, then $\|P((\varphi \mid \chi) \wedge (\chi \mid \psi))\|_{\mathcal{S}} = \|P(\varphi \mid \chi)\|_{\mathcal{S}} \cdot \|P(\chi \mid \psi)\|_{\mathcal{S}}$ in any separable model $\mathcal{S} = \langle W, e, \mu \rangle$, because μ satisfies Equation (1) above, and hence $\mathcal{S}$ satisfies $P((\varphi \mid \chi) \wedge (\chi \mid \psi)) \leftrightarrow_L P(\varphi \mid \chi) \odot P(\chi \mid \psi)$.

As for completeness, we will show that $\not\models_{FP} F$ implies $\not\models_{SFP} F$, for any modal formula F . By either adapting the completeness proofs in [8], [11] or adapting the general result proved in [6, Theorem 20 (1)], together with the fact that LBC is locally finite and $\mathbb{L}\Pi$ is finitely strong standard complete, we can show that deductions in $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ can be translated to deductions in $\mathbb{L}\Pi$ by considering atomic modal formulas $P\Phi$ as new $\mathbb{L}\Pi$ propositional variables p_Φ . Indeed, it holds that, for any modal formula G , $\vdash_{FP} G$ iff $T \vdash_{\mathbb{L}\Pi} G^*$, where G^* is the translation of G with the new variables, and T consists of the following three sets of formulas:

- (i) $T_0 = \{H^* : H \text{ is an instance of axioms (P1), (P2), (P3)}\}$
- (ii) $T_1 = \{p_\Psi : \Psi \text{ is an LBC-theorem}\}$, that translates the rule (Nec), and
- (iii) $T_2 = \{p_{(\varphi|\chi) \wedge (\chi|\psi)} \leftrightarrow_L p_{\varphi|\chi} \odot p_{\chi|\psi} : \vdash_{CPL} (\varphi \rightarrow \chi) \wedge (\chi \rightarrow \psi) \text{ and } \not\vdash_{CPL} \neg\chi\}$, that translates the rule (Sep).

Then by the finite-strong completeness of $\mathbb{L}\Pi$,⁴ if F is not a theorem of $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$, there is a $\mathbb{L}\Pi$ -evaluation v , model of the sets T_0, T_1, T_2 and $v(F) < 1$. Then one can define the probability Kripke model $\mathcal{S} = (W, e, \mu)$, where $W = \text{Seq}_n(\Omega)$, $e(w, p) = \bar{w}_1(p)$ for any propositional variable p , and $\mu(\Phi) = v(\Phi)$ for any conditional formula Φ , and show that $\|\mathcal{F}\|_{\mathcal{S}} = v(F) < 1$, that is, $\mathcal{S}$ is a countermodel of F . Thus, it is left to prove that μ is separable. Indeed, in particular v is a model of T_2 , that means $v(p_{\varphi|\chi} \odot p_{\chi|\psi}) = v(p_{\varphi|\chi}) \cdot v(p_{\chi|\psi}) = \mu(\varphi | \chi) \cdot \mu(\chi | \psi)$, for all those conditionals $(\varphi | \chi), (\chi | \psi)$ such that $\vdash_{CPL} (\varphi \rightarrow \chi) \wedge (\chi \rightarrow \psi)$ and $\not\vdash_{CPL} \neg\chi$. Thus, μ is separable and the claim is settled. ■

We end this section by noticing that in a separable LBC-Kripke model, a formula $P(\varphi | \psi)$ is evaluated by its corresponding conditional probability.

Corollary 3.1: For every basic conditional $(\varphi | \psi)$, the following deduction holds in $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$:⁵

$$\nabla P(\psi | \top) \vdash_{FP} P(\varphi | \psi) \leftrightarrow_L (P(\psi | \top) \rightarrow_{\Pi} P(\varphi \wedge \psi | \top))$$

IV. POSITIVE SEPARABLE MODELS

In this section we will consider a first extension of the logic $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ that allows to deal with positive probabilities on basic conditionals of the form $(\varphi | \top)$.

Definition 4.1: The logic $\text{FP}(\text{LBC}, \mathbb{L}\Pi)^+$ is the schematic extension of $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$ obtained by adding the rule

(Pos) if $\not\vdash_{CPL} \neg\varphi$, derive $\nabla P(\varphi | \top)$.

The effect of axiom (Pos) is to force the probability of non-contradictory classical propositions φ (once identified as conditionals $(\varphi | \top)$) to be strictly positive. Therefore, it is relatively easy to see that the following holds.

Theorem 4.1: The logic $\text{FP}(\text{LBC}, \mathbb{L}\Pi)^+$ is sound and complete w.r.t. the class of *positive* separable LBC-Kripke models, i.e. models $\mathcal{S} = \langle W, e, \mu \rangle$ in which μ is a positive probability

⁴For our purposes, T can be considered to be finite, because the Lindenbaum Boolean algebra of LBC is finite, that is, there are only finitely-many non-logically equivalent basic conditionals $(\varphi | \psi)$, and hence only finitely-many non-logically equivalent formulas in T .

⁵Recall from Section II-B that the interpretation of the connective ∇ is $\nabla(x) = 1$ if $x > 0$, and $\nabla(x) = 0$ otherwise.

measure, that is to say, such that $\mu(\Phi) > 0$ for all conditional formula $\Phi \neq \perp$.

Let us call *basic modal formulas* any combination of atomic modal formulas of the form $P(\varphi_i | \psi_i)$ with $\mathbb{L}\Pi$ connectives. If we restrict ourselves to this sublanguage of $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$, we can in fact consider simpler probabilistic models.

Definition 4.2: A *positive simple model* is a pair $\mathcal{P} = \langle \Omega, \sigma \rangle$ where Ω is the set of Boolean interpretations for the base language $\mathbb{L}$, and σ is a positive probability measure on 2^Ω .

Given a basic modal formula B and a positive model $\mathcal{P} = \langle \Omega, \sigma \rangle$, we interpret B in $\mathcal{P}$ as follows:

- If $B = P(\varphi | \psi)$, then $\|P(\varphi | \psi)\|_{\mathcal{P}} = \frac{\sigma(\varphi \wedge \psi)}{\sigma(\psi)}$;
- If B is compound use again the truth functionality of $\mathbb{L}\Pi$ connectives interpreted in $[0, 1]$.

Theorem 4.2: (1) For every positive separable LBC-Kripke model $\mathcal{S}$ there exists a positive simple model $\mathcal{P}$ such that $\|B\|_{\mathcal{S}} = \|B\|_{\mathcal{P}}$ for every basic modal formula B .

(2) Vice-versa, for every positive simple model $\mathcal{P}$ there exists a positive separable LBC-Kripke model $\mathcal{S}$ such that $\|B\|_{\mathcal{P}} = \|B\|_{\mathcal{S}}$ for every basic modal formula B .

Proof: As for (1), let us prove the claim for $B = (\varphi | \psi)$. The case of compound conditional formulas, indeed, follows by truth-functionality of the connectives of $\mathbb{L}\Pi$. Given a positive separable LBC-Kripke model $\mathcal{S} = \langle W, e, \mu \rangle$, define $\sigma(\varphi) = \mu(\varphi | \top) = \mu(\{w \in W \mid \bar{w} \models (\varphi | \top)\})$. This is a probability on Boolean formulas that can be identified as a probability on 2^Ω . Since μ is positive and separable, we have

$$\mu(\varphi | \psi) = \frac{\mu(\varphi \wedge \psi | \top)}{\mu(\psi | \top)} = \frac{\sigma(\varphi \wedge \psi)}{\sigma(\psi)}.$$

On the other hand, (2) follows by adapting to our logical setting a main result in [5, Theorem 6.13] stated in algebraic terms. Indeed, in that theorem it is proved that, for any positive probability P on an algebra of events $\mathbf{A}$, there is a (plain) probability μ_P on the algebra of conditional events $\mathcal{C}(\mathbf{A})$ such that $\mu_P(a | b) = \frac{\mu(a \wedge b)}{\mu(b)}$ whenever $b \neq \perp$. The proof is rather involved and we refer the reader to [5] for full details. ■

V. A LOGIC FOR (CONDITIONAL) CANONICAL EXTENSIONS

As proved in [5] the atoms of a Boolean algebra of conditionals $\mathcal{C}(\mathbf{A})$ can be fully characterized by the atoms of the original algebra $\mathbf{A}$ and, in particular, if $\alpha_1, \dots, \alpha_n$ are the atoms of $\mathbf{A}$, those of $\mathcal{C}(\mathbf{A})$ are conditional expressions of the form

$$\omega_i = (\alpha_{i_1} | \top) \wedge (\alpha_{i_2} | \neg\alpha_{i_1}) \wedge \dots \wedge (\alpha_{i_{n-1}} | \bigwedge_{j \leq n-2} \neg\alpha_{i_j}).$$

Since the atoms of the Lindenbaum algebra of classical logic (with, say, k variables $x_1, \dots, x_k$) are writable as minterms $\alpha_j = \bigwedge x_i^*$, for $x_i^* \in \{x_i, \neg x_i\}$, the atoms of $\mathcal{C}(\mathbf{L})$ are expressible as above. To ease the reading, we will denote them by $\omega_1, \omega_2, \dots$. Recall from [5] that, if classical logic is defined on k propositional variables, there are $(2^k)!$ atoms of $\mathcal{C}(\mathbf{L})$.

Although not every probability measure on $\mathcal{C}(\mathbf{A})$ satisfies all the axioms of a conditional probability, every positive probability σ on the original algebra $\mathbf{A}$, has an extension to

a positive probability μ_σ on $\mathcal{C}(\mathbf{A})$. These measures, called *canonical* in [5] are such that, for every atom ω_i of $\mathcal{C}(\mathbf{L})$,

$$\mu_\sigma(\omega_i) = \sigma(\alpha_{i_1} \mid \top) \cdot \sigma(\alpha_{i_2} \mid \neg\alpha_{i_1}) \cdots \sigma(\alpha_{i_{n-1}} \mid \bigwedge_{j \leq n-2} \neg\alpha_{i_j}).$$

In this section we will show how to further extend the logic $\text{FP}(\text{LBC}, \mathbb{L}\Pi)^+$ in order for its models to be defined by canonical extensions μ_σ of this kind. In order to do that, let us consider the following $\text{FP}(\text{LBC}, \mathbb{L}\Pi)^+$ -formulas:

$$(\text{Can}_i) \quad P(\omega_i) \leftrightarrow_L P(\alpha_{i_1} \mid \top) \odot \dots \odot P(\alpha_{i_{n-1}} \mid \bigwedge_{j \leq n-2} \neg\alpha_{i_j}),$$

where $\omega_i = (\alpha_{i_1} \mid \top) \wedge \dots \wedge (\alpha_{i_{n-1}} \mid \bigwedge_{j \leq n-2} \neg\alpha_{i_j})$, with $\alpha_{i_1}, \dots, \alpha_{i_n}$ being minterms of the propositional language $\mathbf{L}$.

Definition 5.1: Let $\mathbf{L}$ be a propositional language with k variables. Then, the logic $\text{FP}(\text{LBC}, \mathbb{L}\Pi)_c^+$ is the schematic extension of $\text{FP}(\text{LBC}, \mathbb{L}\Pi)^+$ obtained by adding the axioms (Can_i) for all $i = 1, \dots, (2^k)!$

A separable model $\mathcal{S} = \langle W, e, \mu \rangle$ is *canonical* if there exists a positive probability σ on Ω such that $\mu = \mu_\sigma$, i.e., μ is the canonical extension of some positive σ on Ω .

Finally, we can prove that $\text{FP}(\text{LBC}, \mathbb{L}\Pi)_c^+$ is sound and complete w.r.t. to canonical models.

Theorem 5.1: The logic $\text{FP}(\text{LBC}, \mathbb{L}\Pi)_c^+$ is sound and complete with respect to the class of canonical models.

Proof: Following the lines we sketched in the proof of Theorem 3.1, it is enough to show that a positive separable model satisfies all the axioms (Can_i) iff the model is canonical.

(Left-to-Right). Let $\mathcal{S} = \langle W, e, \mu \rangle$ be positive separable and satisfying (Can_i) for all i . Thus, $\|P(\omega_i)\|_{\mathcal{S}} = \|P(\alpha_{i_1} \mid \top) \odot \dots \odot P(\alpha_{i_{n-1}} \mid \bigwedge_{j \leq i_{n-2}} \neg\alpha_j)\|_{\mathcal{S}} = \|P(\alpha_{i_1} \mid \top)\|_{\mathcal{S}} \cdot \dots \cdot \|P(\alpha_{i_{n-1}} \mid \bigwedge_{j \leq i_{n-2}} \neg\alpha_j)\|_{\mathcal{S}} = \mu(\alpha_{i_1} \mid \top) \cdot \dots \cdot \mu(\alpha_{i_{n-1}} \mid \bigwedge_{j \leq i_{n-2}} \neg\alpha_j).$

Since $\mathcal{S}$ is positive and separable, by Theorem 4.2 there is a positive simple model $\mathcal{P} = (\Omega, \sigma)$ such that, for every basic conditional $(\varphi \mid \psi)$, $\mu(\varphi \mid \psi) = \frac{\sigma(\varphi \wedge \psi)}{\sigma(\psi)}$. In particular, $\mu(\alpha_{i_1} \mid \top) \cdot \dots \cdot \mu(\alpha_{i_{n-1}} \mid \bigwedge_{j \leq i_{n-2}} \neg\alpha_j) = \frac{\alpha_{i_1}}{1} \cdot \frac{\sigma(\alpha_{i_2} \wedge \neg\alpha_{i_1})}{\sigma(\neg\alpha_{i_1})} \cdot \dots \cdot \frac{\sigma(\alpha_{i_{n-1}} \wedge \bigwedge_{j \leq i_{n-2}} \neg\alpha_j)}{\sigma(\bigwedge_{j \leq i_{n-2}} \neg\alpha_j)}$. Thus, $\mu = \mu_\sigma$ and $\mathcal{S}$ is canonical.

(Right-to-Left). Conversely, if $\mathcal{S}$ is canonical, then (Can_i) holds in $\mathcal{S}$ by the very definition of canonical model and the way formulas are interpreted in separable models. ■

In the light of the above argument, we can hence slightly improve the result of Corollary 3.1 as follows.

Corollary 5.1: The following formulas are theorems of $\text{FP}(\text{LBC}, \mathbb{L}\Pi)_c^+$:

- 1) $P(\varphi \mid \psi) \leftrightarrow_L (P(\psi \mid \top) \rightarrow_{\Pi} P(\varphi \wedge \psi \mid \top));$
- 2) $P(\omega) \leftrightarrow_L [P(\top) \rightarrow_{\Pi} P(\alpha_{i_1}) \odot P(\neg\alpha_{i_1}) \rightarrow_{\Pi} P(\alpha_{i_2}) \odot \dots \odot P(\bigwedge_{j \leq i_{n-2}} \neg\alpha_j) \rightarrow_{\Pi} P(\alpha_{i_{n-1}})],$ for $\omega = (\alpha_{i_1} \mid \top) \wedge (\alpha_{i_2} \mid \neg\alpha_{i_1}) \wedge \dots \wedge (\alpha_{i_{n-1}} \mid \bigwedge_{j \leq n-2} \neg\alpha_{i_j})$, where the α_{i_j} 's are pairwise different minterms of the propositional language $\mathbf{L}$.

VI. CONCLUSIONS

In this paper we have introduced fuzzy modal logics for reasoning about the probability of compound conditionals, the latter being Boolean combinations of basic conditionals

$(\varphi \mid \psi)$ and formalized within the recently introduced Logic for Boolean Conditionals LBC [5]. For each of the logics we define, we have proved completeness w.r.t. suitable classes of probability models, where a formula of the kind $P(\varphi \mid \psi)$ is evaluated by a (plain) probability $\mu((\varphi \mid \psi))$ of the conditional formula $(\varphi \mid \psi)$. We have shown that, if $\psi \not\vdash_{\text{CPL}} \perp$, this probability is in fact a conditional probability, and thus evaluated by the ratio $\mu((\varphi \wedge \psi \mid \top)) / \mu((\psi \mid \top))$.

There are a number of issues on this subject left for future work. Among them, we plan to investigate the relationship of our probability logics for compound conditionals with the approach developed by Sanfilippo et al. to probabilistic inference with conjoined and iterated conditionals based on a different notion of conditional, see e.g. [14], [15]. Another topic of interest is the application of these logics to reason with (semi-) fuzzy quantifiers [2]. In addition, we plan to investigate complexity bounds for the SAT problem for $\text{FP}(\text{LBC}, \mathbb{L}\Pi)$. Although it seems reasonable to conjecture that logic to be decidable, while the logics $\mathbb{L}\Pi$ and $\text{FP}(\mathbb{L}\Pi)$ are known to be in PSPACE [12], the complexity of the logic LBC is not known yet, and this latter non-trivial fact needs to be solved first.

Acknowledgments: The authors are grateful to anonymous reviewers for their comments. The authors acknowledge partial support by the Spanish project ISINC (PID2019-111544GB-C21). Flaminio also acknowledges partial support by the Spanish Ramón y Cajal research program RYC-2016-19799.

REFERENCES

- [1] P. Baldi, P. Cintula, C. Noguera. Classical and Fuzzy Two-Layered Modal Logics for Uncertainty: Translations and Proof-Theory. *Int. J. Comput. Intell. Syst.* 13(1): 988-1001, 2020.
- [2] P. Baldi, C.G. Fermüller. From Semi-fuzzy to Fuzzy Quantifiers via Łukasiewicz Logic and Games. *Proc. EUSFLAT/TWIFSGN*, Vol. 1, 112-124, 2017.
- [3] P. Cintula The $\mathbb{L}\Pi$ and $\mathbb{L}\Pi1/2$ propositional and predicate logics. *Fuzzy Sets and Systems*, 124, pp. 289-302, 2001.
- [4] F. Esteva, L. Godo and F. Montagna. The $\mathbb{L}\Pi$ and $\mathbb{L}\Pi1/2$ logics: two complete fuzzy logics joining Łukasiewicz and product logic. *Archive for Mathematical Logic*, 40, 39-67, 2001.
- [5] T. Flaminio, L. Godo, H. Hosni. Boolean algebras of conditionals, probability and logic. *Artificial Intelligence* 286, 103347, 2020.
- [6] T. Flaminio, L. Godo, E. Marchioni. Reasoning about Uncertainty of Fuzzy Events: an Overview. In *Understanding Vagueness - Logical, Philosophical, and Linguistic Perspectives*, P. Cintula et al. (Eds.), College Publications, pp. 367-400, 2011.
- [7] A. Gilio, G. Sanfilippo. Generalized logical operations among conditional events. *Appl. Intell.* 49(1): 79-102, 2019.
- [8] L. Godo, F. Esteva, P. Hájek. Reasoning about probability using fuzzy logic. *Neural Network World*, 10, No. 5, pp. 811-824, 2000.
- [9] L. Godo, E. Marchioni. Coherent Conditional Probability in a Fuzzy Logic Setting. *Log. J. IGPL* 14(3): 457-481, 2006.
- [10] P. Hájek. *Metamathematics of fuzzy logic*, Kluwer 1998.
- [11] P. Hájek, L. Godo, F. Esteva, Probability and Fuzzy Logic. In *Proc. of Uncertainty in Artificial Intelligence UAI'95*, 1995.
- [12] P. Hájek, S. Tulipani, Complexity of fuzzy probability logics. *Fundamenta Informaticae*, 45, 1-7, 2001.
- [13] J. Halpern. *Reasoning about Uncertainty*. MIT Press, 2003.
- [14] G. Sanfilippo, N. Pfeifer, D.E. Over, A. Gilio. Probabilistic inferences from conjoined to iterated conditionals. *Int. J. Approx. Reason.* 93: 103-118, 2018.
- [15] G. Sanfilippo, A. Gilio, D.E. Over, N. Pfeifer. Probabilities of conditionals and previsions of iterated conditionals. *Int. J. Approx. Reason.* 121: 150-173, 2020.

Empleo de contextos L -fuzzy en el diseño de materiales de origen renovable

Itsaso Leceta

Dep. Matemática Aplicada

Universidad del País Vasco - UPV/EHU

San Sebastián

itsaso.leceta@ehu.eus

Cristina Alcalde

Dep. Matemática Aplicada

Universidad del País Vasco - UPV/EHU

San Sebastián

c.alcalde@ehu.eus

Marta Urdanpilleta

Dep. Física Aplicada

Universidad del País Vasco - UPV/EHU

San Sebastián

marta.urdanpilleta@ehu.eus

Pedro Guerrero

Dep. Ing. Química y Medio Ambiente

Universidad del País Vasco - UPV/EHU

San Sebastián

pedromanuel.guerrero@ehu.eus

Koro de la Caba

Dep. Ing. Química y Medio Ambiente

Universidad del País Vasco - UPV/EHU

San Sebastián

koro.delacaba@ehu.eus

Ana Burusco

Dep. Estadística, Infor. y Matemáticas

Universidad Pública de Navarra

Instituto de Smart Cities

Pamplona

burusco@unavarra.es

Resumen—El objetivo principal de este trabajo es utilizar los contextos L -fuzzy para mejorar la eficiencia del diseño de materiales biodegradables de origen renovable fabricados mediante la valorización de residuos con buenas propiedades para su utilización en el sector de los envases alimentarios. El uso de los contextos L -fuzzy permitirá mejorar la eficacia en el proceso de diseño mejorando la eficiencia en la utilización de recursos. Además, mediante la utilización de los mencionados contextos L -fuzzy se pretende optimizar las propiedades de los materiales desarrollados y estimar la composición de los materiales para los requisitos que cumplir en servicio.

Palabras Clave—Contextos L -fuzzy, extracción de conocimiento, materiales renovables, diseño de materiales

I. INTRODUCCIÓN

Cuando se diseña un nuevo material, la selección de la formulación es uno de los principales factores que afecta a las propiedades funcionales del material. En el proceso usual del diseño de nuevos materiales, se varía la cantidad de los componentes de la formulación y se estudia la influencia de esas variaciones en las propiedades del material. Dicho proceso implica un elevado coste de materias primas, energía y recursos humanos. Cuando un/a ingeniero/a diseña un nuevo material, los requerimientos del material en servicio dependiendo de su aplicación determinan las propiedades funcionales que el material debe poseer. Habitualmente, existe un conocimiento existente por parte del equipo investigador que ayuda a establecer unos rangos iniciales de los valores de la formulación con el objetivo de obtener las propiedades deseadas. Muchas veces, el investigador no busca un valor exacto para una propiedad concreta, sino que desea que esté dentro de un intervalo determinado o entorno a un valor exacto. Los valores de esas propiedades pasan a ser difusas, lo que posibilita el estudio mediante variables borrosas [1], [2].

La inexistencia a priori de reglas de asociación que relacionen la composición con las diversas propiedades del

material hace que el uso de herramientas difusas basadas en reglas de asociación no parezca la mejor opción, siendo el empleo de contextos borrosos para la extracción de información una herramienta prometedora para la estimación de la composición del material para unas propiedades determinadas. Las propiedades requeridas van cambiando en función de las aplicaciones. Es decir, según el alimento que se quiera envasar, los requerimientos en servicio del material de envasado varían. El deseo de obtener una modificación en alguna de las propiedades obtenidas o en alguna variable de la formulación conlleva un coste inmenso de recursos debido al proceso de prueba y error que se utiliza actualmente. El empleo de contextos borrosos con información incompleta puede ayudar en la fase de diseño del material agilizando el proceso y ahorrando recursos de todo tipo que dicho diseño conlleva.

II. DESCRIPCIÓN DEL PROBLEMA

II-A. Descripción de las condiciones de diseño de materiales

En el diseño de nuevos materiales, la formulación y las condiciones de procesamiento desempeñan un papel fundamental en el comportamiento en servicio de dichos materiales. Estas propiedades están relacionadas con distintas magnitudes (químicas, físicas, mecánicas, térmicas, ópticas...). Para cada aplicación, sabemos qué valores de salida de estas magnitudes debería tener el material para un desempeño óptimo, y la cuestión es encontrar la formulación que lo permita.

Con vistas a alcanzar este objetivo, no sólo se debe identificar las especies químicas que optimizan las propiedades, sino también sus cantidades. Estas cantidades serán los valores de entrada que se han de determinar en la optimización. Usualmente, existe un conocimiento previo que permite al/la diseñador/a fijar los rangos aproximados en donde es probable que estén los valores de entrada ideales, para obtener los valores de salida asociados al comportamiento que son los deseados.

A la hora de encontrar esos valores de entrada ideales, la aproximación tradicional de prueba y error puede requerir un gran coste de energía, tiempo y materias primas. Hay tantas combinaciones y variaciones a estudiar que éste es un procedimiento ineficiente para analizar la influencia de cada parámetro en el comportamiento del material diseñado. Además de esto, normalmente el análisis se realiza con una variable de cada vez, que no considera las interacciones cruzadas entre distintas variables de entrada en la formulación. A menudo, estas entradas y salidas son borrosas: el/la diseñador/a aceptaría valores de entrada alrededor de ciertos valores de entrada, porque estaría dispuesto/a a admitir valores salida cerca de ciertos valores de salida.

En este caso, les podemos asociar valores borrosos de deseabilidad, y realizar un estudio basado en los conceptos de la teoría borrosa. Dentro de esta aproximación, y con un análisis estricto de las variables implicadas, se puede lograr un conjunto de variables de entrada adecuadas, de una manera eficiente y reduciendo el consumo de tiempo y recursos. Así, se evita un procedimiento ineficiente y sus costos añadidos.

II-B. Propiedades requeridas para aplicaciones del material

Los films basados en gelatina muestran características que los hacen indicados para ser aplicados en envasado [3]–[5]. Las propiedades barrera, el comportamiento mecánico y las propiedades ópticas son de vital importancia al diseñar materiales para envasado de alimentos [6].

Teniendo esto en mente, se proponen las siguientes propiedades de salida para un análisis posterior:

- Ángulo de contacto (CA)
- Tasa de transmisión de vapor de agua (WVTR)
- L^* y b^* (en lo que respecta a propiedades de color)
- Tensión a tracción (TS)
- Elongación a rotura (EB)
- Brillo

Considerando la experiencia en el campo del grupo de investigación *BIOMAT - biopolymeric materials* del que forman parte varios de los autores del presente trabajo, se propuso un conjunto de propiedades salida objetivo para esta aplicación. Este set tiene en cuenta las especificidades de los films basados en gelatina. Este es el punto de partida para lanzar el análisis de conceptos L -fuzzy [7]–[9]. Estas propiedades objetivo están listadas en la Tabla I.

Tabla I
REQUERIMIENTOS PARA MATERIALES DE ENVASADO DE ALIMENTOS GRASOS

CA (°)	L^*	b^*	TS (MPa)	EB (%)	Brillo _{60°} (Unidad de brillo)
90	95	40	60	5	50

En el caso de la tasa de transmisión de vapor de agua, no se seleccionaron valores objetivo por no ser una propiedad clave para la aplicación de envasado de comida grasa. Puede ser que las propiedades objetivo no se puedan alcanzar simultáneamente para un material dado, debido a la interacción

entre ellos. En este caso, el método de estimación propuesto proporciona una solución de compromiso en el que las propiedades estimadas están cerca de las propiedades objetivo.

En la siguiente sección se discuten los conceptos básicos de la teoría de contextos L -fuzzy y la forma de aplicarlos a la selección de formulaciones de materiales.

III. MODELADO DE LOS EXPERIMENTOS UTILIZANDO CONTEXTOS L -FUZZY

III-A. Análisis de contextos L -fuzzy

El análisis formal de conceptos es una herramienta matemática desarrollada por R. Wille en 1982 [10], [11] con el objetivo de procesar datos conceptuales y representarlos de una manera formal. Con el objetivo de reconocer diferentes grados de pertenencia entre objetos y atributos Burusco y Fuentes-González extendieron los conceptos formales de Wille al caso difuso con la definición de contextos L -fuzzy [7]–[9]. Un contexto L -fuzzy se define como una tupla (L, X, Y, R) donde L es un retículo completo, X e Y son respectivamente el conjunto de objetos y atributos, y la relación $R \in L^{X \times Y}$ toma valores en el retículo $(L, \leq)$.

Para extraer información de estos contextos L -fuzzy, el operador derivación se define para todos los conjuntos $A \in L^X$ y $B \in L^Y$ mediante las siguientes expresiones [9]:

$$A_1(y) = \inf_{x \in X} \{I(A(x), R(x, y))\}, \forall y \in Y,$$

$$B_2(x) = \inf_{y \in Y} \{I(B(y), R(x, y))\}, \forall x \in X.$$

Siendo I un operador implicación difuso definido en L .

Los conceptos L -fuzzy constituyen una herramienta que permite visualizar la información almacenada en el contexto. Estos conceptos son pares (M, M_1) donde el conjunto $M \in L^X$ es un punto fijo del operador constructor φ , que se define usando los operadores de derivación $\varphi(A) = (A_1)_2 = A_{12}$ para todo $A \in L^X$.

De manera equivalente, los conceptos L -fuzzy se pueden definir desde el punto de vista de los atributos como pares (N_2, N) siendo $N \in L^Y$ un punto fijo del operador constructor ϕ que se define como $\phi(B) = (B_2)_1 = B_{21}$ para todo $B \in L^Y$.

Dado un subconjunto de objetos $A \in L^X$ (o, un subconjunto de atributos $B \in L^Y$), es posible calcular el concepto L -fuzzy asociado aplicando repetidamente el operador constructor φ (o, equivalentemente, el operador ϕ) hasta que un punto fijo es obtenido. Podemos obtener un punto fijo simplemente aplicando el operador derivación dos veces si el operador de implicación usado es residuo [12], [13]. Así, si I es residuo, entonces el par (A_{12}, A_1) es el concepto L -fuzzy asociado a A . De manera similar, el concepto L -fuzzy asociado a $B \in L^Y$ será el par (B_2, B_{21}) .

El concepto L -fuzzy asociado a un subconjunto de objetos A (o atributos B) proporciona la situación “estable” en el contexto que sea más cercana a las condiciones iniciales fijadas por A (o B). En este trabajo, y con el fin de simplificar los cálculos, el operador de implicación elegido para obtener los conjuntos derivados fue la implicación de Lukasiewicz.

III-B. Metodología propuesta para extraer información de los datos experimentales

La metodología propuesta en este trabajo aprovecha el potencial que tiene el análisis de contextos L -fuzzy para extraer información de tablas de relaciones con doble entrada, y se puede describir en los siguientes pasos:

1. Determinar los conjuntos de objetos y atributos que forman parte del problema. La relación entre ellos vendrá dada por los datos experimentales.
2. Normalizando los valores experimentales, construir el contexto L -fuzzy relacionado con las propiedades del material a analizar.
3. Establecer el conjunto de atributos que representa la situación que queremos estudiar. Para las propiedades para las que no tenemos requerimientos iniciales se tomará el valor de 0.
4. Calcular el concepto L -fuzzy más cercano a la situación de partida, que vendrá dada por el punto fijo del operador constructor.
5. Interpretar los valores obtenidos para extraer la información requerida.

El Algoritmo 1 describe los cálculos a realizar para obtener los valores estimados de las propiedades de los materiales.

IV. DETERMINACIÓN DE LA FORMULACIÓN DEL FILM BASADO EN GELATINA

IV-A. Diseño experimental para la optimización de la formulación

Los valores de entrada de la formulación a ajustar elegidos fueron el contenido de glicerol, pH y contenido de ácido gálico. Los datos para poder aplicar la metodología difusa propuesta, fueron distintas muestras donde los valores de contenido de glicerol analizados fueron 0 %, 5 % y 10 %, los valores de pH seleccionados 4.5, 7.25 y 10, mientras que los valores de contenido de ácido gálico fueron 5 %, 10 % y 15 %. Los valores experimentales de las propiedades de salida obtenidas para alimentar los cálculos se muestran en la Tabla II.

IV-B. Construcción del contexto L -fuzzy que representa los valores experimentales

Con el objetivo de modelar el proceso de diseño de materiales construimos un contexto L -fuzzy (L, X, Y, R) en el retículo $(L = [0, 1]_{100}, \leq)$, considerando el conjunto de objetos $X = \{x_1, x_2, \dots, x_{13}\}$, donde x_i representa la i -ésima formulación. El conjunto de atributos $Y = \{y_1, y_2, \dots, y_{10}\}$ estuvo formado por las diferentes propiedades que se analizaron en cada formulación:

- y_1 : Contenido en glicerol
- y_2 : Contenido en ácido gálico
- y_3 : pH de la solución
- y_4 : Ángulo de contacto de agua
- y_5 : Tasa de transmisión de vapor de agua
- y_6 : Valor de color L^*
- y_7 : Valor de color b^*

Algoritmo 1 Estimación de las propiedades de los materiales

Entradas:

- 1: $\{x_1, x_2, \dots, x_n\}$: conjunto de objetos X .
- 2: $\{y_1, y_2, \dots, y_m\}$: conjunto de atributos Y .
- 3: $E(x_i, y_j)$ para todo $(x_i, y_j) \in X \times Y$: valores experimentales.
- 4: $\{r_1, r_2, \dots, r_k\}$: requerimientos para los atributos $\{y_{j_1}, y_{j_2}, \dots, y_{j_k}\}$.

Salida:

$\{e_1, e_2, \dots, e_m\}$: valores estimados para los atributos.

Pasos:

► Construcción de la relación del contexto.

- 1: **for** $i = 1$ **to** n **do**
- 2: **for** $j = 1$ **to** m **do**
- 3:
$$R(x_i, y_j) \leftarrow \frac{E(x_i, y_j) - \min_{1 \leq h \leq n} \{E(x_h, y_j)\}}{\max_{1 \leq h \leq n} \{E(x_h, y_j)\} - \min_{1 \leq h \leq n} \{E(x_h, y_j)\}}$$
- 4: **end for**
- 5: **end for**

► Construcción del conjunto inicial de atributos para representar la información requerida.

- 6: **for** $j = 1$ **to** m **do**
- 7: **if** $j \in \{j_1, j_2, \dots, j_k\}$ **then**
- 8:
$$B(y_j) \leftarrow \frac{r_j - \min_{1 \leq h \leq n} \{E(x_h, y_j)\}}{\max_{1 \leq h \leq n} \{E(x_h, y_j)\} - \min_{1 \leq h \leq n} \{E(x_h, y_j)\}}$$
- 9: **else**
- 10: $B(y_j) \leftarrow 0$
- 11: **end if**
- 12: **end for**

► Cálculo del concepto asociado.

- 13: **while** $B \neq \phi(B)$ **do**
- 14: $B \leftarrow \phi(B)$
- 15: **end while**

► Obtención de los valores estimados para las propiedades.

- 16: **for** $j = 1$ **to** m **do**
- 17:
$$e_j \leftarrow \min_{1 \leq h \leq n} \{E(x_h, y_j)\} +$$

$$+ B(y_j) \cdot \left(\max_{1 \leq h \leq n} \{E(x_h, y_j)\} - \min_{1 \leq h \leq n} \{E(x_h, y_j)\} \right)$$
- 18: **end for**

- y_8 : Tensión a tracción.
- y_9 : Elongación a rotura
- y_{10} : Valor de brillo a un ángulo de incidencia de 60°

La relación entre objetos y atributos $R \in L^{X \times Y}$ se calculó normalizando los elementos de la matriz E formada por los valores experimentales promedio que se muestran en la Tabla II de la siguiente manera:

$$R(x_i, y_j) = \frac{E(i, j) - \min_{1 \leq h \leq 13} \{E(h, j)\}}{\max_{1 \leq h \leq 13} \{E(h, j)\} - \min_{1 \leq h \leq 13} \{E(h, j)\}}$$

La relación obtenida se muestra en la Tabla III.

Tabla II

VALORES EXPERIMENTALES DE LAS PROPIEDADES FUNCIONALES PARA DIFERENTES FORMULACIONES DE MATERIALES DE ORIGEN RENOVABLE

Formulación			Propiedades de los materiales						
GLY (%)	GA (%)	pH	CA (°)	WVTR ($\frac{g}{m^2 \cdot día}$)	L*	b*	TS (MPa)	EB (%)	Brillo _{60°} (Unidades de brillo)
0	10	4.50	74.8 ± 4.2	1057.7 ± 11.5	95.3 ± 0.5	5.4 ± 0.4	78.7 ± 12.5	3.8 ± 1.5	125.4 ± 3.2
5	15	4.50	74.6 ± 7.8	1147.3 ± 16.5	95.5 ± 0.2	5.4 ± 0.8	83.4 ± 8.4	2.7 ± 0.6	155.4 ± 10.3
5	10	7.25	52.8 ± 2.3	1272.1 ± 13.4	50.2 ± 6.6	44.6 ± 6.9	79.8 ± 3.7	3.1 ± 0.4	101.6 ± 8.5
5	5	4.50	58.9 ± 2.2	1097.3 ± 8.0	95.7 ± 0.4	4.8 ± 0.5	80.5 ± 2.1	3.9 ± 0.4	148.0 ± 9.9
0	10	10.00	38.3 ± 2.5	1256.3 ± 10.2	32.5 ± 1.3	20.7 ± 1.8	56.3 ± 5.4	2.7 ± 0.4	50.1 ± 5.6
10	10	10.00	44.6 ± 1.9	1198.0 ± 22.3	26.6 ± 0.5	8.8 ± 0.5	57.9 ± 1.2	2.6 ± 0.2	29.6 ± 3.7
5	5	10.00	44.7 ± 4.6	1227.0 ± 10.6	69.6 ± 5.7	59.1 ± 3.2	70.8 ± 5.6	3.1 ± 0.4	23.5 ± 2.6
0	5	7.25	56.7 ± 1.8	1206.0 ± 8.5	47.2 ± 0.7	38.6 ± 0.7	77.9 ± 10.4	3.5 ± 0.5	28.6 ± 2.5
10	15	7.25	99.3 ± 2.5	1310.3 ± 19.4	46.6 ± 1.8	42.8 ± 1.5	72.1 ± 6.2	2.7 ± 0.4	100.0 ± 6.4
5	15	10.00	40.8 ± 2.5	1001.7 ± 15.3	23.8 ± 0.9	4.0 ± 0.8	62.3 ± 4.0	2.6 ± 0.3	70.9 ± 3.2
10	10	4.50	101.0 ± 3.7	960.3 ± 22.5	95.2 ± 0.2	6.5 ± 0.9	81.8 ± 8.0	3.1 ± 0.6	148.4 ± 6.1
10	5	7.25	68.2 ± 4.7	1220.0 ± 20.7	52.8 ± 3.9	48.6 ± 1.3	77.5 ± 4.1	3.6 ± 0.4	97.6 ± 1.6
0	15	7.25	67.7 ± 4.2	1304.3 ± 9.3	54.1 ± 1.0	49.8 ± 1.5	86.4 ± 7.9	3.2 ± 0.6	100.0 ± 7.0

Tabla III

RELACIÓN DEL CONTEXTO *L*-FUZZY

<i>R</i>	<i>y</i> ₁	<i>y</i> ₂	<i>y</i> ₃	<i>y</i> ₄	<i>y</i> ₅	<i>y</i> ₆	<i>y</i> ₇	<i>y</i> ₈	<i>y</i> ₉	<i>y</i> ₁₀
<i>x</i> ₁	0.00	0.50	0.00	0.58	0.28	0.99	0.03	0.66	0.52	0.77
<i>x</i> ₂	0.50	1.00	0.00	0.58	0.53	1.00	0.02	0.80	0.02	1.00
<i>x</i> ₃	0.50	0.50	0.50	0.23	0.89	0.37	0.74	0.70	0.21	0.59
<i>x</i> ₄	0.50	0.00	0.00	0.33	0.39	1.00	0.01	0.72	0.56	0.94
<i>x</i> ₅	0.00	0.50	1.00	0.00	0.85	0.12	0.30	0.00	0.02	0.20
<i>x</i> ₆	1.00	0.50	1.00	0.10	0.68	0.04	0.09	0.04	0.00	0.05
<i>x</i> ₇	0.50	0.00	1.00	0.10	0.76	0.64	1.00	0.43	0.21	0.00
<i>x</i> ₈	0.00	0.00	0.50	0.29	0.70	0.32	0.63	0.64	0.35	0.04
<i>x</i> ₉	1.00	1.00	0.50	0.97	1.00	0.32	0.70	0.47	0.02	0.58
<i>x</i> ₁₀	0.50	1.00	1.00	0.04	0.12	0.00	0.00	0.18	0.00	0.36
<i>x</i> ₁₁	1.00	0.50	0.00	1.00	0.00	0.99	0.04	0.76	0.19	0.95
<i>x</i> ₁₂	1.00	0.00	0.50	0.48	0.74	0.40	0.81	0.63	0.40	0.56
<i>x</i> ₁₃	0.00	1.00	0.50	0.47	0.98	0.42	0.83	0.89	0.24	0.58

IV-C. Conceptos *L*-fuzzy asociados a los requerimientos iniciales

Una vez definido el contexto *L*-fuzzy (*L*, *X*, *Y*, *R*), se estimaron las características del material a partir de los conceptos *L*-fuzzy asociados con el conjunto de atributos que representan los requisitos deseados que se muestran en la Tabla I.

Los valores requeridos para las propiedades se normalizaron para obtener valores en *L*, y se supuso inicialmente que los valores de pertenencia de los atributos que representan las propiedades desconocidas eran 0.

Por lo tanto, las propiedades del material para la aplicación de envases de alimentos grasos fueron representadas por el siguiente conjunto difuso de atributos:

$$B = \{y_1/0, y_2/0, y_3/0, y_4/0.82, y_5/0, y_6/0.99, y_7/0.65, y_8/0.11, y_9/1, y_{10}/0.20\}$$

Y, después calcular el concepto *L*-fuzzy asociado, nos centramos en su intensidad:

$$B_{21} = \{y_1/0.63, y_2/0.60, y_3/0.63, y_4/0.89, y_5/0.81, y_6/0.99, y_7/0.65, y_8/0.98, y_9/1, y_{10}/0.70\}$$

La intensidad de un concepto se puede interpretar como un conjunto de atributos que siempre se encuentran juntos en el contexto. Por lo tanto, el conjunto de atributos B_{21} obtenido se

puede interpretar como el conjunto de atributos más cercano al conjunto requerido *B* que se dan al mismo tiempo.

A partir de los valores de pertenencia de la intensidad de los conceptos *L*-fuzzy obtenidos, invirtiendo el proceso de normalización, se estimó que, para el envasado de alimentos grasos, las propiedades del material deberían ser las que se muestran en la Tabla IV.

V. RESULTADOS Y DISCUSIÓN

Teniendo en cuenta los datos experimentales para el análisis de contexto *L*-fuzzy, las formulaciones estimadas para un rendimiento óptimo son las dadas en la Tabla IV. Para la aplicación de envasado de alimento graso, la metodología sugiere valores alrededor del 6.3 % de contenido en glicerol, 11 % de contenido en ácido gálico, y un pH de 8.

Con estos valores se realizaron medidas experimentales de las propiedades funcionales en el laboratorio, obteniéndose los valores presentados en la Tabla V.

Como puede observarse en la Tabla V, los valores estimados fueron muy similares a los valores objetivo de las propiedades funcionales, excepto para los valores de brillo y tensión a tracción, y los valores propuestos por el análisis de conceptos *L*-fuzzy fueron mayores que los valores objetivo seleccionados. En lo que respecta a valores de brillo, no es posible obtener valores cercanos a 50 teniendo en cuenta las restricciones

Tabla IV
FORMULACIÓN Y PROPIEDADES ESTIMADAS PARA LOS MATERIALES DE ENVASADO DE ALIMENTOS GRASOS

GLY (%)	GA (%)	pH	CA (°)	WVTR $\left(\frac{g}{m^2 \cdot día}\right)$	L*	b*	TS (MPa)	EB (%)	Brillo _{60°} (Unidades de brillo)
6.28	10.96	7.95	93.94	1243.66	95.00	40.00	89.38	5.00	116.27

Tabla V
VALORES REQUERIDOS PARA EL MATERIAL, ESTIMACIÓN, VALORES EXPERIMENTALES Y ERROR OBTENIDO.

Formulación				Propiedades de los materiales					
GLY (%)	GA (%)	pH	CA (°)	WVTR $\left(\frac{g}{m^2 \cdot día}\right)$	L*	b*	TS (MPa)	EB (%)	Brillo _{60°} (Unidades de brillo)
Valores requeridos			90		95	40	60	5	50
Valores estimados	6.3	11	93.9	1243.7	95	40	89.4	5	116.3
Valores experimentales			74.3	1370.3	60.9	42.1	83.3	4.8	95.3
Error (%)			20.9	10.2	35.9	5.3	6.7	4.9	18

en el resto de propiedades; esto indica que los valores de brillo deben ser mayores. Como los valores de brillo están relacionados inversamente con la rugosidad de la superficie del film [14], valores mayores de brillo indican superficies más lisas. Algo similar sucedió con los valores de tensión a tracción: los valores estimados fueron mayores que el objetivo. Es interesante señalar que el comportamiento mecánico del material propuesto por el método de estimación es incluso mejor que el de la especificación, pues tiene mayor tensión a tracción, sin disminuir los valores de la elongación a rotura. Para envasado de alimentos grasos, la tasa de transmisión de vapor de agua no es una propiedad crucial, y es por ello que no se seleccionó ningún valor objetivo; pero el concepto *L-fuzzy* estima un valor de $1243 \frac{g}{m^2 \cdot día}$ para esta propiedad. Para verificar el método de estimación empleado, se desarrolló un material con la formulación propuesta, y se midieron las propiedades funcionales respuesta del mismo. Como se muestra en la Tabla V, para la mayoría de propiedades, los valores propuestos por el contexto *L-fuzzy* son razonablemente similares a los valores experimentales. Los errores obtenidos son menores al 21 % en todos los casos, excepto para el valor de color L*, para el que el error es mayor y el valor experimental obtenido es menor que la estimación. Esto produce una luminosidad mayor para el film, mayor que la requerida en la especificación, lo cual es adecuado para la aplicación de envasado de alimentos. Además, se puede confirmar, tal y como predijo el contexto *L-fuzzy*, que los valores de brillo fueron mayores que los valores objetivo y el comportamiento mecánico del material es incluso mejor que el requerido para la formulación estimada.

VI. CONCLUSIONES

En el presente trabajo, se empleó el análisis de conceptos *L-fuzzy* como una herramienta de toma de decisiones para obtener una formulación del material con las propiedades funcionales requeridas en envasado de alimentos grasos. Los resultados obtenidos después de implementar la metodología propuesta revelaron que los valores predichos por los conceptos *L-fuzzy* fueron muy similares a los valores objetivo

en la mayoría de propiedades. Además, después de analizar el comportamiento de las formulaciones estimadas, los resultados experimentales resultaron ser muy cercanos a los valores aproximados propuestos. El error entre valores estimados por los conceptos *L-fuzzy* y los valores experimentales de las formulaciones propuestas fue inferior al 20 % en la mayoría de propiedades funcionales. Por otro lado, en los casos en que la diferencia entre valor predicho y valor experimental fue mayor del 20 %, la predicción proporciona una mejora de las propiedades requeridas, obteniendo un mejor comportamiento del material.

Definitivamente, las formulaciones propuestas por el análisis de conceptos *L-fuzzy* son el mejor punto de partida para estudiar el efecto de cada parámetro en las propiedades funcionales del material. En vez de tener una enorme matriz de formulaciones que estudiar, utilizar este método de estimación ayuda a identificar aproximadamente los valores de los parámetros de la formulación que llevan a obtener las propiedades deseadas, y permite saber si es posible conseguir las propiedades funcionales para un sistema específico.

AGRADECIMIENTOS

Los autores agradecen a los grupos de investigación del Gobierno Vasco (IT1256-19) y de la Universidad del País Vasco UPV/EHU (GIU18/154).

REFERENCIAS

- [1] A. Balakrishna, G. R. Chandra, B. Gogulamudi, C. Someswararao: Fuzzy Approach to the Selection of Material Data in Concurrent Engineering Environment, *Engineering* 3 (9), pp. 921–927, 2011.
- [2] T. W. Liao: A fuzzy multicriteria decision-making method for material selection, *Journal of Manufacturing Systems* 15 (1), pp. 1–12, 1996.
- [3] H. Haghighi, S. K. Leugoue, F. Pfeifer, H. W. Siesler, F. Licciardello, P. Fava, A. Pulvirenti: Development of antimicrobial films based on chitosan-polyvinyl alcohol blend enriched with ethyl lauroyl arginate (LAE) for food packaging applications, *Food Hydrocolloids* 100, pp. 105419, 2020.
- [4] V. K. Pandey, S. N. Upadhyay, K. Niranjan, P. K. Mishra: Antimicrobial biodegradable chitosan-based composite Nano-layers for food packaging, *International Journal of Biological Macromolecules* 157, pp. 212–219, 2020.
- [5] R. Priyadarshi, J. W. Rhim: Chitosan-based biodegradable functional films for food packaging applications, *Innovative Food Science & Emerging Technologies* 62, pp. 102346, 2020.

- [6] I. Leceta, P. Guerrero, K. de la Caba: Functional properties of chitosan-based films, *Carbohydrate Polymers* 93 (1) , pp. 339–346, 2013.
- [7] A. Burusco, R. Fuentes-González: The study of the L-fuzzy concept lattice, *Mathware and Soft Computing* 1 (3), pp. 209–218, 1994.
- [8] A. Burusco, R. Fuentes-González: Construction of the L-fuzzy Concept Lattice, *Fuzzy Sets and Systems* 97 (1), pp. 109–114, 1998.
- [9] A. Burusco, R. Fuentes-González: Concept lattices defined from implication operators, *Fuzzy Sets and Systems* 114 (3), pp. 431–436, 2000.
- [10] B. Ganter, R. Wille: Formal concept analysis: Mathematical foundations, Springer, Berlin - New York, 1999.
- [11] R. Wille: Restructuring lattice theory: an approach based on hierarchies of concepts, in: *Rival I. (Ed.), Ordered Sets, Reidel, Dordrecht-Boston*, pp. 445–470, 1982.
- [12] R. Bělohlávek: Fuzzy Galois Connections, *Mathematical Logic Quarterly* 45 (4), pp. 497–504, 1999.
- [13] S. Pollandt: Fuzzy-Begriffe: Springer Berlin Heidelberg, 1997.
- [14] L. Sánchez-González, M. Cháfer, A. Chiralt, C.o González-Martínez: Physical properties of edible chitosan films containing bergamot essential oil and their inhibitory action on *Penicillium italicum*, *Carbohydrate Polymers* 82 (2), pp. 277–283, 2010.

Una versión ternaria del algoritmo de Quine McCluskey para la minimización de base de reglas difusas

Leonardo Jara

Departamento de Ciencias de la Computación e Inteligencia Artificial, Instituto Interuniversitario de Investigación en Data Science and Computational Intelligence (DaSCI) Universidad de Granada 18071-Granada, España E-mail: ljara@correo.ugr.es

Antonio González

Departamento de Ciencias de la Computación e Inteligencia Artificial Instituto Interuniversitario de Investigación en Data Science and Computational Intelligence (DaSCI) Universidad de Granada 18071-Granada, España E-mail: A.Gonzalez@decsai.ugr.es

Raúl Pérez

Departamento de Ciencias de la Computación e Inteligencia Artificial Instituto Interuniversitario de Investigación en Data Science and Computational Intelligence (DaSCI) Universidad de Granada 18071-Granada, España E-mail: fgr@decsai.ugr.es

Resumen—Los sistemas de clasificación basados en reglas difusas permiten en muchos casos obtener un conocimiento fácil de interpretar, pero esto no siempre ocurre. El objetivo de este trabajo consiste en mejorar la interpretabilidad de estos sistemas, utilizando procesos de minimización de conjuntos de reglas difusas basados en una modificación del método de Quine McCluskey, que permite trabajar con valores ternarios. De este modo, se pretende mejorar la capacidad de eliminación de variables irrelevantes en los antecedentes de las reglas y por lo tanto, mejorar la interpretabilidad del modelo sin alterar en gran medida la precisión.

Palabras Clave—Reglas difusas, Clasificación, Interpretabilidad, Simplicidad, Minimización.

I. INTRODUCCION

Tradicionalmente la gran ventaja de los sistemas de clasificación basados en reglas difusas es que son capaces de obtener un conocimiento interpretable a la vez que consiguen buenos niveles de precisión. Sin embargo, esto no es siempre así, y en algunos casos el conocimiento obtenido no es tan fácil de interpretar. Este punto es importante ya que un algoritmo que obtenga conocimiento interpretable permite al usuario entender el proceso seguido en la toma de decisiones. La interpretabilidad es un tema de discusión frecuente [1] y existe una gran cantidad de estudios y métricas que nos permiten comparar modelos. Uno de los parámetros frecuentemente utilizado para estudiar la interpretabilidad de un sistema es el número de reglas, de forma que cuanto mayor sea menor es la interpretabilidad.

Existen estudios que buscan mejorar la interpretabilidad de la base de reglas como el CFM-BD [2] que utiliza técnicas de probabilidad e inducción de reglas, creando así una base de reglas más compacta. También tenemos los que buscan reducir

el número de reglas difusas utilizando los mapas de Karnaugh [3]. El mapa de Karnaugh es un método gráfico que se utiliza para simplificación de funciones algebraicas de forma canónica. El principal inconveniente con esta técnica radica en que solo puede usar un máximo de 6 variables binarias, las cuales solo pueden representar dos variables difusas con 3 etiquetas lingüísticas, lo que limita bastante su uso.

Este problema también lo encontramos en uno de los algoritmos de clasificación de reglas difusas basados en el algoritmo de Wang y Mendel para clasificación [4], conocido como el algoritmo de CHI [5]. Este algoritmo tiene como ventaja el ser muy eficiente en el proceso aprendizaje a la vez que obtiene niveles de precisión aceptables, pero como inconveniente genera un gran número de reglas difusas [6], lo que dificulta la interpretabilidad del modelo obtenido.

Debido a estos inconvenientes se propuso un nuevo método para minimizar la base de reglas difusas [7], basado en el algoritmo de Quine McCluskey [8], en donde se consiguieron buenos resultados, pero sus principales inconvenientes se encontraban en la adaptación de las etiquetas lingüísticas a una codificación binaria, lo que creaba tablas de la verdad muy extensas y no lograba eliminar variables.

Quine McCluskey es un método de minimización de funciones booleanas ampliamente utilizado para la minimización de circuitos lógicos, debido a que obtiene el mínimo valor de una función booleana con un método tabular. La lógica binaria es eficiente y potente, pero es uno de los inconvenientes para cierto tipo de problemas como el propuesto en este trabajo donde queremos trabajar con un número de tres etiquetas lingüísticas.

En este trabajo proponemos el uso de una lógica ternaria [9] que reemplaza los valores de 0 y 1 de la lógica binaria por valores 0, 1 y 2, lo que nos permitirá realizar una reducción más eficiente del conjunto de reglas y poder mejorar por tanto la interpretabilidad del sistema.

Este documento los dividimos de la siguiente forma, en

Trabajo cofinanciado por el Programa Nacional de Becas de Postgrado en el Extranjero "Don Carlos Antonio López, BECAL", otorgado por el Gobierno de la República del Paraguay y por el proyecto de investigación RTI2018-098460-B-I00 financiado a su vez con fondos FEDER (Unión Europea).

la siguiente sección describimos las herramientas preliminares del trabajo, en la sección III se explicará en detalle el funcionamiento del algoritmo. La sección IV muestra los resultados obtenidos con el método de minimización ternario comparándolo con el modelo de minimización binario, y por último en la sección V se realiza un análisis de la experimentación realizada.

II. PRELIMINARES

Describimos en esta sección las herramientas básicas que se requieren en la realización de nuestra propuesta.

II-A. Modelo de regla

La propuesta de minimización que realizamos en este trabajo parte de un conjunto de reglas difusas básicas que podrían obtenerse de múltiples formas, sin embargo, con la idea de poder realizar una experimentación completa del modelo, asumiremos que el conjunto de reglas inicial se ha obtenido a partir del algoritmo Chi.

La estructura habitual de un conjunto de reglas difusas utilizando el algoritmo Chi es [10]:

$$\text{Rule } R_j : \text{ If } x_1 \text{ is } A_{j1} \text{ and } \dots \text{ and } x_n \text{ is } A_{jn} \text{ then} \\ \text{Class} = C_j \text{ with } RW_j \quad (1)$$

Donde R_j es la etiqueta j-ésima de la regla, $x = (x_1, \dots, x_n)$ es un vector n dimensional que representa un ejemplo, A_{ji} es conjunto difuso representado por una función de pertenencia triangular, C_j es la etiqueta de la clase y RW_j es el peso de la regla. El peso de la regla utilizado en este trabajo es el factor de certeza penalizado [11] PCF:

$$RW_j = PCF = \frac{\sum_{x_p \in \text{Class } C_j} \mu_{A_j}(x_p) - \sum_{x_p \notin \text{Class } C_j} \mu_{A_j}(x_p)}{\sum_{p=1}^P \mu_{A_j}(x_p)} \quad (2)$$

Donde $\mu_{A_j}(x_p)$ es el grado de pertenencia del ejemplo x_p con la parte del antecedente de la regla difusa R_j y se calcula como sigue.

$$\mu_{A_j}(X_p) = \prod_{i=1}^n \mu_{A_{ji}}(X_{pi}) \quad (3)$$

En donde $\mu_{A_{ji}}(X_{pi})$ es el grado de pertenencia del valor X_{pi} del conjunto difuso A_{ji} de la regla R_j .

Denominaremos *modelo de reglas difusas extendido* [12] al modelo de regla difusa que permite que el valor asignado a una variable sea un subconjunto de las etiquetas difusas de su dominio. Este modelo de regla es interesante, ya que cuando una variable toma todas las etiquetas de su dominio, la variable se puede considerar irrelevante y puede eliminarse del antecedente de la regla. En la propuesta realizada en este trabajo, las reglas resultantes de la simplificación serán de este tipo.

II-B. Método de Quine McCluskey

En esta sección presentamos las definiciones básicas y los pasos para aplicar el método Quine McCluskey [6].

Definiciones

- Literal: Es una variable lógica o su negación (q or $\bar{q}$).
- Minterm: es una expresión algebraica booleana de n variables booleanas, que solamente se evalúa como verdadera para una única combinación de esas variables.
- Implicante Primo: es el producto que no se puede combinar con otro término para eliminar una variable y una mayor simplificación.
- Implicante Primo esencial: es un implicante primo que es capaz de cubrir una salida de la función que no está cubierta por ninguna combinación de implicantes primos.

El método de Quine McCluskey (QM) utiliza las siguientes tres leyes básicas de simplificación.

- $q + \bar{q} = 1$ (Complemento)
- $q + q = q$ (Idempotente)
- $q(w + z) = qw + qz$ (Distributiva)

Donde q , w y z son literales.

El método QM tiene los siguientes pasos:

- Encontrar el implicante primo: En este paso, se sustituye el literal en forma de 0 y 1 y generamos una tabla. Inicialmente el número de filas de la tabla es igual al número total de minterms de la función original no simplificada. Si dos términos sólo se diferencian en un bit, una variable aparece en ambas formas (variable y negación), entonces se podrá utilizar la ley del complemento. Iterativamente, se comparan todos los términos y se genera el implicante primo.
- Encontrar el implicante primo esencial: Usando los implicantes primos del paso anterior, se genera una tabla para encontrar los implicantes primos esenciales. Algunos implicantes primos pueden ser redundantes y pueden ser omitidos, pero si aparecen sólo una vez, no pueden ser omitidos y proporcionan implicantes primos esenciales.
- Encontrar otro implicante primo: No es necesario que el implicante primo esencial cubra todos los términos mínimos. En ese caso, se considerara otros implicantes primos para asegurarnos de que todos los términos menores han sido cubiertos.
- Cuando no podemos obtener ni un implicante primo esencial estamos ante un problema de cobertura cíclica, que puede ser resuelto analíticamente por el método de Petrick's [13].

En general, el método QM proporciona un método mejor para la simplificación de una función booleana que el mapa de Karnaugh, pero sigue siendo un problema NP-Duro, por lo que la complejidad de vuelve exponencial.

III. MÉTODO PROPUESTO PARA LA MINIMIZACIÓN DE REGLAS

En [7] se propone una técnica que toma como base el modelo de minimización de circuitos digitales de Quine McCluskey para la simplificación y reducción de bases de reglas difusas para clasificación. Dicha técnica considera que los antecedentes de todas las reglas que describen una clase se pueden considerar como un único circuito y la idea es aplicar los criterios de reducción que se describen en el modelo de Quine McCluskey. En las conclusiones del trabajo, los autores describen que el uso de la base binaria produce unos resultados de reducción aceptables, pero que el sistema tiene problemas para eliminar variables completas en el antecedente de la regla y eso limita la capacidad de simplificación y por tanto la mejora en la interpretabilidad del modelo resultante.

En este trabajo proponemos un método de minimización de conjuntos de reglas difusas extendiendo el algoritmo Quine McCluskey de minimización de funciones booleanas para trabajar con lógica ternaria. El algoritmo se extiende para trabajar con tablas de verdad ternarias. Al trabajar con lógica ternaria y con conjuntos de reglas difusas que codifican sus variables difusas usando 3 etiquetas, se mejora la capacidad de eliminación de variables completas en el antecedente de la regla.

El proceso que se sigue es el siguiente: (A) partimos de un conjunto de reglas difusas. En general, entenderemos que las reglas siguen el modelo conocido como "Mamdani", como se describe en la ecuación (1). Además, vamos a suponer que el dominio asociado a cada variable difusa está compuesto por 3 etiquetas uniformemente distribuidas sobre el universo de discurso de la variable. (B) Se transforma el antecedente de cada una de las reglas en un vector con codificación ternaria y a cada uno de estos vectores se le asocia la clase de la regla. (C) Se toma una de las posibles clases C_j y todos los vectores transformados asociados a esa clase y se aplica sobre ellos el algoritmo de minimización de Quine McCluskey ternario. (D) El nuevo conjunto de vectores obtenidos, asociados a la clase C_j , son transformados en reglas y sustituyen a las reglas originales para esa clase. (E) El proceso vuelve al paso 3 hasta que el método se aplica a todas las clases.

Cada variable lingüística tiene un dominio difuso asociado y el algoritmo de CHI utiliza 3 etiquetas lingüísticas impares, en este caso usaremos las siguientes, *small*, *medium* y *large* que se observan en la Figura 1.

III-A. Codificación de la base de reglas.

Como dijimos anteriormente, para utilizar el método de Quine McCluskey es necesario transformar la base de reglas en un formato compatible con el método de Quine McCluskey, donde se utilizan tablas de la verdad binarias. En este trabajo, adaptamos el método utilizando una codificación ternaria y trabajaremos con conjuntos de reglas que usan variables difusas que tienen asociados dominios difusos con 3 etiquetas.

Teniendo como referencia general las tres etiquetas difusas *small*, *medium*, *large*, se usará la codificación que se observa en la Figura 2.

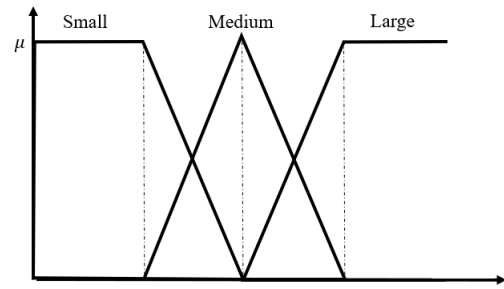


Figura 1. Dominio difuso de una variable lingüística $x_1, \dots, x_n$.

Etiqueta	Código
small	0
medium	1
large	2

Figura 2. Codificación ternaria

III-B. Método de minimización Quine McCluskey ternario

En esta sección describimos la extensión a base ternaria del método de Quine McCluskey. Dicha extensión se puede describir con los siguientes pasos:

1. Convertimos el antecedente de cada regla de la base de reglas a una fila de una tabla de verdad de valores ternarios tomando el sistema de codificación de la Figura 2. Se construye una tabla distinta para cada clase.
2. Sumamos los valores de todas las variables correspondientes a cada fila de esa tabla (cada fila representa el antecedente de cada una de las reglas) y se reordenan las filas en orden ascendente.
3. Creamos una columna con los *minterms* asociados al número de fila a la que pertenece cada regla.
4. Comparamos reglas cuya suma solo difiera en una unidad para buscar elementos redundantes. Si tres reglas difieren en una unidad entre sí y tienen la misma variable con distintas etiquetas difusas, marcaremos esta variable con el símbolo "*" que significa que dicha variable es redundante. Asociamos la nueva regla con los *minterms* pertenecientes a las reglas combinadas.
5. Los términos que no se pueden combinar con tres reglas donde solo una variable tiene las tres etiquetas difusas, se analizan en pares. Comparamos reglas cuya suma difiera en una unidad para buscar elementos redundantes. En caso de que 2 reglas cuya suma difiera en una unidad y tengan una misma variable con distinta etiqueta lingüística, marcaremos esta variable como la unión de las etiquetas difusas. Asociamos la nueva regla con los *minterms* pertenecientes a las reglas combinadas. Los nuevos valores codificados se muestran en la Figura 3.
6. Los términos que no se pueden combinar, los apartamos en otra tabla como implicantes primos con su respectivo *minterm* asociado.

7. Eliminamos los duplicados en la columna del *minterm*.
8. Repetimos el proceso 4,5,6 y 7, agregando las nuevas variables correspondientes a la redundancia “*” del paso 4 y la unión de las etiquetas difusas correspondiente al paso 5, hasta que ninguna fila se puede combinar. Creando así una tabla con los implicantes primos.
9. Para encontrar los implicantes esenciales creamos una tabla de cobertura. Donde las filas corresponden a los *minterm* de las reglas combinadas denominadas implicantes primos del paso 8 y las columnas corresponden a los valores individuales de cada *minterm*. Los *minterms* correspondientes a las reglas redundantes son omitidos en este paso, no se colocan en las columnas. En la tabla de cobertura los *minterms* cubren los implicantes primos. Un implicante primo esencial es aquel implicante primo que cubre un *minterm* y ningún otro implicante primo cubre el mismo *minterm*.
10. Si los implicantes primos no esenciales no tienen un *minterm* en común con los implicantes primos esenciales, se vuelve a crear la tabla de cobertura del paso 9 de modo a extraer los nuevos implicantes primos esenciales. Esto se repite hasta no obtener más implicantes primos esenciales.
11. En el caso de que no exista ningún implicante primo esencial se aplica el método de Petrick’s.

Etiqueta	Código
Small	0
Médium	1
Large	2
Small + Medium + Large	*
Small + Medium	01
Medium + Large	12
Small + Large	02

Figura 3. Codificación ternaria ampliada.

Veamos un ejemplo para aclarar como funciona el método. Supongamos el siguiente conjunto de reglas difusas, todas de la clase C_0 ,

- R_1 : If x_1 is SMALL and x_2 is SMALL and x_3 is MEDIUM and x_4 is LARGE then Class = C_0
 R_2 : If x_1 is SMALL and x_2 is MEDIUM and x_3 is MEDIUM and x_4 is LARGE then Class = C_0
 R_3 : If x_1 is SMALL and x_2 is LARGE and x_3 is MEDIUM and x_4 is LARGE then Class = C_0
 R_4 : If x_1 is SMALL and x_2 is SMALL and x_3 is SMALL and x_4 is LARGE then Class = C_0
 R_5 : If x_1 is MEDIUM and x_2 is SMALL and x_3 is SMALL and x_4 is LARGE then Class = C_0
 R_6 : If x_1 is MEDIUM and x_2 is SMALL and x_3 is SMALL and x_4 is MEDIUM then Class = C_0

Convertimos la base de reglas en una tabla de codificación ternaria.

Regla	Minterm	Variable 1	Variable 2	Variable 3	Variable 4	Suma	Clase
R_4	1	0	0	0	2	2	0
R_1	2	0	0	1	2	3	0
R_5	3	1	0	0	2	3	0
R_6	4	1	0	0	1	3	0
R_2	5	0	1	1	2	4	0
R_3	6	0	2	1	2	5	0

Podemos apreciar que las reglas correspondientes a los *minterm* 2, 5 y 6 tienen todas las etiquetas difusas posibles en la variable 2, por lo que podemos simplificar las reglas 2,5 y 6. Al igual que las reglas 1 y 3 que tienen 2 etiquetas difusas en la variable 1.

Minterm	Variable 1	Variable 2	Variable 3	Variable 4	Clase
2,5,6	0	*	1	2	0
1,3	0-1	0	0	2	0
4	2	0	0	1	0

En la tabla ya no se pueden realizar agrupaciones, a esta tabla la denominamos implicantes primos. Completamos la tabla de cobertura con la idea de encontrar los implicantes primos esenciales.

Minterm	1	2	3	4	5	6
2,5,6		x			x	x
1,3	x		x			
4				x		

Podemos apreciar que cada implicante primo tiene un único valor en alguna columna correspondiente a los *minterms*. Por esto todos los implicantes primos son esenciales y la base de reglas simplificada queda de la siguiente forma.

If x_1 is SMALL and x_3 is MEDIUM and x_4 is LARGE then Class = C_0

If x_1 is SMALL or MEDIUM and x_2 is SMALL and x_3 is SMALL and x_4 is MEDIUM then Class = C_0

If x_1 is LARGE and x_2 is SMALL and x_3 is SMALL and x_4 is MEDIUM then Class = C_0

donde en la primera regla eliminamos la variable x_2 y en la segunda regla usamos el modelo de reglas difusas extendidas en la variable x_1 .

IV. ESTUDIO EXPERIMENTAL

En esta sección queremos comprobar si el método propuesto permite reducir el número de reglas difusas, manteniendo el nivel de precisión, con respecto a los resultados obtenidos por el algoritmo de CHI original, y a la vez si permite obtener reglas más simples que las obtenidas con la propuesta descrita en [7]. Para ello utilizamos las bases de datos descritas en la Tabla I obtenidas del repositorio de la UCI [14]. Usaremos el algoritmo de CHI [5] para generar los conjuntos de reglas.

Serán esos conjuntos de reglas los que usaremos para evaluar el comportamiento de la propuesta y serán comparados con los obtenidos por el método usando codificación binaria [7].

La implementación del código se ha hecho en Python y se han utilizado dominios difusos compuestos por 3 etiquetas difusas con función de pertenencia triangular simétrica.

El estudio se centra en la comparación de tres parámetros: capacidad de predicción P , número de reglas $\#R$ y la simplicidad de cada base de conocimiento *Simp* en el caso binario *BIN* y en el caso ternario *TER*. En este trabajo se ha considerado la simplicidad como la reducción en las condiciones en los antecedentes de las reglas. Los resultados que se muestran son los obtenidos tras realizar una validación cruzada de 10.

Tabla I
BASES DE DATOS UTILIZADAS EN EL ESTUDIO

Base de datos	Ejemplos	Atributos	Número de Clases
Shuttle	57999	9	9
abalone	3756	8	28
iris	150	4	3
newthyroid	215	5	3
banana	5300	2	2
heart	270	13	2
saheart	462	9	2
bupa	345	6	2
led7digit	500	7	10
page-blocks	5472	10	5
mammographic	830	5	2
flare	1066	11	6
pima	768	8	2
phoneme	5404	5	2
australian	690	14	2
crx	653	15	2
balance	625	4	3
german	1000	20	2
thyroid	7200	21	3
magic	19020	10	2
wine	178	13	13
cleveland	297	13	5
automobile	159	25	6
bands	365	19	2
ionosphere	351	33	2
spectfheart	267	44	2
sonar	208	60	2
movement-libras	360	90	15
penbased	10992	16	10
vehicle	846	18	4
vowel	990	11	11
segment	2310	19	7
winequality-white	4898	11	11
spambase	4597	57	2
ring	7400	20	2

La Tabla II muestra una comparativa de las bases de reglas obtenidas por el algoritmo de aprendizaje (CHI) y los dos métodos de simplificación estudiados, la minimización de la base de reglas utilizando el método de Quine McCluskey Binario (CHI + BIN) y la minimización de la base de reglas utilizando el método de Quine McCluskey Ternario (CHI + TER). Donde P representa la precisión y R la cantidad de reglas generadas por el algoritmo.

En esta tabla podemos observar que ambos métodos proporcionan una reducción importante en el número de reglas

final (quedándose aproximadamente en un 58 % del tamaño del conjunto de reglas original). Comparando los métodos binario y ternario en su capacidad de reducción de reglas, se puede observar que no aparecen diferencias significativas entre ambos métodos. Ambos métodos reducen en valores muy parecidos los conjuntos de reglas como se puede observar por los valores medios obtenidos sobre este parámetro. En cuanto a la capacidad de predicción, podemos observar en la Tabla II que los valores medios de P en CHI con respecto a (CHI + BIN) y (CHI + TER), presentan una mínima variabilidad.

Tabla II
RESULTADOS OBTENIDOS

Base de Datos	CHI		CHI + BIN		CHI + TER	
	P	$\#R$	P	$\#R$	P	$\#R$
Shuttle	0,801	28,7	0,802	17,5	0,802	16,7
abalone	0,231	69,4	0,230	42,5	0,228	41,5
iris	0,926	14,7	0,940	7,9	0,940	7,9
newthyroid	0,846	20,4	0,810	7,9	0,833	9,5
banana	0,602	7,9	0,580	4	0,558	4
heart	0,515	217,2	0,493	165,3	0,504	156,9
saheart	0,727	168,7	0,721	78,5	0,721	74,6
bupa	0,578	43,3	0,573	21,8	0,570	20,9
led7digit	0,670	93,9	0,618	92,9	0,676	50,2
page-blocks	0,918	49,7	0,918	26,7	0,916	24,5
mammographic	0,808	45,5	0,812	19	0,816	17
flare	0,567	127,6	0,580	81	0,492	62,2
pima	0,725	101	0,695	43,9	0,697	41,8
phoneme	0,718	50,1	0,712	17,5	0,709	17,5
australian	0,798	313,9	0,783	179	0,783	139,6
crx	0,813	258,2	0,808	154,6	0,805	113,8
balance	0,893	82	0,840	23,6	0,851	24,1
german	0,196	885,2	0,196	836,1	0,196	824,8
thyroid	0,920	463,6	0,910	268,8	0,919	190,8
magic	0,764	313,1	0,744	103,7	0,744	95,4
wine	0,926	124,3	0,921	86	0,926	81,8
cleveland	0,380	239,2	0,354	194,2	0,374	186,7
automobile	0,611	109,7	0,602	86,4	0,577	81,4
bands	0,686	258,2	0,683	194	0,680	197,9
ionosphere	0,652	228,8	0,655	215,6	0,655	213,3
spectfheart	0,663	235,2	0,663	235,2	0,663	235,2
sonar	0,595	187,2	0,596	185,6	0,596	185,6
movement libras	0,733	260,6	0,731	257,3	0,731	257,3
penbased	0,975	3281	0,973	1352,7	0,972	1427,5
vehicle	0,630	378,6	0,617	211,8	0,618	218,4
vowel	0,639	285,9	0,628	193,9	0,622	189,3
segment	0,843	318,5	0,849	139,3	0,852	136,8
winequality-white	0,495	294	0,494	133,4	0,481	132
spambase	0,728	357	0,693	210,2	0,688	211,1
ring	0,552	573	0,545	428,1	0,536	423,9
Media	0,689	299,580	0,679	180,454	0,678	174,626

En la Tabla III podemos observar una comparación de los métodos de minimización binario y ternario, donde la columna *eliminadas* corresponde a la cantidad media de variables irrelevantes que fueron detectadas por el algoritmo de minimización ternario y las columnas *Simp BIN* y *Simp TER* muestran la medida de simplicidad de cada base de datos para los métodos binario y ternario respectivamente de acuerdo a la ecuación (4).

$$S = \frac{(\text{variables.reglas}_{(BIN,TER)}) - \text{eliminadas}}{\text{variables.reglas}_{(CHI)}} \quad (4)$$

Tabla III
ESTUDIO DE INTERPRETABILIDAD

Base de datos	Eliminadas	Simp BIN	Simp TER
Shuttle	1,0	0,610	0,578
abalone	0,9	0,612	0,596
iris	0,0	0,537	0,537
newthyroid	1,8	0,387	0,448
banana	1,0	0,506	0,443
heart	0,7	0,761	0,722
saheart	7,1	0,465	0,438
bupa	2,3	0,503	0,474
led7digit	0,0	0,989	0,535
page-blocks	3,1	0,537	0,487
mammographic	6,5	0,418	0,345
flare	8,0	0,635	0,482
pima	8,7	0,435	0,403
phoneme	7,6	0,349	0,319
australian	8,4	0,570	0,443
crx	10,3	0,599	0,438
balance	25,0	0,288	0,218
german	0,7	0,945	0,932
thyroid	13,1	0,580	0,410
magic	16,1	0,331	0,300
wine	1,6	0,692	0,657
cleveland	0,0	0,812	0,781
automobile	0,1	0,788	0,742
bands	2,8	0,751	0,766
ionosphere	1,6	0,942	0,932
spectfheart	0,0	1,000	1,000
sonar	0,0	0,991	0,991
movement libras	0,0	0,987	0,987
penbased	29,2	0,412	0,435
vehicle	6,4	0,559	0,576
vowel	0,5	0,678	0,662
segment	10,6	0,437	0,428
winequality-white	14,2	0,454	0,445
spambase	19,6	0,589	0,590
ring	23,4	0,747	0,738
Media	6,6	0,626	0,579

En esta última tabla se observa que el método ternario efectivamente es capaz de detectar variables irrelevantes en los antecedentes de las reglas, aunque esa detección no es homogénea. En bases de datos como *iris* o *cleveland* o *sonar* no es capaz de encontrar ninguna variable irrelevante. La razón de esto puede deberse a que el método prioriza la reducción en el número de reglas que la eliminación de variables a la hora de realizar las agrupaciones. Por otro lado, podemos observar que la simplicidad medida como número medio de condiciones, es mejor la aportada por el método ternario. Se han comparado ambos modelos usando el test de Wilcoxon para probar que la simplicidad del modelo (CHI + TER) es mayor que (CHI + BIN), obteniéndose un p-value de 0.9999 y por tanto aceptándose la hipótesis. Este resultado pone de manifiesto que el método ternario permite detectar con más facilidad variables irrelevantes.

V. CONCLUSIONES

En este trabajo se ha presentado una extensión de una técnica para la minimización de conjuntos de reglas difusas basado en la técnica de simplificación de circuitos digitales de Quine McCluskey. Tomando como base una versión previa que trabaja con circuitos binarios. Dicha versión ofrece un buen resultado en cuanto a la reducción del número de reglas,

pero no permite, de forma fácil, la detección de variables irrelevantes en los antecedentes de las reglas. Este problema se debe al uso de la codificación binaria. Cambiado a codificación ternaria y tomando los dominios con tres etiquetas lingüísticas, la técnica no sólo reduce las reglas, sino que adquiere mayor capacidad para reducir las condiciones en el antecedente de la regla.

Se ha realizado un estudio experimental, usando el algoritmo de CHI como método para extraer las reglas y comparando la versión binaria con la ternaria que se ha propuesto en este trabajo. Los resultados muestran que la nueva propuesta reduce el número de reglas a niveles semejantes a como lo hacía la binaria. Además, frente a la binaria, las bases de reglas obtenidas usando la minimización ternaria presentan un mayor grado de simplicidad cuando se combina la reducción de las reglas y el número de condiciones en el antecedente de las reglas.

REFERENCIAS

- [1] M. J. Gacto, R. Alcalá, y F. Herrera, «Interpretability of linguistic fuzzy rule-based systems: An overview of interpretability measures», *Information Sciences*, vol. 181, n.º 20, pp. 4340-4360, oct. 2011, doi: 10.1016/j.ins.2011.02.021.
- [2] M. Elkano, J. A. Sanz, E. Barrenechea, H. Bustince, y M. Galar, «CFM-BD: A Distributed Rule Induction Algorithm for Building Compact Fuzzy Models in Big Data Classification Problems», *IEEE Transactions on Fuzzy Systems*, vol. 28, n.º 1, pp. 163-177, ene. 2020, doi: 10.1109/TFUZZ.2019.2900856.
- [3] M. Karnaug, «The map method for synthesis of combinational logic circuits», *Transactions of the American Institute of Electrical Engineers, Part I: Communication and Electronics*, vol. 72, n.º 5, pp. 593-599, nov. 1953, doi: 10.1109/TCE.1953.6371932.
- [4] L.X. Wang y J. M. Mendel, «Generating Fuzzy Rules by Learning from Examples», vol. 22, n.º 6, pp. 14, 1992.
- [5] Z. Chi, H. Yan y T. Pham, T, «Fuzzy algorithms with applications to image processing and pattern recognition». World Scientific, 1996.
- [6] D. Álvarez Estevez y V. Moret Bonillo, «Revisiting the Wang-Mendel algorithm for fuzzy classification», *Expert Systems*, vol. 35, n.º 4, p. e12268, 2018, doi: 10.1111/essy.12268
- [7] L. Jara, A. González, y R. Pérez, «A preliminary study to apply the Quine McCluskey algorithm for fuzzy rule base minimization», en 2020 IEEE International Conference on Fuzzy Systems (FUZZ-IEEE), jul. 2020, pp. 1-6. doi: 10.1109/FUZZ48607.2020.9177739.
- [8] P. W. C. Prasad, A. Beg, y A. K. Singh, «Effect of Quine-McCluskey simplification on Boolean space complexity», en 2009 Innovative Technologies in Intelligent Systems and Industrial Applications, Kuala Lumpur, Malaysia, jul. 2009, pp. 165-170. doi: 10.1109/CITI-SIA.2009.5224219.
- [9] A. Dhande, V. Ingole, y V. Ghiye, *Ternary Digital System: Concepts and Applications*. 2014
- [10] M. Elkano, M. Galar, J. Sanz, y H. Bustince, «CHI-BD: A fuzzy rule-based classification system for Big Data classification problems», *Fuzzy Sets and Systems*, vol. 348, pp. 75-101, oct. 2018, doi: 10.1016/j.fss.2017.07.003.
- [11] H. Ishibuchi y T. Yamamoto, «Rule weight specification in fuzzy rule-based classification systems», *IEEE Trans. Fuzzy Syst.*, vol. 13, n.º 4, pp. 428-435, ago. 2005, doi: 10.1109/TFUZZ.2004.841738.
- [12] D. García, A. González, y R. Pérez, «Overview of the SLAVE learning algorithm: A review of its evolution and prospects», *International Journal of Computational Intelligence Systems*, vol. 7, n.º 6, pp. 1194-1221, nov. 2014, doi: 10.1080/18756891.2014.967008.
- [13] S.K. Petrick, «On the Minimization of Boolean Functions», in *Proceedings of the International Conference Information Processing*, pp. 422-423, 1959.
- [14] D. Dua, y C. Graff, *UCI Machine Learning Repository* [http://archive.ics.uci.edu/ml]. Irvine, CA: University of California, School of Information and Computer Science, 2019

Prometheus: Harnessing Fuzzy Logic and Natural Language for Human-centric Explainable Artificial Intelligence

Ettore Mariotti, Jose M. Alonso-Moral

Centro Singular de Investigación en Tecnoloxías Intelixentes (CiTIUS)
Universidade de Santiago de Compostela
 15782, Santiago de Compostela, SPAIN
 {ettore.mariotti, josemaria.alonso.moral}@usc.es

Albert Gatt

Institute of Linguistics and
Language Technology, University of Malta (UM)
 Msida, Malta
 albert.gatt@um.edu.mt

Abstract—Prometheus is an interpretable model which is suitable for the generation of visual and textual explanations grounded in common sense knowledge. This model can be seen as a special case of generalized additive models, which can be also interpreted as a list of (fuzzy) rules. The goodness of the model is illustrated with one benchmark dataset from the medical domain. Reported results are encouraging. They suggest that Prometheus exhibits a good balance between understandability and classification performance in comparison with other well-known models (e.g., linear models, decision trees or fuzzy rule-based classifiers) which are deemed as interpretable.

Index Terms—Explainable Artificial Intelligence, Interpretable Machine Learning, Shapley Values, Generalized Additive Models, Fuzzy Rule-based Classifiers

I. INTRODUCTION

In Greek mythology, Prometheus was a semi-god who stole a spark of fire from the gods and created humanity from clay [1]. Moreover, he is said to have brought fire (intelligence) to humankind. In this paper, we introduce a new Artificial Intelligence (AI) modelling method whose name takes inspiration from this myth, in that it constitutes a step ahead towards the creation of self-explaining AI models in real-world applications.

Given a classification or regression task, the objective of machine learning has been to mathematically formalise algorithms for maximising some accuracy measure for the given data. A wide variety of models have been created, some of them achieving impressive performance on benchmarking datasets. However, there is a class of problems for which having good performance metrics is not enough. In particular, a system that supports decision-making for fault-critical situations, or where someone has to take responsibility for the action, has limited utility if it does not provide any meaningful insight into its reasoning process. For example, if an advanced AI in a hospital suggests a diagnosis and a cure, the well-being of the patient is ultimately the responsibility of the doctor in charge and he/she is the person who in the end will have the final word about the diagnosis and treatment. Thus a doctor is going to use such an AI-based diagnostic system only if its recommendations are transparently explained, allowing

the professional to evaluate their soundness and possibly communicate them to the patient.

It is within this frame of responsibility, trust, fairness, accountability and liability that eXplainable AI (XAI) plays a key role [2]. An AI system that can explain its reasoning is more likely to inspire trust in decision-makers. Recently, XAI has received lots of attention, with proposals for different approaches and tools aimed to answer different questions. In this context, if we want to lower the barrier to access to these tools, then the use of the verbal modality and textual generation can sparkle an interactive communication between humans and AI that will help understandability and the transfer of knowledge [3]–[5]. This is because humans tend to vastly favour the modality of language when explaining their decisions and the reason behind them.

In this paper, we introduce Prometheus, an interpretable-by-design AI model that supports both visual explanations and textual descriptions of its reasoning. We have evaluated Prometheus with the Breast Cancer benchmark dataset [6]. Moreover, we have shown how Prometheus works in comparison with alternative methods in an illustrative example. The rest of the manuscript is organized as follows. Section II introduces related work. Section III describes how to build and interpret Prometheus XAI models. Section IV presents our experimental settings and reports achieved results. Section V concludes the paper and points out future work.

II. RELATED WORK

One important feature that distinguishes between different XAI approaches [7] is whether an explanation is generated after the prediction is done (post-hoc), or whether one tries to design an interpretable model in the first place.

A. Interpretable by design models

In the literature we find three broad classes of interpretable models; several variations exist within each class.

- **Rule-based Systems (RBSs):** a list of rules in the form IF X THEN Y. The more rules there are, the harder it becomes to understand the model.

- **Decision Trees (DTs):** a list of rules nested into a tree structure. At each step of the reasoning process, the model is looking at just one “test” at a time and based on the result decides which branch to descend along the tree until it reaches a leaf that corresponds to the prediction.
- **Generalized Additive Models (GAMs):** a class of models that are only concerned with first-order interactions between the features and the target. This allows the study of the effect of each feature one at a time.

Both RBSs and DTs are relatively easy to interpret and describe in natural language [8], [9]. However, as far as we know, it is rare to find linguistic explanations associated with GAMs (see e.g., [10]). GAMs are models of the form $g(y) = \sum_i f_i(x_i)$ where x_i represents the i -th feature of $\mathbf{x}$ and $g(\cdot)$ is the *link function* (e.g., logistic sigmoid for classification tasks). Each f_i is called a *shape function* [11]. GAMs are interpretable due to the additive nature of the modelling, allowing a user to focus on the contribution of each feature without accounting for interaction effects. In other words, they are a class of models for which partial dependency plots and their generalisation, Individual Conditional Expectations [12], can be computed exactly, they let users follow the process of computation in a sequential manner, and they are not subject to variation effects (because interactions are restricted to first-order, by design, but see below for details on GA2Ms which incorporate interaction terms). The main types of GAM are:

- **Linear Models (LMs):** where $f_i(x) = \beta_i$ and β_i is a static coefficient.
- **Scoring Systems:** similar to LMs, but where β coefficients are discrete and small [13].
- **Spline Systems:** each partial f_i is estimated with a spline.
- **GA2M:** introduces additional terms that explicitly deal with second-order interactions $f_i(x_j, x_k)$ which are described by heatmaps [14].
- **Explainable Boosting Machine (EBM)**¹: combines small trees and also deals with second-order interactions.
- **GAMs with Neural Networks:** each f_i is estimated with a neural network. They were first proposed by [15].

B. Post-hoc explanations

With black-box models, it is possible to some extent to generate post-hoc explanations of their behaviour, using different techniques. For example, LIME [16] makes a linear approximation to the black-box so that it is locally equivalent. On the other hand, SHAP [17] is inspired by the game-theoretic Shapley Values that assign to the feature x_j a contribution $\phi_j(x_j)$ such that

$$f(\mathbf{x}) - \mathbb{E}_{\mathbf{x}}[f(\mathbf{x})] = \sum_i \phi_i(x_i)$$

SHAP is in general NP-hard to compute, but for some models (like tree ensembles [18]) can be computed efficiently. Moreover, it is accessible and usable thanks to its open-source implementation.²

The interpretation of these methods must nonetheless be conducted carefully as they have been shown to be potentially vulnerable to some adversarial attacks [19].

III. DESCRIPTION OF THE PROMETHEUS XAI MODEL

The architecture of Prometheus is depicted in Fig. 1. This is another step toward a human-centric approach to knowledge discovery and human-machine interaction. The user can benefit from flexibility in the customisation of the processing of data according to common-sense expert knowledge. Also, he/she can exploit the insights learnt from the model communicated with visualisations and textual explanations.

The first component of Prometheus is a *discretiser* that discretises continuous variables into k -bins (using a k -means or quantiles strategy) transforming the original instance into a wider vector of 0s and 1s (this can be interpreted as a way of building rules IF x_i IS IN BIN j). Each bin has a semantic meaning that can be linguistically refined by an expert, just by tuning the associated strong fuzzy partition. As an example let's imagine the task of recommending a certain movie to someone given that we have a feature that represents their age: the discretiser would produce a tuple of features stating whether age is *low*, *medium* or *high*. This initial partition can be semantically refined later by the user with personalized labels like *young*, *mid-age* or *old*.

The output of the discretiser is then combined linearly by an LM with zero bias (regularised by an L1 penalty as it incentives a sparse representation)³. In this way, a partial score is attributed to each partition and the final prediction (in log-odds space) will be the global aggregation (sum) of the fired partial scores. Going back to our example, a *young* person could contribute with a score of +1.2 (thus representing evidence of the class “Do recommend”) while an *old* one could contribute with -1.1 (thus being counter-evidence for the class, thus “Do not recommend”). If the *mid-age* contributes with 0, then it can be omitted for the sake of reducing the cognitive load to the user.

The generated model can be seen both as a GAM where each f_i is piece-wise constant, and as a fuzzy rule list aggregated by the sum operator (e.g., “ R_i : IF *feature_i* is *label₃* THEN *output₁* with $w = s$ ”). This means that we can use the visualisations typically used for GAMs but also the linguistic explanations associated with fuzzy rule-based systems. In particular, the interpretation of the model as a rule list can support the generation of simple linguistic explanations. This can be achieved for example by ad-hoc mapping rules to templates of the form “Instance i has a high value of feature f and this contributes s to the total log-odd score of class C .” Naturally, the language generated could also be tailored to the user, or made more or less precise depending on the context. Thus, we can produce a textual description of the process that leads to the final prediction.

We adopt a relatively simple generation strategy for textual descriptions in the present work, noting that more sophisticated

¹Available at <https://github.com/interpretml/interpret>

²Available at <https://github.com/slundberg/shap>

³The number of bins and the coefficient for regularisation are automatically chosen with cross-validation on the training data.

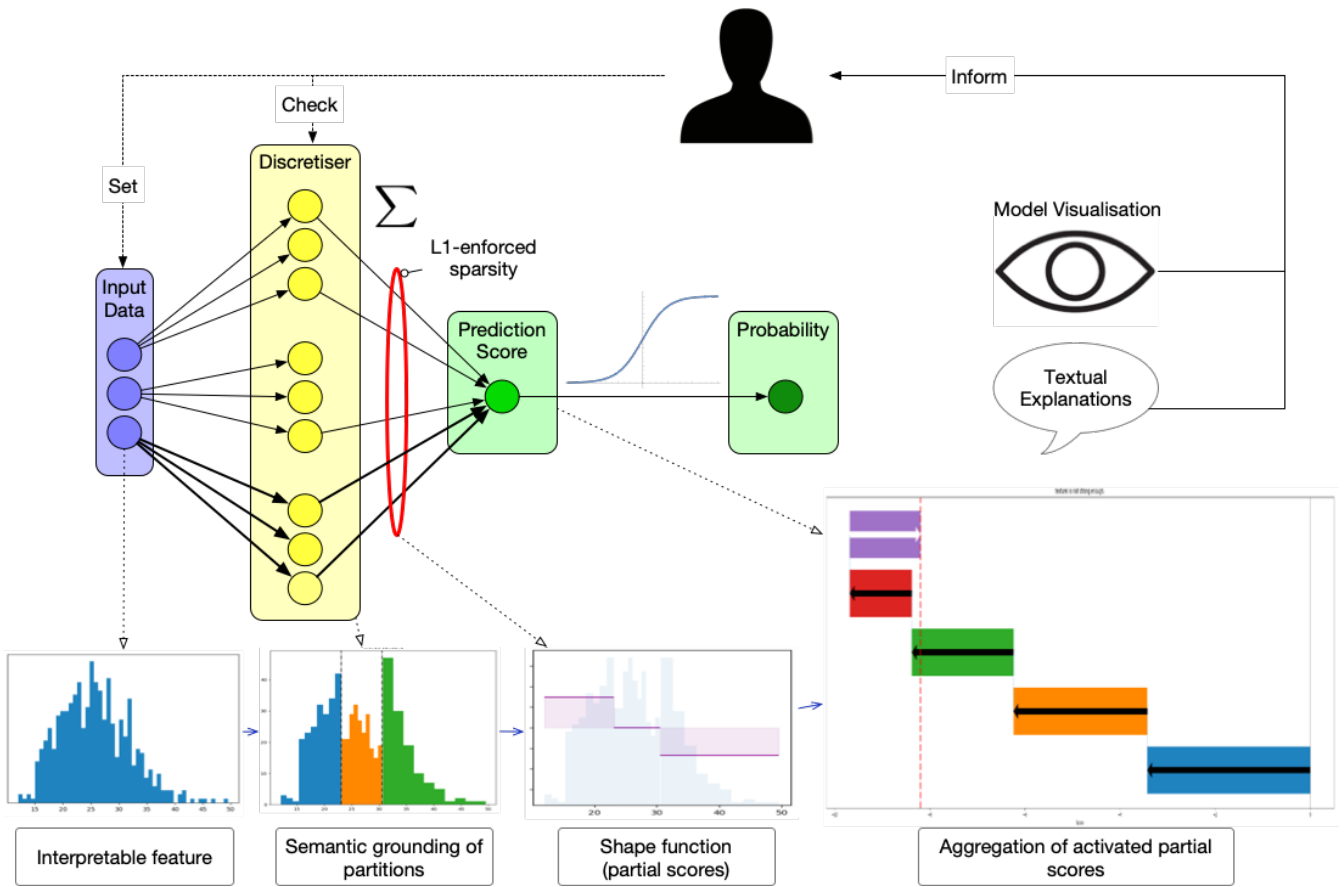


Fig. 1. A pictorial description of the proposed architecture.

methods have been developed in the field of Natural Language Generation [20]. Here, the textual generation process follows a template-based approach. Contributions are first divided into evidence and counter-evidence for the target class. They are then presented in an ordered way sorted by the absolute magnitude of their partial score so that more important features are presented first (see an illustrative example in Fig. 4, which will be described in depth in the next section). Then, the discretiser component gives us the option of using semantic labels to give a qualitative description of the magnitude of the feature. For example, if a feature had three bins (like in Fig. 1) we could establish that *for that particular feature* we assign one semantic label out of *[low, medium, high]*; with the meaning of this semantic label tied to the distribution of data. However, a given explanation can be excessively long due to many features contributing very little to the overall score. Accordingly, to keep the explanation shorter and relevant for the user, we have implemented a compression algorithm that groups small partial scores into a score called “other reasons”. The values are sorted by absolute value and then are progressively accumulated, until a threshold is met (while ensuring that the absolute magnitude of the resulting aggregation is smaller than the smallest remained partial score). The rate of

compression can be set by the user, allowing the user to focus first on what is most relevant and then giving him/her the option of delving into the details if needed.

It is important to note that even though the additive nature of the Prometheus model might qualitatively resemble SHAP, it is significantly different from it. This is because additive coefficients in Prometheus constitute in themselves the interpretable model. By contrast, SHAP provides these coefficients as a post-hoc explanation. Moreover Shapley values describe $f(\mathbf{x}) - \mathbb{E}_{\mathbf{x}}[f(\mathbf{x})]$ while for Prometheus they describe precisely $f(\mathbf{x})$. It is nonetheless possible to make use of the visualisation toolbox provided by the SHAP package (while keeping in mind that we are visualising something radically different).

IV. EXPERIMENTS

We have empirically assessed both the accuracy and complexity (as a proxy for explainability) of the Prometheus XAI model. As humans we have a limited cognitive capacity, thus between two systems of equal performances we would tend to find more interpretable a simple one. We considered the Breast Cancer dataset (taken from the UCI machine learning data repository [21]). This is a relatively small dataset (569 instances, 30 features) where explainability, as is expected in

medical diagnosis, is an important concern. All features are numerical and there are no missing values. The task consists of predicting if a given cancer is benign (357 instances) or malign (212 instances) [6].

For comparison purposes, in addition to Prometheus, the following models are generated⁴: (i) EBMs; (ii) LM-L1 (with L1 regularisation chosen with 4-fold cross-validation); (iii) J48 decision tree; (iv) REPTree; (v) RandomTree; (vi) Fuzzy Hoeffding Decision Tree (FHDT) (vii) Fuzzy Unordered Rule Induction Algorithm (FURIA); and (viii) Random Forest (RF).

Figures 2 and 3 show the interpretability-accuracy trade-off of models which are deemed as interpretable⁵. Every measure is estimated with stratified 10-fold cross-validation. Classification performance is measured in terms of the area under the receiver operating characteristic curve (ROC AUC) in Fig. 2. Classification performance is measured in terms of F1 Score in Fig. 3. In both cases, the structural complexity of models is computed as the total rule length (TRL)⁶. In these figures, Prometheus emerges as the most accurate model, at the cost of higher structural complexity. Moreover, it performs reasonably well (ROC AUC = 0.988; F1 Score = 0.949) compared to the black-box model RF (ROC AUC = 0.991; F1 Score = 0.968). However, EBM (ROC AUC = 0.992; F1 Score = 0.971) achieves even better performance, which can be attributed to its higher flexibility due to the high number of shallow trees, which allow the modelling of a more complex shape function, and to the capability of modelling higher-order interactions. In addition, this high performance comes at the cost of increased model complexity (TRL=7680), i.e., the user has to deal with a more complex shape function representation for inspecting the model, what in practice jeopardizes intelligibility.

Regarding explainability, the prediction of a single instance consists of the summation of partial scores activated when a feature lies in a certain bin. This can be displayed in different ways (see Fig. 4). For example, the list of activated rules is given on top, a waterfall plot is given in the centre, and a linguistic explanation is given at the bottom. In addition, Fig. 5 shows the same example with a higher compression rate, i.e., with a less detailed and rougher explanation.

On the other hand, Figs. 6, 7, and 8 show the Shapely values associated with models LM-L1, RF, and EBM, respectively. These graphs provide similar information to that given by Prometheus: the final output is laid down as a sum of terms. The main difference is that this sum of terms amounts to the difference between the expected value of the function and the specific value of the instance in SHAP, while the sum directly

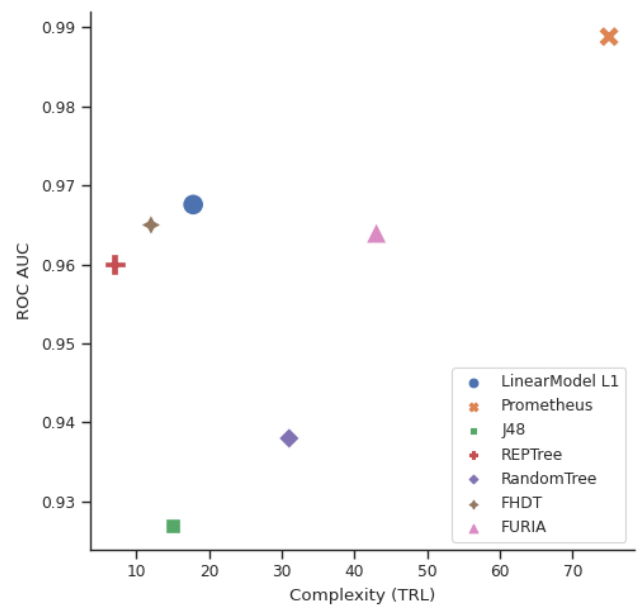


Fig. 2. The trade-off between classification performance (ROC AUC) and complexity (TRL); computed using stratified 10-fold cross-validation.

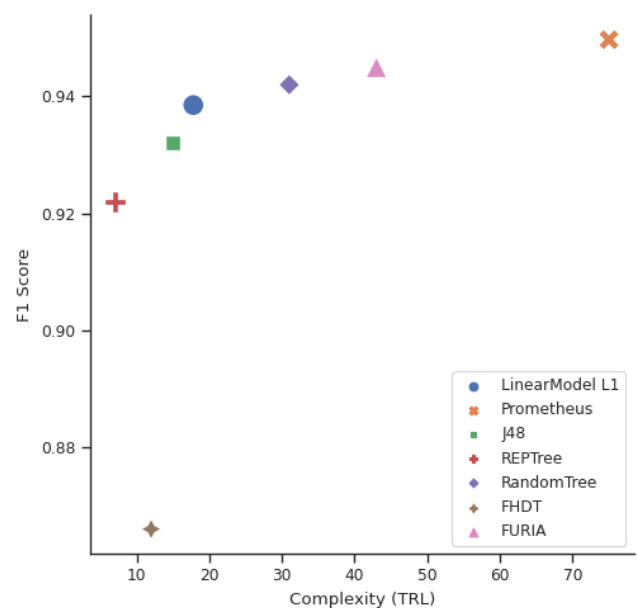


Fig. 3. The trade-off between classification performance (F1 Score) and complexity (TRL); computed using stratified 10-fold cross-validation.

⁴We use the implementation of EBMs and LM-L1 in Python InterpretML and scikit-learn packages, respectively. We use the implementation of J48, REPTree, RandomTree, FHDT, FURIA and RF in ExpliClas [8], [9].

⁵EBM is not included in the figures because it uses a discretiser with 128 bins what yields huge structural complexity (TRL=7680) and makes impossible any linguistic interpretation of the model.

⁶TRL counts the number of premises and conclusions in a rule list like the one provided by FURIA or FHDT; TRL counts the number of nodes in a tree (like J48, REPTree or RandomTree); TRL counts the number of non-zero coefficients multiplied by 2 in an LM; TRL equals 2 x number of bins x number of features in Prometheus and EBM.

explains the $f(x)$ in Prometheus. Moreover, it is important to notice how SHAP is a post-hoc method, thus it does not directly reflect the underlying computation. The explanation of Prometheus is instead exactly the description of its internal rationale. The explanations given by these different models show some inconsistencies, and this can be misleading and jeopardise user confidence. For example, *worst perimeter* is

Activated Rules

R1: IF high worst concave points THEN **malign** with weight 2.4
 R2: IF medium radius error THEN **malign** with weight 1.3
 R3: IF medium worst area THEN **malign** with weight 1.1
 R4: IF high mean radius THEN **malign** with weight 0.4
 R5: IF medium mean area THEN **malign** with weight 0.2
 R6: IF high mean smoothness THEN **malign** with weight 0.06
 R7: IF low mean texture THEN **benign** with weight 0.7
 R8: IF low worst texture THEN **benign** with weight 0.9

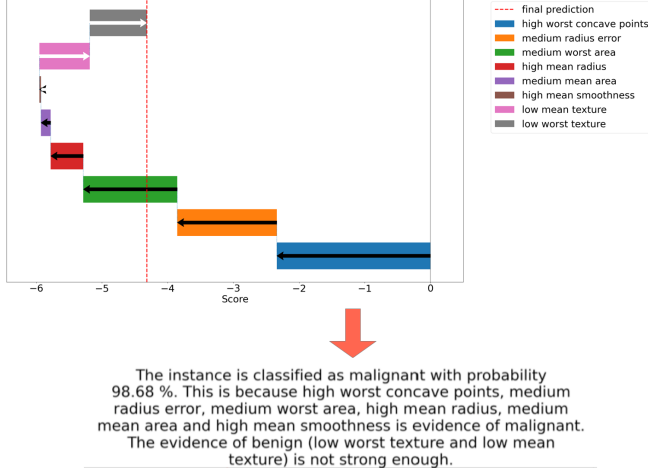


Fig. 4. A full local factual explanation provided by Prometheus.

Activated Rules

R1: IF high worst concave points THEN **malign** with weight 2.4
 R2: IF medium radius error THEN **malign** with weight 1.3
 R3: IF medium worst area THEN **malign** with weight 1.1
 R4: IF other reasons THEN **benign** with weight 0.9

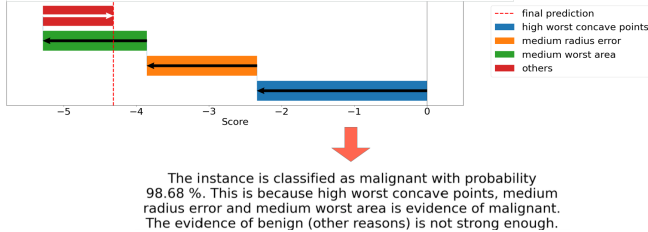


Fig. 5. Illustrative example of compressed explanation.

evidence for *benign* in case of LM-L1 but it is for *malign* in case of RF and EBM. Moreover, this feature is not highlighted among the most relevant ones by Prometheus.

In addition, the linguistic explanations provided by Prometheus can be compared to those given by ExpliClas:

- **J48**: Diagnosis is *malign* because *mean concavity* and *worst area* are *medium*.
- **REPTree**: Diagnosis is *malign* because *worst perimeter* is *medium*.
- **RandomTree**: Diagnosis is *malign* because *mean area*, *mean smoothness*, *area error* and *worst concavity* are *medium*.
- **FHDT**: We have a high confidence in the classification result. It is very likely that diagnosis is *malign*. There is also a medium chance that it is *benign*. On balance *malign* is more likely, because in accordance with rule 7 *concavity error* is *low* and *worst concave points* is *high*.
- **FURIA**: We have a high confidence in the classification

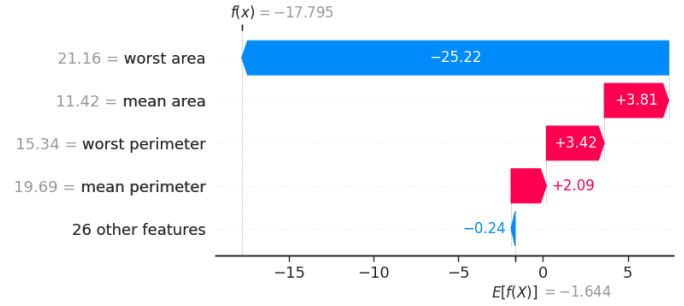


Fig. 6. Local explanation given by SHAP for LM-L1 model.

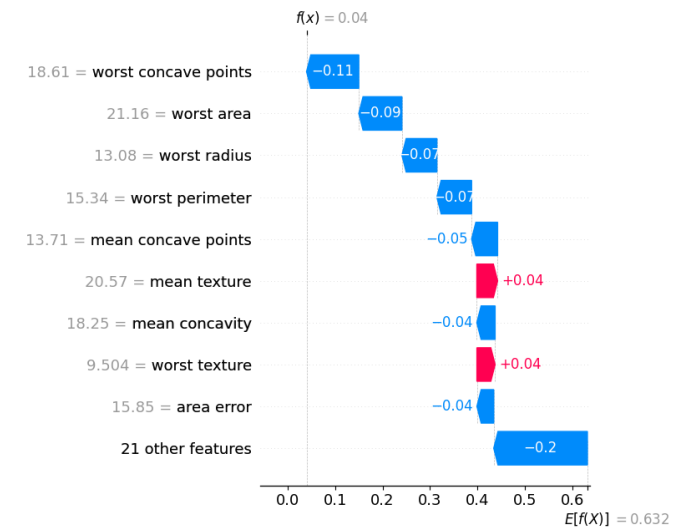


Fig. 7. Local explanation given by SHAP for RF model.

result. It is very likely that diagnosis is *malign* because *worst concave points* is *high* and *worst radius* is *medium*.

As we can see above, different models are supported by different rationale and provide consistent but different complementary explanations. In the case of fuzzy classifiers (FHDT and FURIA), ExpliClas begins explanations with a sentence related to the confidence in the classification result which emerges from the rule firing degree associated with the winner rule. This is similar to how Prometheus begins each textual explanation (see Fig. 5) with a sentence verbalising the probability associated with the given classification.

ExpliClas also provides users with factual explanations of the inferred class but nothing is said about alternative classes. This is because explanations verbalise the information contained in the activated branch of a tree or the winner rule in a fuzzy classifier. On the contrary, Prometheus provides users with contrastive explanations, which verbalise which features are evidence for and against the inferred class.

Another commonality shared by ExpliClas and Prometheus is that they provide users with a graphical output along with the list of related rules. In the case of the illustrative example under consideration, FURIA fires rules:

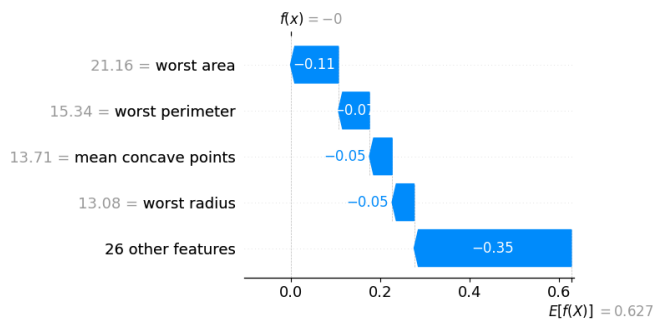


Fig. 8. Local explanation given by SHAP for EBM model.

- R1: IF *worst radius* in [16.25, 16.82, inf, inf] and *worst concave points* in [0.1452, 0.1456, inf, inf] THEN *Diagnosis* is *malign* with CF = 0.99
- R2: IF *worst area* in [947.9, 967, inf, inf] and *worst fractal dimension* in [0.06469, 0.06515, inf, inf] THEN *Diagnosis* is *malign* with CF = 0.99

However, FURIA rules are not easily verbalised, so they need to be linguistically approximated by ExpliClas as:

- R1: IF *worst radius* is *medium* and *worst concave points* is *high* THEN *Diagnosis* is *malign*
- R2: IF *worst area* is *medium* and *worst fractal dimension* is *medium* THEN *Diagnosis* is *malign*

In the case of ExpliClas, explanations only pay attention to the winner rule (R1 in this example; because the two rules are activated with the same degree and then the first rule is selected). On the contrary, all listed rules contribute to the explanation elaborated by Prometheus.

V. CONCLUSIONS

In this paper, we have introduced Prometheus, an interpretable model that can produce both textual and visual explanations related to tabular data. This model can be seen as a weighted RBS with sum aggregation at inference level. We measured reasonable performances on classification metrics, beating LMs and DTs. In addition, the model compares favourably to RFs and EBMs, though its accuracy is slightly lower. Prometheus was designed for supporting linguistic explanations with different degrees of details. As future work, we plan exhaustive experimentation with more benchmark datasets. In addition, we plan to explore more sophisticated natural language generation techniques, multiclass classification and measuring the impact of explainability with intrinsic and extrinsic human evaluation. For the sake of reproducibility, Prometheus (as well as other related XAI tools) is available at <https://gitlab.citius.usc.es/jose.alonso/xai>.

ACKNOWLEDGEMENTS

Jose M. Alonso-Moral is a *Ramon y Cajal* Researcher (RYC-2016-19802). This work is funded by the Spanish Ministry of Science, Innovation and Universities (grant RTI2018-099646-B-I00), the Galician Ministry of Education, University

and Professional Training (grants ED431F 2018/02, ED431C 2018/29, and ED431G2019/04), and the NL4XAI project which has received funding from the European Union's Horizon 2020 research and innovation programme under the Marie Skłodowska-Curie grant agreement No 860621.

REFERENCES

- [1] E. Beall, "Hesiod's prometheus and development in myth," *Journal of the History of Ideas*, vol. 52, no. 3, pp. 355–371, 1991.
- [2] T. Miller, "Explanation in artificial intelligence: Insights from the social sciences," *Artificial Intelligence*, vol. 267, pp. 1–38, 2019.
- [3] J. M. Alonso et al., "Interactive natural language technology for explainable artificial intelligence," in *1st Workshop on Foundations of Trustworthy AI integrating Learning, Optimisation and Reasoning (TAILOR)*, *European Conference on Artificial Intelligence (ECAI)*, *LNAI 12641*. Springer, 2021, pp. 63–70.
- [4] E. Mariotti, J. M. Alonso, and A. Gatt, "Towards Harnessing Natural Language Generation to Explain Black-box Models," in *2nd Workshop on Interactive Natural Language Technology for Explainable Artificial Intelligence*. Dublin, Ireland: ACL, 2020, pp. 22–27.
- [5] J. M. Alonso, C. Castiello, L. Magdalena, and C. Mencar, *Explainable Fuzzy Systems. Paving the Way from Interpretable Fuzzy Systems to Explainable AI Systems*. Springer International Publishing, 2021.
- [6] M. Patricio et al., "Using Resistin, glucose, age and BMI to predict the presence of breast cancer," *BMC Cancer*, vol. 18, pp. 1–8, 2018.
- [7] A. Barredo Arrieta et al., "Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI," *Information Fusion*, vol. 58, pp. 82–115, 2020.
- [8] J. M. Alonso and A. Bugarín, "Expliclas: Automatic generation of explanations in natural language for weka classifiers," in *IEEE International Conferences on Fuzzy Systems*, 2019, pp. 1–6.
- [9] J. M. Alonso, P. Ducange, R. Pecori, and R. Vilas, "Building explanations for fuzzy decision trees with the ExpliClas software," in *IEEE World Congress on Computational Intelligence*, 2020, pp. 1–8.
- [10] F. Hohman, A. Srinivasan, and S. M. Drucker, "TeleGam: Combining Visualization and Verbalization for Interpretable Machine Learning," in *IEEE Visualization Conference (VIS)*, 2019, pp. 151–155.
- [11] Y. Lou, R. Caruana, and J. Gehrke, "Intelligible models for classification and regression," in *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, 2012, pp. 150–158.
- [12] A. Goldstein et al., "Peeking Inside the Black Box: Visualizing Statistical Learning With Plots of Individual Conditional Expectation," *Journal of Computational and Graphical Statistics*, vol. 24, pp. 44–65, 2015.
- [13] B. Ustun and C. Rudin, "Supersparse Linear Integer Models for Optimized Medical Scoring Systems," *Machine Learning*, vol. 102, no. 3, pp. 349–391, 2016.
- [14] Y. Lou et al., "Accurate intelligible models with pairwise interactions," in *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, 2013, pp. 623–631.
- [15] W. J. E. Potts, "Generalized additive neural networks," in *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, 1999, pp. 194–200.
- [16] M. T. Ribeiro, S. Singh, and C. Guestrin, "Why Should I Trust You?: Explaining the Predictions of Any Classifier," in *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, 2016, pp. 1135–1144.
- [17] S. M. Lundberg and S.-I. Lee, "A Unified Approach to Interpreting Model Predictions," in *Advances in Neural Information Processing Systems (NIPS)*, 2017, pp. 1–10.
- [18] S. M. Lundberg et al., "From local explanations to global understanding with explainable AI for trees," *Nature Machine Intelligence*, vol. 2, pp. 56–67, 2020.
- [19] D. Slack et al., "Fooling LIME and SHAP: Adversarial Attacks on Post hoc Explanation Methods," in *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, 2020, pp. 180–186.
- [20] A. Gatt and E. Krahmer, "Survey of the state of the art in natural language generation: Core tasks, applications and evaluation," *Journal of Artificial Intelligence Research*, vol. 61, pp. 65–170, 2018.
- [21] D. Dua and C. Graff, "UCI Machine Learning Repository," 2017, University of California, Irvine, School of Information and Computer Sciences. [Online]. Available: <http://archive.ics.uci.edu/ml>

Modelando el comportamiento del consumidor con el modelo de 2-tuplas lingüísticas y una heurística basada en el conocimiento de marca

Jesús Giráldez-Cru*, Manuel Chica*[†], Oscar Cordon* and Francisco Herrera*

* *Departamento de Ciencias de la Computación e Inteligencia Artificial (DECSAI)
Instituto Andalúz Interuniversitario en Data Science and Computational Intelligence (DaSCI)
Universidad de Granada*

{jgiralde,manuelchica}@ugr.es {ocordon,herrera}@decsai.ugr.es

[†]*School of Electrical Engineering and Computing, The University of Newcastle (Australia)*

Resumen—Los modelos basados en agentes (ABM) son un paradigma de simulación para modelar sistemas complejos mediante la definición del comportamiento heterogéneo de cada uno de sus individuos en una aproximación *bottom-up*. En este trabajo, se presenta un ABM para marketing donde las percepciones de los consumidores son modeladas usando el modelo lingüístico 2-tuplas. Estas variables representan las opiniones que los consumidores tienen sobre las diferentes características de cada producto, las cuales guían sus decisiones (p. ej., precio o calidad). Al contrario que los valores exactos, las variables lingüísticas difusas son una representación realista de estos aspectos cualitativos. En nuestro ABM, los agentes usan una heurística de toma de decisiones para seleccionar un producto, la cual está basada en estas percepciones y en una regla probabilística de maximización de utilidad. Este proceso requiere una agregación difusa de las percepciones de cada producto, guiada por una serie de pesos asociados a los *drivers* del agente consumidor e implementada mediante el operador promedio ponderado de 2-tuplas. Adicionalmente, hay productos en el mercado que no son conocidos por algunos consumidores. En nuestro ABM, esta información se modela aplicando un filtro de conocimiento de marca en la heurística de toma de decisiones. De esta manera, los agentes consumidores sólo pueden elegir aquellos productos que les son conocidos. Los resultados de nuestro análisis experimental muestran que nuestra representación realista de las percepciones del consumidor es más precisa que otros métodos existentes.

Index Terms—comportamiento del consumidor, marketing, heurísticas de compra, reconocimiento de marca, información y toma de decisiones lingüística difusa, 2-tuplas, modelos basados en agentes.

I. INTRODUCCIÓN

El principal objetivo cuando se modela un mercado es entender las reglas que lo gobiernan, para analizar posteriormente el resultado de posibles escenarios. Por tanto, es crucial entender y predecir las compras de los consumidores. En los enfoques clásicos, estas decisiones se infieren comúnmente de variables globales en un esquema *top-down*. El principal inconveniente de este paradigma es la incapacidad de representar comportamientos heterogéneos de los consumidores, y los eventos emergentes de éstos. Por esta razón, estos enfoques clásicos normalmente resultan en representaciones imprecisas de la realidad del mercado [1], [2].

Un enfoque alternativo a estos enfoques clásicos consiste en estudiar el comportamiento complejo del mercado como resultado de una agregación *bottom-up* de las decisiones de cada uno de sus consumidores [2], [3]. Para ello, los modelos basados en agentes (ABM) [4], [5] proveen una infraestructura adecuada. ABM es una técnica de modelado descriptivo (una agregación de muchas decisiones individuales de cada agente) que ayuda al diseñador y al profesional de marketing a entender mejor el mercado y su comportamiento. En la mayoría de los casos, modelar comportamientos individuales es más simple y más preciso que modelar el comportamiento del sistema completo por reglas globales. Los ABM se han aplicado con éxito en otras áreas tan diversas como economía [6], política [7] y sistemas de confianza social [8].

La mayoría de los ABM para marketing existentes representan las opiniones del consumidor usando valores numéricos [9]. Ésta es una representación poco realista de este tipo de información cualitativa. Además, las opiniones del consumidor se definen normalmente usando datos de cuestionarios disponibles para la compañía, los cuales son comúnmente respondidos en términos lingüísticos. Por tanto, manejar valores numéricos requiere el preproceso adicional para transformar las respuestas lingüísticas de los cuestionarios en datos exactos, con lo que el diseño del modelo es más complejo y suele provocar una pérdida de información en la representación.

En este trabajo, se desarrolla un ABM para modelar mercados virtuales con una representación realista de las percepciones del consumidor, la cual está basada en variables lingüísticas difusas [10]–[12]. Esta información permite dar valores difusos (p.ej., a cada una de las características de los productos (p.ej., precio, calidad o gusto) [13]. En concreto, representamos estas percepciones usando el modelo lingüístico 2-tuplas [14], que consisten en un par formado por una etiqueta lingüística y un desplazamiento simbólico. Como se explica en [13], esta representación solventa los inconvenientes de otros enfoques lingüísticos difusos ordinales [15] ya que permite asignar valores diferentes a dos variables que tengan la misma etiqueta lingüística (teniendo dos valores diferentes en el desplazamiento simbólico). Además, esta representación

es sustancialmente más realista que otras representaciones existentes y estándares como los valores numéricos u otro tipo de valores exactos.

Dado que cada agente tiene una percepción sobre cada característica de los productos del mercado, se propone una heurística de toma de decisiones difusa que agrega dichas percepciones lingüísticas difusas para calcular una valoración global de cada producto para el consumidor y elegir uno de ellos, lo cual simula el proceso de compra de los consumidores. Dado que las percepciones del consumidor se representan con variables lingüísticas difusas y tienen un vector de pesos asociado a cada variable (*driver*), su agregación se puede realizar con el operador promedio ponderado para 2-tuplas [14].

De este modo, el ABM para el análisis de comportamiento de consumidor con representación lingüística difusa presentado en [13] representa adecuadamente varias características del comportamiento real del consumidor, pero no considera mecanismos más complejos como procesos de boca a boca en redes sociales u otro tipo de interacciones complejas entre agentes. Para avanzar en el diseño de un ABM realista para marketing, en este trabajo se introducen dos importantes modificaciones con respecto a dicho modelo. Primero, las decisiones del consumidor raramente son completamente deterministas. En nuestro sistema, este comportamiento se modela por medio de una función probabilística de maximización de la utilidad proporcional a las valoraciones agregadas de los productos, a diferencia de la estrategia plenamente determinista usada en [13]. Segundo, los consumidores no siempre conocen todas las marcas en el mercado. En el ABM propuesto en este trabajo, se incluye esta información definiendo las marcas de las que cada agente tiene conocimiento. Este tipo de información está disponible normalmente en datos de encuestas sobre el consumidor facilitado por agencias de marketing. Nuestra heurística de toma de decisiones implementa un filtro de conocimiento de marca, de forma que los agentes consumidores sólo pueden comprar productos que conocen.

El modelo propuesto se evalúa nuestro sistema en un caso de estudio de marketing real, comparando su rendimiento con el de otras representaciones tradicionales de las percepciones del consumidor, analizando su precisión con respecto a datos reales que representan las ventas en este mercado. El análisis experimental muestra que nuestro modelo es más preciso.

El resto de este trabajo se organiza así: en la Sec. II se describen conceptos preliminares; el ABM propuesto se presenta en Sec. III; en la Sec. IV se presenta una evaluación experimental del modelo; y finalmente se concluye en la Sec. V.

II. CONTEXTO

II-A. Representación numérica de las preferencias del consumidor usando datos de encuestas

Los sistemas ABM para marketing normalmente requieren de una definición de las percepciones del consumidor para cada marca del mercado, con el fin de simular un mercado virtual realista. Estas percepciones se obtienen de datos de encuestas y estudios de salud de marca de consumidores

reales, facilitados por renombradas consultoras de marketing como Kantar Millward-Brown [16]. Estos datos se estructuran en forma de encuestas incluyendo conjuntos de respuestas a una serie de preguntas sobre las marcas [17].

En algunos casos, las respuestas de las encuestas se recogen directamente en la misma escala requerida por el ABM, un valor real en el intervalo $[0, 10]$, representando 0 la percepción más negativa, 5 una percepción neutral, y 10 la percepción más positiva. Sin embargo, la situación más habitual es que las respuestas de las encuestas muestren una naturaleza lingüística. En estos casos, se requiere un preprocesado manual de las respuestas para traducirlas a la escala $[0, 10]$, con la consecuente *pérdida de información*. Por tanto, este problema se aborda mejor trabajando directamente con las valoraciones lingüísticas siguiendo un enfoque lingüístico difuso en lugar de transformarlas a valores numéricos.

II-B. Variables lingüísticas

Las variables lingüísticas [10]–[12] toman como valores palabras o frases en el lenguaje natural. Se usan en *enfoques lingüísticos difusos*, donde el problema requiere manejar aspectos cualitativos [15]. Éste es un requerimiento típico en muchos contextos, donde la representación de la información más directa y realista es claramente el lenguaje natural.

En el *enfoque lingüístico difuso ordinal*, un caso especial de enfoques lingüísticos difusos, las variables lingüísticas toman valores de un conjunto predefinido y totalmente ordenado de *etiquetas lingüísticas* $S = \{s_0, \dots, s_g\}$ de tamaño finito $|S| = g + 1$. Se considera la definición convencional de conjunto ordenado donde $\forall s_i, s_j \in S. s_i \leq s_j \Leftrightarrow i \leq j$.

En este trabajo, consideramos funciones de pertenencia triangulares para las variables lingüísticas. En la Figura 1 se representa un ejemplo de esta pertenencia difusa. En concreto, se usa una función de pertenencia triangular definida en el intervalo $[0, 1]$ para un conjunto de 5 etiquetas lingüísticas $\{muyMala, mala, neutra, buena, muyBuena\}$ de la variable lingüística.

El manejo de variables lingüísticas difusas normalmente requiere agregar su información, es decir, los valores de estas variables. Un enfoque común es aproximar estos valores lingüísticos con valores numéricos correspondientes a su índice en el conjunto ordenado de etiquetas $\{0, 1, \dots, g\}$, agregarlos con métodos comunes, obteniendo un valor real intermedio $\beta \in [0, g]$, y finalmente aproximar este valor intermedio a una etiqueta lingüística del conjunto original [18]. Para ello, se definen dos operadores para aproximar variables lingüísticas a números y *vice versa*:

Definición 1 (Aproximaciones Numérico-Lingüísticas del modelo simbólico lingüístico computacional ordinal). *Dado un número real $\beta \in [0, g]$ y una etiqueta lingüística s_k , siendo k la posición de la etiqueta s_k en el conjunto ordenado de etiquetas lingüísticas $S = \{s_0, \dots, s_g\}$ del que la variable lingüística v toma valores, se definen las siguientes*

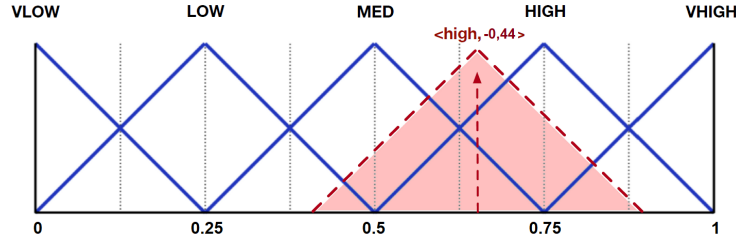


Figura 1. Función de pertenencia triangular para variables lingüísticas en el intervalo $[0, 1]$ para el conjunto de etiquetas lingüísticas $\{muyMala, mala, neutra, buena, muyBuena\}$. Se incluye también la 2-tupla correspondiente a la aproximación $\Delta(2,56) = \langle buena, -0,44 \rangle$.

aproximaciones numérico-lingüísticas $\Delta' : S \rightarrow \{0, \dots, g\}$ y $\Delta : [0, g] \rightarrow S$ como:

$$\Delta'(s_k) = k$$

$$\Delta(\beta) = s_k \text{ s.t. } \forall s_i \in S \wedge i \neq k \quad |\Delta'(s_k) - \beta| < |\Delta'(s_i) - \beta| \vee \exists s_i \in S \wedge i = k+1 \quad |\Delta'(s_k) - \beta| = |\Delta'(s_i) - \beta|$$

En resumen, el valor numérico asociado a una etiqueta lingüística s_k corresponde con su posición en el orden del conjunto y la etiqueta asociada al valor numérico β es la más cercana de acuerdo a los valores numéricos de las etiquetas en la escala $[0, g]$. En el caso de que β quede justo en el punto medio entre dos etiquetas, por convenio se asigna a la menor en el orden.

II-C. El operador de agregación promedio ponderado para variables lingüísticas

Existe una amplia variedad de operadores de agregación en la literatura que han sido adaptados para trabajar con variables lingüísticas [18], destacando el operador OWA [19] debido a su gran versatilidad. En nuestro caso, cada agente consumidor tiene un peso asociado a cada *driver*, previamente especificado, que representa sus preferencias de compra, por lo que emplearemos el operador promedio ponderado.

Definición 2 (Promedio ponderado para variables lingüísticas [18]). Sean $(a_1, \dots, a_m)$ una serie de etiquetas lingüísticas a agregar, con $a_i \in S$, y sea $W = \{w_1, \dots, w_m\}$ el conjunto de pesos asociado. El operador promedio ponderado de variables lingüísticas se define como:

$$\phi(A, W) = \Delta\left(\frac{\sum_{i=1}^m \Delta'(a_i) \cdot w_i}{\sum_{i=1}^m w_i}\right)$$

donde las funciones Δ y Δ' son las aproximaciones numérico-lingüísticas definidas anteriormente.

En nuestro caso el vector de pesos W cumple que $w_i \in [0, 1]$ y $\sum_i w_i = 1$, por lo que tenemos una combinación lineal y el denominador vale 1.

II-D. El modelo de representación de 2-tuplas lingüísticas

Un inconveniente del enfoque anterior es la pérdida de información causada por la agregación de etiquetas lingüísticas. Para resolver este problema, en [14] se propuso la *representación de 2-tuplas lingüísticas*.

Definición 3 (2-tupla lingüística [14]). Una 2-tupla lingüística es un par $\langle s_k, \alpha \rangle$, en el que $s_k \in S$ es una etiqueta lingüística y $\alpha \in [-0,5, 0,5]$ es un desplazamiento simbólico que especifica la traslación de la función de pertenencia que representa la etiqueta lingüística s_k más cercana si la información lingüística resultante de un cálculo simbólico no corresponde exactamente con una etiqueta del conjunto. El conjunto de 2-tuplas asociado con S se define como $\bar{S} = S \times [-0,5, 0,5]$.

Definición 4 (Aproximaciones Numérico-Lingüísticas de 2-tuplas lingüísticas [14]). Dado un número real $\beta \in [0, g]$ y una etiqueta lingüística $s_k \in S$, se definen las siguientes aproximaciones numérico-lingüísticas $\Delta' : S \rightarrow \{0, \dots, g\}$ y $\Delta : [0, g] \rightarrow \bar{S}$ como:

$$\Delta'(\langle s_k, \alpha \rangle) = k + \alpha$$

$$\Delta(\beta) = \langle s_k, \alpha \rangle, \text{ donde } k = \text{round}(\beta) \text{ y } \alpha = \beta - k$$

La función *round* devuelve el entero $k \in \{0, \dots, g\}$ más cercano a β .

De este modo, una 2-tupla lingüística se vincula a un valor numérico equivalente β en el intervalo de granularidad de S , $[0, g]$, y ese valor se obtiene a partir del valor numérico de la etiqueta y el del desplazamiento. Por otro lado, un valor numérico β se asocia a una 2-tupla compuesta por la etiqueta lingüística más cercana según los valores numéricos de las etiquetas en la escala $[0, g]$ y al desplazamiento necesario para hacer coincidir el valor de dicha etiqueta con el de β .

Por tanto, el operador de agregación promedio ponderado de la Definición 2 se puede aplicar directamente a 2-tuplas lingüísticas usando las aproximaciones numérico-lingüísticas de la Definición 4 [14]. Un ejemplo de estas aproximaciones se muestra en la Figura 1, donde se ve como $\Delta(2,56) = \langle buena, -0,44 \rangle$.

III. MODELO BASADO EN AGENTES CON INFORMACIÓN LINGÜÍSTICA DIFUSA Y CONOCIMIENTO DE MARCA

En esta sección se define el ABM que simula el comportamiento de un mercado con percepciones del consumidor lingüísticas y una heurística de toma de decisión lingüística. En nuestro modelo, usamos variables lingüísticas difusas para representar los diferentes aspectos de cada marca o producto (p.ej., *precio, calidad, confort, ...*). A estos aspectos se les llama *drivers* dado que guían las elecciones de los

consumidores. En concreto, usamos 2-tuplas lingüísticas para almacenar las percepciones sobre los *drivers* que gobiernan el comportamiento del mercado.

En nuestro modelo, los agentes representan consumidores que llevan a cabo un proceso de toma de decisiones para seleccionar un producto entre las marcas disponibles. Este proceso se lleva a cabo según sus propias percepciones y la valoración que le dan a cada marca. La población de agentes se organiza en segmentos, grupos de agentes muy similares en términos de comportamiento. Todo ello nos permite simular el comportamiento del mercado y realizar predicciones en él.

III-A. Marcas

En nuestro ABM, las estrategias de toma de decisión de los consumidores únicamente se llevan a cabo entre un conjunto finito $B = \{b_1, \dots, b_n\}$ de n marcas disponibles en el mercado. Para modelar los atributos de cada marca, se considera también un conjunto $D = \{d_1, \dots, d_m\}$ de m *drivers*. Estos *drivers* son fijos para todas las marcas en el mercados.

III-B. Percepciones del consumidor

Todo consumidor está representado por un agente en el sistema. Cada agente tiene sus propias percepciones (positiva, neutra o negativa) sobre cada *driver* de cada marca. Para representar las preferencias de los *drivers*, se define para cada agente x un vector de pesos $W^x = [w_1^x, \dots, w_m^x]$, tal que todos los pesos deben estar en el intervalo $[0, 1]$ y su suma debe ser igual a 1. Estos pesos representan la importancia de cada *driver* cuando un agente consumidor x realiza una decisión de compra.

Las percepciones del consumidor se modelan definiendo, para cada agente x del ABM, una matriz de percepciones P^x de dimensiones $n \times m$, donde cada elemento $p_{i,j}^x \in P^x$ representa la percepción del agente x sobre el *driver* $d_j \in D$ de la marca $b_i \in B$. En nuestro modelo, estas percepciones se representan usando 2-tuplas lingüísticas, todas ellas tomando valores de un conjunto ordenado de etiquetas lingüísticas común (ver la Definición 3). Esto nos permite representar la visión cualitativa del consumidor sobre cada marca.

III-C. El conocimiento de marca del consumidor

En un mercado real, cada consumidor puede no tener conocimiento de ciertas marcas. Para modelar esta información, se define en nuestro modelo el conocimiento de marca de cada agente. En concreto, para cada agente x , se define un vector A^x de n variables Booleanas, donde $a_i^x \in A^x$ representa si el agente x tiene conocimiento de la marca $b_i \in B$.

III-D. Agentes consumidores

Basados en la caracterización de las marcas y de los consumidores presentada anteriormente, ahora pasamos a definir los agentes consumidores. Nótese que esta definición representa el estado mental de los consumidores, es decir, su conocimiento sobre el mercado y sus percepciones sobre los productos disponibles.

Definición 5 (Agente consumidor). *Un agente consumidor x se define como la tupla $\langle A^x, W^x, P^x \rangle$, donde A^x , W^x y P^x son respectivamente el vector de conocimiento de marca, el vector de pesos de *drivers* y la matriz de percepciones en 2-tuplas lingüísticas descritos anteriormente.*

III-E. Heurística de toma de decisión

El proceso de toma de decisión difusa de cada agente consiste en seleccionar una de las marcas disponibles en el ABM en función de sus percepciones sobre los *drivers*. Estas decisiones simulan decisiones de compras de los consumidores en el mercado. El proceso se divide en dos pasos: (i) la agregación de las valoraciones de cada marca, y (ii) la selección de una marca. En el primer paso, el agente necesita agregar las percepciones de todos los *drivers* de cada marca. Esta agregación se calcula usando el operador de agregación promedio para 2-tuplas de la forma:

Definición 6 (Valoración de marca). *Dados un agente consumidor x y una marca b_i , se define la valoración de marca $as(x, b_i)$ como la agregación de sus percepciones para este marca calculadas con el operador de agregación promedio ϕ :*

$$as(x, b_i) = \phi(P_i^x, W^x) = \Delta(W^x \cdot \Delta'((P_i^x)^T))$$

donde $P_i^x = [p_{i,1}^x, \dots, p_{i,m}^x]$ es la fila i -ésima de la matriz P^x , W^x es el vector de pesos de *drivers* del agente x , y Δ y Δ' son las funciones de aproximación numérico-lingüísticas para 2-tuplas de la Definición 4¹.

El segundo es la selección de una marca en función de la valoración que el agente hace sobre cada marca. En este trabajo se usa una función de maximización de la utilidad probabilística $maxUtil^P$, que asigna a cada marca una probabilidad proporcional a su valoración y elige aleatoriamente una marca utilizando una ruleta probabilística. Se usa esta función no determinista inspirada por trabajos previos en análisis de marketing y comportamiento del consumidor [20]–[22].

Definición 7 (Maximización de la utilidad probabilística). *Para un agente consumidor x , la función de maximización de utilidad probabilística $maxUtil^P$ es la función probabilística de selección de marca que elige aleatoriamente una marca usando una ruleta probabilística con las siguientes probabilidades p_x para cada marca $b_i \in B$:*

$$p_x(b_i) = e^{\Delta'(as(x, b_i))} = e^{W^x \cdot \Delta'((P_i^x)^T)}$$

Sin pérdida de generalidad, estas probabilidades no normalizadas se pueden normalizar simplemente como $p_x(b_i) = p_x(b_i) / \sum_{b_j \in B} p_x(b_j)$.

IV. EVALUACIÓN EXPERIMENTAL DEL MODELO

En esta sección se presenta una evaluación experimental del ABM para marketing con representación lingüística difusa en un caso de estudio real. En concreto, se presenta una comparativa entre nuestro ABM y otros dos modelos que

¹Por simplicidad, se han sobrecargado las funciones Δ y Δ' para vectores de la forma: $\Delta(B) = [\Delta(b_i)]_{1 \leq i \leq n}$ y $\Delta'(A) = [\Delta'(a_i)]_{1 \leq i \leq n}$.

únicamente difieren en la representación de las percepciones (el resto del modelo se mantiene sin ningún otro cambio).

Por una parte, se compara nuestro ABM frente a un modelo con representación numérica de las percepciones, en el intervalo $[0, 10]$. Por otra, se compara nuestro ABM frente a un modelo simbólico lingüístico ordinal donde las percepciones se representan usando etiquetas lingüísticas (usando el mismo conjunto de etiquetas que en las 2-tuplas).

Los resultados de estos tres modelos se validan usando datos reales que representan las ventas en un mercado concreto y se analiza el rendimiento de los tres modelos midiendo su precisión en términos de predicción en ventas (agregadas para todos los agentes consumidores en el ABM).

IV-A. Descripción de las condiciones de simulación del ABM

Para la inicialización de los agentes se usan segmentos de consumidores, es decir, grupos de consumidores con un comportamiento *similar*. En nuestro ABM, todos los agentes del mismo segmento se caracterizan por los mismos pesos W^x y similares percepciones P^x , generadas aleatoriamente siguiendo una distribución normal con media igual al promedio del segmento y pequeña desviación estándar.² Los segmentos no introducen ningún cambio en el ABM, únicamente afectan a las percepciones de los agentes.

Para reducir la influencia de valores atípicos producidos por la inicialización aleatoria de las percepciones, cada ejecución del ABM está compuesta por un número de simulaciones Monte Carlo (MC).³

IV-B. Validación en un caso de estudio real para marketing

En esta subsección se presentan los resultados de la ejecución del ABM. Los resultados representan el número de elecciones en los que cada marca es elegida, considerando que cada agente únicamente lleva a cabo un único proceso de toma de decisión.

En este caso de estudio, se ejecuta nuestro ABM con 1000 agentes, inicializando sus pesos de *drivers* y sus percepciones usando estudios de marketing existentes. En los dos enfoques lingüísticos difusos, se usan como etiquetas lingüísticas el conjunto $S = \{mala, neutra, buena\}$. Este caso de estudio contiene 5 marcas y 6 *drivers*.⁴

En nuestro estudio, se realizan dos experimentos diferentes, uno usando el filtro de conocimiento de marca en la heurística de toma de decisiones y el otro sin usarlo. Cuando este filtro está activo, un agente sólo es capaz de elegir aquellas marcas sobre las que tiene conocimiento. Para desactivarlo, sencillamente se establece al 100 % el conocimiento de marca para todos los agentes y todas las marcas.

La Figura 2 muestra la comparativa de los tres modelos sin usar el filtro de conocimiento de marca (izquierda) y usándolo (derecha). Para ello, se presentan los diagramas de caja donde se representan los valores máximo, mínimo, mediana, y

cuartiles primero y tercero del número de elecciones de cada marca.

Se observa que los resultados en ambos experimentos son diferentes, para cualquier representación de las percepciones del consumidor. Se puede ver, por ejemplo, el número de ventas de la *marca2* usando la representación lingüística 2-tuplas. Cuando no se considera el conocimiento de marca, este número es mucho mayor que el número de ventas cuando este filtro de conocimiento de marca está activo. Además, se pueden observar diferencias en los resultados de los tres modelos. Son especialmente significativas las diferencias de nuestro ABM con respecto al modelo con etiquetas lingüísticas. Por ejemplo, el número de consumidores que eligen la *marca5* es mucho menor en el ABM que modela las percepciones con etiquetas lingüísticas (con y sin conocimiento de marca).

Para medir la precisión de cada modelo, se calcula como estimadores del error el Error Absoluto Medio (MAE, *Mean Absolute Error*) y la Raíz del Error Cuadrático Medio (RMSE, *Root Mean Squared Error*) con respecto a datos reales (es decir, ventas reales).

Cuadro I
MAE Y RMSE DE LOS TRES MODELOS CON RESPECTO A DATOS REALES, CON Y SIN CONOCIMIENTO DE MARCA. MEJOR RESULTADO EN NEGRITA.

Estimador del error	Conocimiento de marca	Representación de las percepciones		
		Númerica	2-tuplas	Et. ling.
MAE	No	2.83	2.95	5.25
	Sí	1.78	1.60	2.46
RMSE	No	3.08	3.20	6.23
	Sí	2.14	1.95	2.98

Los valores de los estimadores MAE y RMSE para los tres modelos se reportan en el Cuadro I. Se puede observar que los modelos con representaciones de las percepciones del consumidor en formatos numéricos y 2-tuplas son mucho más precisos que el modelo con etiquetas lingüísticas. Esto es consecuencia de la representación menos expresiva usada en el modelo lingüístico difuso ordinal. Se hace énfasis en que este problema no aparece en nuestro ABM basado en la representación lingüística 2-tuplas. Además, se puede observar que el rendimiento de cada modelo es mucho peor cuando el filtro de conocimiento de marca no está activo. Esto sugiere que el conocimiento de marca es otra componente fundamental para simular un mercado con precisión. De hecho, el modelo más preciso para ambos estimadores del error (MAE y RMSE) es aquel que representa las percepciones del consumidor usando 2-tuplas lingüísticas y usando el filtro de conocimiento de marca, es decir, el modelo presentado en este trabajo.

V. CONCLUSIONES Y TRABAJO FUTURO

En este trabajo se ha presentado un modelo lingüístico difuso para representar las percepciones de los consumidores, así como una heurística de toma de decisiones que las maneja y que se usa para simular las compras de los consumidores. Todo ello se integra en un ABM para el análisis en marketing, donde los agentes representan a los consumidores del mercado.

²En nuestros experimentos, se usa una desviación estándar de 1,5.

³Cada ejecución del ABM está compuesta de 100 simulaciones MC.

⁴Los nombres de las marcas se han omitido por razones de anonimato.

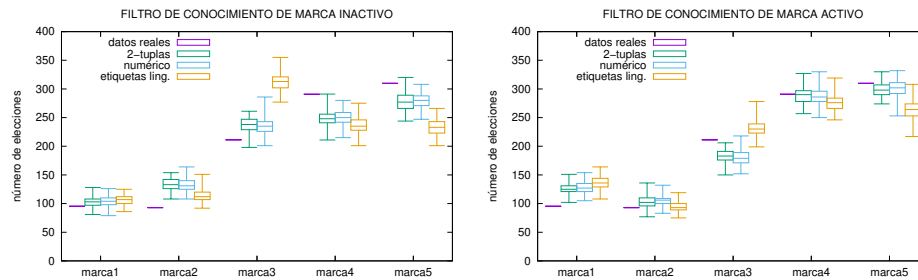


Figura 2. Comparación de las representaciones de percepciones sin usar el filtro de conocimiento de marca (izquierda) y usándolo (derecha).

Las percepciones de los consumidores son normalmente aspectos cualitativos. En nuestro modelo, se usan 2-tuplas lingüísticas, las cuales no sufren ninguna pérdida de información, incluso en el proceso de agregación de la información. En la vida real, los consumidores tienen un conocimiento de marca limitado, por lo que hemos incorporado también esta información en nuestro sistema para hacerlo más realista. Hemos definido una heurística no determinista basada en las valoraciones que cada consumidor realiza sobre cada producto para simular el proceso de toma de decisiones de compra. Se han analizado experimentalmente los resultados de nuestro ABM en un caso de estudio de marketing real a gran escala, mostrando diferencias notables en los diferentes escenarios donde únicamente se modifica la representación de las percepciones del consumidor (sin alterar sus valores), y además se ha estudiado la influencia de incluir información sobre el conocimiento de marca del consumidor en el modelo.

Como trabajo futuro, se planea extender este ABM para marketing basado en 2-tuplas lingüísticas en dos direcciones. Por una parte, se planea investigar otras heurísticas de tomas de decisión en nuestro sistema [20]–[22], adaptándolas para manejar información lingüística difusa. Por otra, se planea extender nuestro ABM para marketing incorporando comportamiento temporal para construir simulaciones con eventos discretos [3]–[5]. De esta forma, las percepciones de cada consumidor pueden cambiar a lo largo del tiempo, y por tanto, se puede analizar como esto afecta a los resultados de la toma de decisiones.

AGRADECIMIENTOS

Trabajo financiado por MICIU, JdA, UGR, AEI, y ERDF (PGC2018-101216-B-I00, A-TIC-284-UGR18, PY18-4475, FJCI-2017-32420). Agradecemos a la compañía ZIO Analytics por los datos aportados, usados en este trabajo.

REFERENCIAS

- [1] G. Overgoor, M. Chica, W. Rand, and A. Weishampel, "Letting the computers take over: Using AI to solve marketing problems," *California Management Review*, p. 0008125619859318, 2019.
- [2] W. Rand and R. T. Rust, "Agent-based modeling in marketing: Guidelines for rigor," *International Journal of Research in Marketing*, vol. 28, no. 3, pp. 181–193, 2011.
- [3] M. Chica and W. Rand, "Building agent-based decision support systems for word-of-mouth programs. a freemium application," *Journal of Marketing Research*, vol. 54, pp. 752–767, 2017.
- [4] C. M. Macal and M. J. North, "Tutorial on agent-based modeling and simulation," in *Proceedings of the 37th Conference on Winter Simulation*. ACM, 2005, pp. 2–15.
- [5] J. M. Epstein, *Generative social science: Studies in agent-based computational modeling*. Princeton University Press, 2006.
- [6] L. Hamill and N. Gilbert, *Agent-Based Modelling in Economics*, 1st ed. Wiley Publishing, 2016.
- [7] I. Moya, M. Chica, J. L. Sáez-Lozano, and O. Córdón, "An agent-based model for understanding the influence of the 11-M terrorist attacks on the 2004 spanish elections," *Knowledge-Based Systems*, vol. 123, pp. 200–216, 2017.
- [8] M. Chica, R. Chiong, M. Adam, and T. Teubner, "An evolutionary game model with punishment and protection to promote trust in the sharing economy," *Scientific Reports*, vol. 9, p. 19789, 2019.
- [9] M. Janssen and W. Jager, "An integrated approach to simulating behavioural processes: A case study of the lock-in of consumption patterns," *Journal of Artificial Societies and Social Simulation*, vol. 2, no. 2, 1999.
- [10] L. A. Zadeh, "The concept of a linguistic variable and its application to approximate reasoning - I," *Information Sciences*, vol. 8, no. 3, pp. 199–249, 1975.
- [11] —, "The concept of a linguistic variable and its application to approximate reasoning - II," *Information Sciences*, vol. 8, no. 4, pp. 301–357, 1975.
- [12] —, "The concept of a linguistic variable and its application to approximate reasoning - III," *Information Sciences*, vol. 9, no. 1, pp. 43–80, 1975.
- [13] J. Giráldez-Cru, M. Chica, O. Córdón, and F. Herrera, "Modeling agent-based consumers decision-making with 2-tuple fuzzy linguistic perceptions," *International Journal of Intelligent Systems*, vol. 35, pp. 283–299, 2020.
- [14] F. Herrera and L. Martínez-López, "A 2-tuple fuzzy linguistic representation model for computing with words," *IEEE Transactions on Fuzzy Systems*, vol. 8, no. 6, pp. 746–752, 2000.
- [15] E. Herrera-Viedma, "Modeling the retrieval process for an information retrieval system using an ordinal fuzzy linguistic approach," *Journal of the American Society for Information Science and Technology*, vol. 52, no. 6, pp. 460–475, 2001.
- [16] M. Brown. (2019) Brand tracking studies. [Online]. Available: <http://www.millwardbrown.com/mb-global/what-we-do/brand/brand-tracking>
- [17] A. Charan, *Marketing analytics: a practitioner's guide to marketing analytics and research methods*. World Scientific Publishing Company, 2015.
- [18] M. Delgado, J. L. Verdegay, and M. A. Vila, "On aggregation operations of linguistic labels," *International Journal of Intelligent Systems*, vol. 8, no. 3, pp. 351–370, 1993.
- [19] R. R. Yager, "On ordered weighted averaging aggregation operators in multicriteria decisionmaking," *IEEE Transactions on Systems, Man, and Cybernetics*, vol. 18, no. 1, pp. 183–190, 1988.
- [20] J. R. Bettman, M. F. Luce, and J. W. Payne, "Constructive Consumer Choice Processes," *Journal of Consumer Research*, vol. 25, no. 3, pp. 187–217, 1998.
- [21] G. Gigerenzer and W. Gaissmaier, "Heuristic decision making," *Annual Review of Psychology*, vol. 62, no. 1, pp. 451–482, 2011.
- [22] C. del Campo, S. Pauser, E. Steiner, and R. Vetschera, "Decision making styles and the use of heuristics in decision making," *Journal of Business Economics*, vol. 86, no. 4, pp. 389–412, 2016.

Generación de descripciones lingüísticas de la contaminación acústica diaria en zonas urbanas

Juan Moreno-García

*Escuela de Ingeniería Industrial y Aeroespacial
Universidad de Castilla-La Mancha
Toledo, Spain
juan.moreno@uclm.es*

Luis Jimenez-Linares

*Escuela Superior de Informática
Universidad de Castilla-La Mancha
Ciudad Real, Spain
luis.jimenez@uclm.es*

Luis Rodriguez-Benitez

*Escuela Superior de Informática
Universidad de Castilla-La Mancha
Ciudad Real, Spain
luis.rodriguez@uclm.es*

Abstract—Uno de los mayores problemas que preocupan a la sociedad actual es la contaminación, distinguiéndose diferentes tipos, por ejemplo, acústica, ambiental, térmica, etc. Entre ellas, la contaminación acústica causa serios problemas a los ciudadanos porque es continua durante gran parte del día, debido a que en su mayoría es causada por el tráfico. Por otro lado, las grandes ciudades aportan una gran cantidad de datos obtenidos diariamente gracias a la sensorización derivada del concepto de “ciudades inteligentes”, lo que permite visualizar la información de las zonas sensorizadas y alertar a las instituciones de los problemas y, a los ciudadanos, conocer la situación de la contaminación acústica en base a datos para poder realizar las quejas y denuncias pertinentes a las instituciones. Una forma universalmente comprensible de mostrar la información contenida en los datos capturados es la generación de descripciones lingüísticas que sintetizan la información que reside en los datos. Este trabajo presenta un método para generar descripciones lingüísticas a partir de los datos de contaminación acústica capturados por las estaciones de medición de ruido. Se presentará un método de generación de descripciones de un día que considera los periodos diarios en los que se estructuran los datos tomados de las estaciones (diurno, vespertino, nocturno y diurno completo). Para probar el método propuesto, se han utilizado los datos disponibles de la ciudad de Madrid para generar descripciones que permitan analizar la influencia de Covid-19 en la contaminación acústica.

Index Terms—descripción lingüística, lógica difusa, contaminación acústica, Covid-19

I. INTRODUCCIÓN

Hoy en día existe una gran concienciación sobre los problemas medioambientales y esto hace que las instituciones lleven a cabo iniciativas para mejorar la situación. La contaminación es uno de los problemas ambientales que más atención recibe debido a su impacto en la salud de las personas y otros seres vivos. Entre ellos se encuentra la contaminación acústica, que llamaremos ruido, y que puede definirse como el exceso de sonido que altera las condiciones ambientales de una zona determinada. El ruido tiene graves consecuencias de todo tipo para quienes lo padecen, siendo los problemas más destacados: los problemas auditivos por estar sometidos de forma habitual a un exceso de ruido en el ambiente, los problemas de sueño (alteración del ciclo del sueño, insomnio, somnolencia

durante el día, cansancio, etc.), los problemas psicológicos (irritabilidad, estrés, ansiedad, etc.) y los fisiológicos (aumento del ritmo cardíaco y respiratorio o de la presión arterial). En las grandes ciudades, las fuentes de ruido son muy diversas, siendo las más importantes el tráfico y la construcción, aunque también se debe a la actividad de las personas en la ciudad. La Organización Mundial de la Salud (OMS) recomienda no superar los 65 decibelios durante el día y los 55 durante la noche.

Además, las ciudades están aplicando el concepto de “ciudades inteligentes” para lograr una mayor sostenibilidad económica, social y medioambiental mediante el uso de las tecnologías de la información y la comunicación (TIC). La aplicación de las TIC crea infraestructuras que pretenden garantizar un desarrollo sostenible, un aumento de la calidad de vida de los ciudadanos, una mayor eficiencia de los recursos disponibles y una participación ciudadana activa. Por ello, las ciudades disponen de datos diarios procedentes de muy diversas fuentes, entre las que se encuentran las estaciones de medición de ruido, que pueden utilizarse para lograr la sostenibilidad deseada. A partir de ellos, se puede obtener información sobre el ruido de las zonas sensorizadas, lo que permite alertar a las instituciones para que tomen las medidas oportunas, y se puede informar a los ciudadanos de la situación de la contaminación acústica en base a los datos para que puedan realizar protestas ante las instituciones en base a información contrastada.

Para que la información contenida en los datos sea comprensible para cualquiera, se utilizan descripciones lingüísticas generadas automáticamente. La lógica difusa se puede utilizar para el tratamiento de los datos de ruido, encontrando en la literatura trabajos que evalúan el riesgo [1], [2], realizan un análisis del ruido [3], [4], predicen el ruido [5], crean diseños tratando de evitar el ruido [6], etc. Este trabajo presenta un método para generar descripciones lingüísticas diarias a partir de los datos de ruido captados por las estaciones de medición que considera los cuatro periodos diarios para los que se dispone de información (día, tarde, noche y día completo). Para probar el método propuesto, se han utilizado los datos disponibles de la ciudad de Madrid para generar descripciones que permitan analizar la influencia de Covid-19 en la contaminación acústica.

Este trabajo ha sido financiado por el Departamento de Tecnologías y Sistemas de Información de la Universidad de Castilla-La Mancha, por la Escuela de Ingeniería Industrial y Aeroespacial de Toledo y por la Escuela Superior de Informática de Ciudad Real.

El resto del documento está organizado como sigue: la Sección II describe algunos conceptos relevantes en relación con el sonido y el ruido. La Sección III presenta el proceso llevado a cabo para representar y clasificar el ruido utilizando conjuntos difusos. La Sección IV muestra nuestra propuesta para generar descripciones lingüísticas de un día. La Sección V presenta algunas pruebas realizadas con datos de ruido en cinco zonas de Madrid para realizar una primera aproximación a la detección del efecto del encierro domiciliario debido a Covid-19. Finalmente, las conclusiones y los trabajos futuros se detallarán en la Sección VI.

II. CONTAMINACIÓN ACÚSTICA

Las definiciones de sonido, acústica y ruido [7] que se dan a continuación son necesarias para la correcta comprensión de este trabajo. El sonido se define como la variación de presión producida en un medio (sólido, líquido o gaseoso) por un elemento vibrante que puede ser detectado por el oído humano. La acústica puede definirse como la ciencia que se ocupa de la producción, control, transmisión, recepción y efectos del sonido. El ruido se define como un conjunto de sonidos inarmónicos o desafinados que resultan desagradables para el oído humano, es decir, un sonido molesto.

La medición del ruido está influida por la distancia a la fuente de ruido, que puede ser una fuente puntual, una fuente espacial o una fuente lineal. En las ciudades, la principal fuente de ruido es el tráfico, por ejemplo, en España es la causa del 99% del ruido urbano. El estudio del ruido utiliza las mismas magnitudes que para el sonido, que en su expresión más simple produce la formación de una onda sinusoidal con las siguientes magnitudes: velocidad, longitud de onda, periodo y amplitud.

Dos conceptos importantes en la medición del ruido que a veces se confunden, probablemente porque ambos se miden en decibelios (dB), son el “nivel de potencia sonora” y el “nivel de presión sonora”. El nivel de potencia sonora emitido por una fuente sonora determina la cantidad de ruido que produce, mientras que el nivel de presión sonora determina la cantidad de sonido que llega a un punto determinado. El nivel de presión sonora depende de factores como la distancia de la fuente al foco, la dirección o la existencia de otros ruidos en el entorno (ruido de fondo).

Las estaciones de medición de ruido, conocidas como estaciones NMT, suelen utilizar el nivel de presión sonora ponderado A (dBA) para medir el ruido en una zona. En este trabajo hemos utilizado datos de ruido extraídos de la página web del “Portal de datos abiertos del Ayuntamiento de Madrid”¹. Los conjuntos de datos proporcionan seis mediciones de presión sonora para cada uno de los cuatro periodos de tiempo en los que se puede dividir un día. El primer periodo corresponde al día completo (T) mientras que los otros tres periodos dividen el día en el periodo diurno (D) de 7:00 a 19:00, el periodo vespertino (E) de 19:00 a 23:00 y el periodo nocturno (N) de 23:00 a 7:00. Para cada periodo, la medida L_{Aeq} es el nivel de presión sonora continuo equivalente ponderado A

determinado a lo largo de todo el día (periodo T). Además, se proporcionan otras cinco medidas de presión sonora con ponderación de frecuencia A y ponderación de tiempo lento para indicar el nivel que se supera durante un tiempo de observación. Estas medidas son L_{AS01} , L_{AS10} , L_{AS50} , L_{AS90} y L_{AS99} para indicar el ruido superado durante 1%, 10%, 50%, 90% y 99% respectivamente. Los valores de L_{AS01} y L_{AS99} se aproximan al ruido mínimo y máximo alcanzado durante el periodo estudiado respectivamente. En la Tabla I se detalla una muestra correspondiente al 9 de abril de 2020 de la estación situada en la Plaza del Emperador Carlos V de Madrid, donde se muestran los cuatro periodos en las filas y las seis mediciones para cada uno de estos periodos en las columnas.

TABLA I
MUESTRAS DE UN DÍA EN LA ESTACIÓN NMT CARLOS V DE MADRID

Periodo	L_{Aeq}	L_{AS01}	L_{AS10}	L_{AS50}	L_{AS90}	L_{AS99}
D	62.1	69.5	64.8	58.4	52.6	48.7
E	63.5	71.2	66.8	60.2	54.3	50.6
N	59.5	68.0	61.7	54.1	46.4	43.0
T	61.7	69.8	64.8	57.5	50.2	44.0

III. DESCRIPCIÓN LINGÜÍSTICA DEL RUIDO

En esta sección se muestra cómo se pueden representar los datos de ruido para ser utilizados en la generación de descripciones, también se detallan algunos conceptos utilizados por el método presentado para generar las descripciones. En concreto, se utiliza la lógica difusa, ya que permite representar lingüísticamente valores numéricos interrelacionados, como es el caso de las mediciones de presión sonora. Para representar los valores de un periodo se utiliza un conjunto difuso con una función de pertenencia gaussiana (Sección III-A). Para categorizar estos conjuntos difusos, hemos creado un conjunto de etiquetas lingüísticas con una función de pertenencia gaussiana que representa las clases en las que se clasifica el ruido (Sección III-B). Finalmente, para clasificar el ruido de un periodo, se compara el conjunto difuso gaussiano con las etiquetas del conjunto de etiquetas (las clases) seleccionando la clase que ofrece la mayor intersección con el conjunto (Sección III-C).

A. Obtención de un conjunto difuso con función de pertenencia gaussiana a partir de datos con ruido

Las medidas L_{AS01} , L_{AS10} , L_{AS50} , L_{AS90} y L_{AS99} se utilizarán para generar un conjunto difuso que represente una muestra de un periodo. La información obtenida de las estaciones permite conocer la distribución de los valores que superan un umbral durante un 1%, 10%, 50%, 90% y 99% que pueden ser representados por un conjunto difuso con una función de pertenencia gaussiana. La función gaussiana está definida por la Ecuación 1.

$$a * e^{-\frac{(x-b)^2}{2*c^2}} \quad (1)$$

¹<https://cutt.ly/Lj6O0xP>

donde a , b y c son constantes reales y $c > 1$. a es el valor máximo que toma la función, b es la posición central de la función y c la desviación típica que riga la amplitud de la función.

En este trabajo se han utilizado los siguientes valores para estos tres parámetros: $a = 1$ (una función de pertenencia da valores de salida en el intervalo $[0, 1]$), $b = \rho$ (media aritmética) y $c = \sigma$ (desviación estándar). Como puede verse, la función gaussiana está determinada por la media y la desviación estándar. Por tanto, para obtener el conjunto difuso que la representa, hay que calcular estos dos valores para L_{AS01} , L_{AS10} , L_{AS50} , L_{AS90} y L_{AS99} . Por ejemplo, para el día completo que se muestra en la última fila de la Tabla I (periodo T) se utiliza como entrada $[69.8, 64.8, 57.5, 50.2, 44.0]$, entonces $\rho = 57.26$ y $\sigma = 9.38$, aplicando la ecuación se obtiene el conjunto difuso mostrado en la Figura 1. Como se puede observar, establece el centro en 57.26 y para $x = 44$ y $x = 69.8$ tiene valores muy bajos.

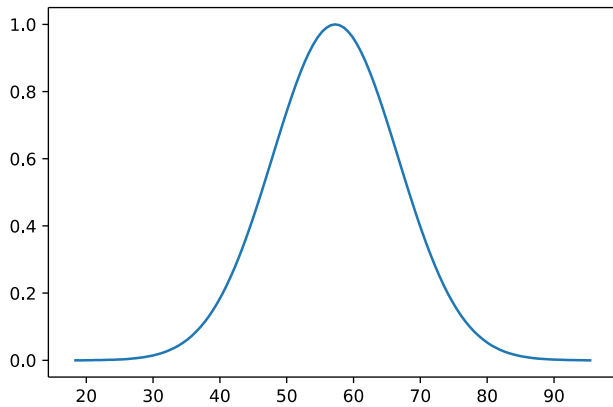


Fig. 1. Conjunto obtenido para el periodo T mostrado en la última fila de la Tabla I.

B. Clasificación de la contaminación ambiental mediante conjuntos difusos

Para clasificar los conjuntos difusos obtenidos en el apartado anterior, se necesita un conjunto de etiquetas lingüísticas que definan las clases a considerar, que se denominarán *CLASSES*. El proceso para obtener las etiquetas que representan las clases es el mismo que el utilizado para calcular los conjuntos difusos. Con la ayuda de un experto, se han definido los valores L_{AS01} , L_{AS10} , L_{AS50} , L_{AS90} y L_{AS99} para cada etiqueta del conjunto que definen la clase (Tabla II), generando posteriormente la etiqueta con función de pertenencia gaussiana mediante el método mostrado anteriormente. Se han seguido las directrices de la OMS, definiendo el “ruido molesto” en un valor cercano a 60 dBA y clasificando el “ruido perjudicial” en 70 dBA o superior, mientras que las otras tres etiquetas dividen el ruido no molesto para el oído humano en tres clases (nulo, imperceptible y aceptable). La Figura 2

muestra gráficamente las etiquetas obtenidas para cada clase. Obsérvese que las etiquetas podrían hacer uso de cualquier otro tipo de función de pertenencia, por ejemplo, trapezoidal, triangular, etc.

TABLA II
VALORES DE L_{AS01} , L_{AS10} , L_{AS50} , L_{AS90} Y L_{AS99} UTILIZADOS PARA DETERMINAR LAS CLASES

Periodo	L_{AS01}	L_{AS10}	L_{AS50}	L_{AS90}	L_{AS99}
nulo	20.0	25.0	30.0	35.0	40.0
imperceptible	30.0	35.0	40.0	45.0	50.0
aceptable	40.0	45.0	50.0	55.0	60.0
molesto	50.0	55.0	60.0	65.0	70.0
perjudicial	60.0	65.0	70.0	75.0	80.0

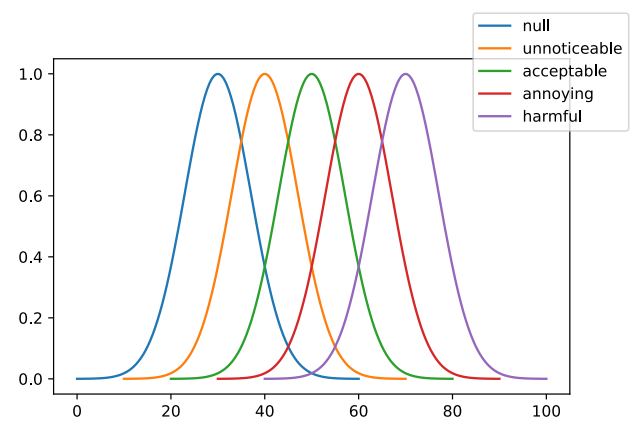


Fig. 2. Conjunto de etiquetas lingüísticas para clasificar el ruido (*CLASSES*).

C. Clasificación de un conjunto difuso mediante el conjunto de etiquetas que clasifica el ruido

Una vez que la información de entrada se puede representar como un conjunto difuso y se dispone del conjunto *CLASSES* para su clasificación, queda detallar la forma en que se realiza el proceso de clasificación, que se basa en el concepto de intersección entre conjuntos difusos. La intersección entre dos conjuntos difusos A y B se define como otro conjunto difuso C cuya función de pertenencia es $\mu_C(x) = \mu_A(x) \cap \mu_B(x)$. La Figura 3 muestra varios ejemplos de la intersección de un conjunto difuso y una etiqueta con el conjunto de intersección resultante entre los dos conjuntos resaltados en verde.

Para clasificar un conjunto difuso se realiza la intersección de las etiquetas *CLASSES* con el conjunto difuso, seleccionando la clase que ofrece la mayor intersección de todas, es decir, se asigna la clase de *CLASSES* cuya intersección con el conjunto difuso cubre el mayor porcentaje de la función de pertenencia del conjunto difuso a comparar. Por ejemplo, en la Figura 3 se muestra la comparación del conjunto difuso T con las etiquetas de clase *aceptable*, *molesto* y *perjudicial*. Dado que la intersección cubre 58%, 73% y 38% del área

de la función de pertenencia de T para estas tres clases se determina que la clase del conjunto difuso T es *molesto*. De este modo, la clase de un ruido puede obtenerse mediante un conjunto de etiquetas.

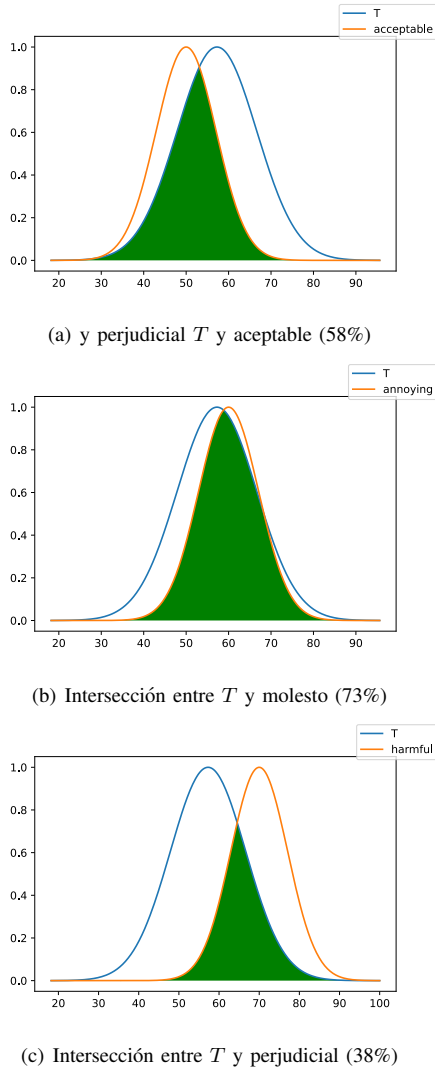


Fig. 3. Ejemplos de selección de clases de un conjunto difuso T

IV. GENERACIÓN DE DESCRIPCIONES LINGÜÍSTICAS

A continuación se presenta el método propuesto para generar una descripción de los datos diarios que hace uso de los conceptos mostrados en el apartado anterior. Para generar las descripciones diarias se utilizarán los cuatro periodos diarios (D , E , N , T) y las cinco medidas L_{AS01} , L_{AS10} , L_{AS50} , L_{AS90} y L_{AS99} de cada periodo. El Algoritmo 1 muestra los pasos seguidos para la generación de la descripción. Necesita como entrada un vector que contenga las medidas del día para cada periodo, llamado $DATA$. Las sentencias que componen el algoritmo son las siguientes:

- Línea 1: calcula el conjunto de etiquetas que representan las clases como se muestra en la sección III-B y genera

Algorithm 1 Descripción del comportamiento de un día

```

1:  $CLASSES \leftarrow$  ObteniendoClases() {Sección III-B.}
   {El siguiente bucle calcula las gaussianas (Sección III-A).}
2: for  $period(p)$  in  $[T, D, E, N]$  do
3:    $GAUSS[p] \leftarrow$  CalcularGaussiana( $DATA[p]$ ) { $DATA$  es un vector que contiene los datos de entrada.  $GAUSS$  es un vector que contiene las gaussianas de cada periodo.}
4: end for
   {El siguiente bucle calcula los vectores  $LABEL$  y  $PERC$ .}
5: for  $period(p)$  in  $[T, D, E, N]$  do
6:    $LABEL[p] \leftarrow$  SelecClass( $GAUSS[p]$ ,  $CLASSES$ ) {Sección III-C}
7:    $PERC[LABEL[p]] \leftarrow$  Intersección( $GAUSS[p]$ ,  $CLASSES[LABEL[p]]$ )
8: end for
   { $description$  es la descripción generada utilizando  $MAXIMA$ .}
9:  $description \leftarrow$  GenerarPlantilla( $MAXIMA$ )

```

el conjunto de etiquetas $CLASSES$ que contiene las etiquetas de las clases (Figura 2).

- Líneas 2-4: Este bucle obtiene los conjuntos difusos con función de pertenencia gaussiana que representan los datos para cada uno de los periodos de las mediciones del nivel de presión sonora ($DATA$) y almacena el resultado en el vector $GAUSS$. El método para completar este paso se describió en la Sección III-A.
- Líneas 5-8: Este bucle calcula la clase que representa cada uno de los periodos junto con el porcentaje de área que cubre la función de pertenencia del conjunto difuso de entrada por esa clase. Esto se hace realizando la intersección entre cada conjunto difuso en $GAUSS$ con las etiquetas en $CLASSES$. Se selecciona la clase que cubre el mayor porcentaje del área de la función de membresía del conjunto difuso (Sección III-C).
- Línea 6: obtiene la clase asignada al conjunto difuso para ese periodo ($GAUSS[class]$) almacenándola en $LABEL$.
- Línea 7: calcula el porcentaje del conjunto $GAUSS[class]$ que cubre la intersección con la etiqueta seleccionada de $CLASSES$ ($CLASSES[LABEL[period]]$) y lo almacena en el vector $PERC$.
- Línea 9: genera la descripción lingüística utilizando $LABEL$ y $PERC$ para completar la siguiente plantilla:

En general, fue un día con ruido $LABEL[T]$ durante el $PERC[T]\%$ del tiempo. Durante el día hubo un ruido $LABEL[D]$ que ocupó el $PERC[D]\%$ del tiempo. Por la noche, durante el $PERC[E]\%$ del tiempo, se produjo un ruido $LABEL[E]$. Por último, fue una noche con un ruido

$LABEL[N]$ durante $PERC[N]\%$ de su tiempo

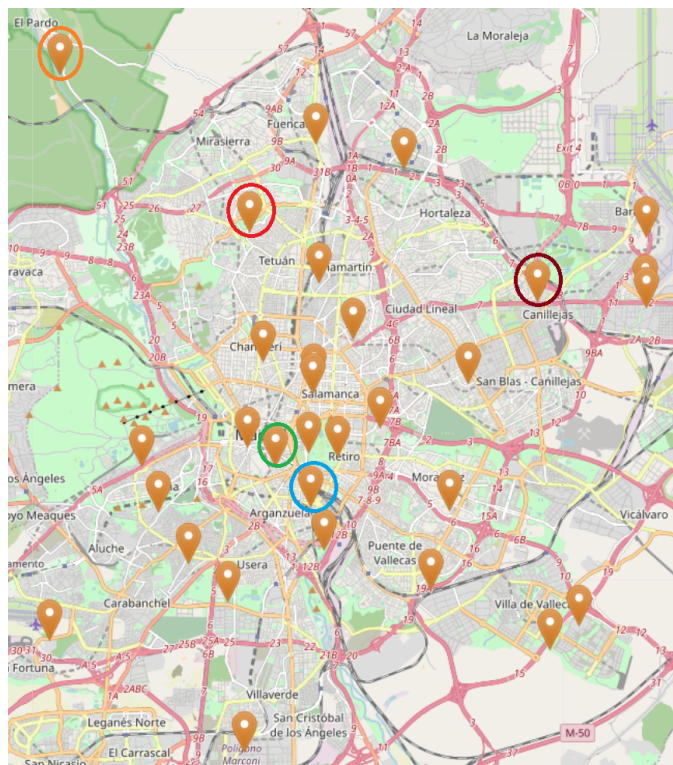


Fig. 4. Ubicación de las zonas de medición utilizadas en las pruebas.

V. TESTS

Para probar el método presentado se han generado descripciones para cinco áreas de Madrid situadas en zonas con diferentes características identificadas en la Figura 4. Situadas en el centro de Madrid se encuentran las estaciones de Plaza del Carmen (círculo verde) y Carlos V (círculo azul), que son zonas con bastante vida urbana y tráfico, especialmente la Plaza del Emperador Carlos V, cercana al nudo de comunicaciones de Atocha. Algo alejadas del centro están las estaciones de El Barrio del Pilar (círculo rojo), situada en la Avenida de Monforte Lemos junto al Parque de la Vaguada, y Campo de las Naciones (círculo granate), con la estación situada en la rotonda del monumento a Don Juan de Borbón, que tiene un tráfico constante. Lejos del centro está la estación de El Pardo (círculo naranja), un barrio situado en un bosque, una zona muy tranquila. Se han elegido dos miércoles para generar las descripciones, concretamente el 5 de febrero de 2020 y el 1 de abril de 2020. El día 5-2-2020 corresponde a una fecha anterior a la del encierro domiciliario de España por Covid-19 (15 de marzo de 2020), es decir, representa la actividad normal, mientras que el día 1-4-2020 es un día de encierro domiciliario. De este modo, las descripciones nos permitirán evaluar el efecto del confinamiento domiciliario en la contaminación acústica, aunque en el futuro se deberá realizar un estudio más amplio para poder extraer conclusiones válidas. En el apartado de bibliografía, existen algunos estudios sobre

el efecto del Covid-19 en la contaminación acústica en grandes ciudades, como Dublín [8], Roma [9] o Madrid [10].

TABLA III
DESCRIPCIÓN LINGÜÍSTICA DE LA ESTACIÓN DE MEDIDA CARLOS V

Carlos 5 (1-02-2020)
En general, fue un día con ruido perjudicial durante el 71,0% del tiempo. Durante el día hubo ruido perjudicial el 98,0% del tiempo. Por la tarde, durante el 99,0% del tiempo, se alcanzó un ruido perjudicial. Por último, por la noche hubo ruido molesto durante el 84,0% del tiempo.
Carlos 5 (5-04-2020) - confinamiento domiciliario
En general, fue un día con ruido molesto durante el 80,0% del tiempo. Durante el día hubo ruidos molestos que ocuparon el 84,0% del tiempo. Por la tarde, durante el 94,0% del tiempo, se alcanzó un ruido molesto. Por último, por la noche hubo ruido molesto durante el 75,0% del tiempo.

Tabla III muestra las descripciones obtenidas para la estación situada en Carlos V. Como se puede observar, se trata de una zona muy ruidosa en día laborable, obteniendo una clasificación de “ruido perjudicial” durante todos los periodos y persistiendo durante la mayor parte del tiempo. Se observa que durante el encierro recibe una clasificación de ruido menor, aunque se siguen manteniendo valores superiores a las recomendaciones de la OMS durante casi todos los periodos completos. Esto puede deberse a la concentración de tráfico continuo e intenso en la zona.

TABLA IV
DESCRIPCIÓN LINGÜÍSTICA DE LA ESTACIÓN DE MEDIDA DE LA PLAZA DEL CARMEN

Plaza del Carmen (1-02-2020)
En general, fue un día con ruido molesto durante el 76,0% del tiempo. Durante el día hubo ruido molesto el 91,0% del tiempo. Por la tarde, durante el 95,0% del tiempo hubo ruido molesto. Por último, por la noche hubo ruido molesto durante el 68,0% del tiempo.
Plaza del Carmen (5-04-2020) - confinamiento domiciliario
En general, fue un día con un ruido aceptable durante el 73,0% del tiempo. Durante el día hubo un ruido aceptable durante el 60,0% del tiempo. Por la tarde, durante el 75,0% del tiempo, se obtuvo un ruido aceptable. Por último, por la noche hubo un ruido aceptable durante el 76,0% del tiempo.

En la Tabla IV se detallan las descripciones obtenidas para la estación situada en la Plaza del Carmen. Durante un día laborable es una zona con ruido molesto durante todo el día, aunque se observa que por la noche el ruido molesto se mantiene durante menos tiempo que el resto del día. El confinamiento causó un efecto positivo en esta zona obteniendo un ruido aceptable durante más del 60% del tiempo de todos los periodos.

La Tabla V muestra las descripciones de la estación situada en el Barrio del Pilar. En un día laborable es un barrio ruidoso todo el tiempo en los periodos diurnos y vespertinos, pero por la noche se consigue un ruido aceptable. En confinamiento la calidad de vida mejoró con respecto a la contaminación acústica, convirtiéndose en un barrio con un ruido aceptable durante todo el día. En este caso, esto se debe al fuerte des-

TABLA V
DESCRIPCIÓN LINGÜÍSTICA DE LA ESTACIÓN DE MEDIDA DEL BARRIO PILAR

Barrio Pilar (1-02-2020)
En general, fue un día con ruido molesto durante el 66,0% del tiempo. Durante el día hubo ruido molesto el 100,0% del tiempo. Por la tarde, durante el 100,0% del tiempo, hubo ruido molesto. Por último, durante la noche hubo un ruido aceptable durante el 93,0% del tiempo.
Barrio Pilar (5-04-2020) - confinamiento domiciliario
En general, fue un día con un ruido aceptable durante el 63,0% del tiempo. Por la mañana hubo un ruido aceptable durante el 74,0% del tiempo. Por la tarde, el ruido fue aceptable el 67,0% del tiempo. Por último, fue una noche con un ruido aceptable durante el 66,0% de su tiempo.

censo del tráfico en la zona provocado por el confinamiento, ya que los residentes no pueden circular libremente.

TABLA VI
DESCRIPCIÓN LINGÜÍSTICA DE LA ESTACIÓN DE MEDIDA DE CAMPO DE LAS NACIONES

Campo de las Naciones (1-02-2020)
En general, fue un día con un ruido aceptable durante el 63,0% del tiempo. Durante el día hubo un ruido molesto que ocupó el 100,0% del tiempo. Por la tarde, durante el 91,0% del tiempo, hubo ruido molesto. Por último, por la noche hubo un ruido aceptable durante el 70,0% del tiempo.
Campo de las Naciones (5-04-2020) - confinamiento domiciliario
En general, fue un día con un ruido aceptable durante el 76,0% del tiempo. Durante el día hubo un ruido aceptable durante el 76,0% del tiempo. Por la tarde, durante el 77,0% del tiempo, se obtuvo un ruido aceptable. Por último, por la noche hubo un ruido aceptable durante el 80,0% del tiempo.

La estación de Campo de las Naciones está situada en una rotonda con tráfico, lo que provoca un nivel de ruido molesto durante el día y la tarde, con un nivel de ruido aceptable por la noche (Tabla VI). Cuando se produjo el confinamiento, se alcanzó un nivel de ruido aceptable durante más del 75% de las veces.

TABLA VII
DESCRIPCIÓN LINGÜÍSTICA DE LA ESTACIÓN DE MEDICIÓN DE EL PARDO

El Pardo (1-02-2020)
En general, fue un día con un ruido aceptable durante el 54,0% del tiempo. Durante el día hubo un ruido molesto que ocupó el 76,0% del tiempo. Por la tarde, durante el 72,0% del tiempo, hubo un ruido aceptable. Por último, por la noche hubo un ruido imperceptible durante el 58,0% del tiempo.
El Pardo (5-04-2020) - confinamiento domiciliario
En general, fue un día con un ruido aceptable durante el 56,0% de su tiempo. Durante el día hubo un ruido aceptable durante el 64,0% del tiempo. Por la tarde, durante el 63,0% del tiempo, se obtuvo un ruido aceptable. Por último, por la noche hubo un ruido imperceptible durante el 64,0% del tiempo.

Por último, se ha comprobado que en general El Pardo es una zona con un nivel de ruido aceptable, destacando que el confinamiento ha modificado a mejor los niveles de ruido de la mañana, manteniendo el descanso. Se puede concluir que es una zona sin mucho ruido de forma habitual, por lo que el

confinamiento no ha tenido tanto impacto como en las otras zonas estudiadas.

VI. CONCLUSIONES

En este trabajo se ha presentado un método para generar descripciones lingüísticas de la contaminación acústica, más concretamente, se ha detallado un método para generar descripciones de un día que considera cuatro periodos diarios (día, tarde, noche y día completo). Para probar el método propuesto, se han utilizado los datos disponibles de la ciudad de Madrid, seleccionando cinco zonas con diferentes características y considerando dos días, uno con actividad normal y otro con confinamiento domiciliario debido a la Covid-19. De las descripciones obtenidas se puede concluir que el confinamiento ha reducido la contaminación acústica, excepto en un lugar donde ya era aceptable (El Pardo). También se ha comprobado que estas descripciones son válidas para el estudio de la contaminación acústica, ya que son más fáciles de analizar que los datos brutos.

Como trabajo futuro estamos considerando el uso de otros tipos de funciones de pertenencia, tanto para el conjunto de etiquetas lingüísticas como para los conjuntos difusos utilizados para representar los datos tomados de las estaciones. También pretendemos investigar la descripción de intervalos de tiempo más largos, como semanas, meses o años.

REFERENCES

- [1] R. Golmohammadi, M. Eshaghi, and M. Reyahi Khoram, "Fuzzy Logic Method for Assessment of Noise Exposure Risk in an Industrial Workplace", *Int J Occup Hyg*, vol. 3, no. 2, pp. 49-55, 2011.
- [2] S.K. De, B.K. Swain, S. Goswami, and M. Das, "Adaptive noise risk modelling: fuzzy logic approach", *Systems Science & Control Engineering*, vol. 5, no. 1, pp. 129-141, 2017.
- [3] M. Ghaderi, H. Javadikia, L. Naderloo, M. Mostafaei, and H. Rabbani, "Analysis of noise pollution emitted by stationary MF285 tractor using different mixtures of biodiesel, bioethanol, and diesel through artificial intelligence", *Environmental Science and Pollution Research*, vol. 26, pp. 21682-21692, 2019.
- [4] M. F. I. Mohd Idris, N. A. Zainol Abidin, and K. A. Abd Aziz, "The Determination of Factors that Influence the Noise Pollution in Malaysia using Fuzzy Logic", *Journal of Computing Research and Innovation*, vol. 5, no. 2, pp. 1-10, Oct. 2020.
- [5] V. Nouraniac, H. Gökçekuş, and I.K. Umar, "Artificial intelligence based ensemble model for prediction of vehicular traffic noise", *Environmental Research*, vol. 10, 108852, 2020.
- [6] P.T. Nguyen, "Construction site layout planning and safety management using fuzzy-based bee colony optimization model", *Neural Comput & Applic*, 2020 (<https://doi.org/10.1007/s00521-020-05361-0>).
- [7] R J Peters, B.J. Smith, and M. Hollins, "Acoustics and Noise Control", Third Edition, 2011.
- [8] B. Basu, E. Murphy, A. Molter, A. Sarkar Basu, S. Sannigrahi, M. Belmonte, and F. Pilla, "Investigating changes in noise pollution due to the COVID-19 lockdown: The case of Dublin, Ireland", *Sustainable Cities and Society*, vol. 65, 102597, February 2021.
- [9] F. Aletta, S. Brinchi, S. Carrese, A. Gemma, C. Guattari, L. Mannini, and S.M. Patella, "Analysing urban traffic volumes and mapping noise emissions in Rome (Italy) in the context of containment measures for the COVID-19 disease", *Noise Mapping*, vol. 7, pp. 114-122, 2020.
- [10] C. Asensio, I. Pavón, and G. de Arcas, "Changes in noise levels in the city of Madrid during COVID-19 lockdown in 2020", *The Journal of the Acoustical Society of America*, vol. 148, pp. 1748, 2020.

Un sistema para la descripción automática en lenguaje natural de gráficas de sectores: aplicación en datos de calidad del aire

Andrea Cascallar-Fuentes*, Javier Gallego-Fernández*, Alejandro Ramos-Soto*, Anthony Saunders†, Alberto Bugarín-Diz*

*Centro Singular de Investigación en Tecnoloxías Intelixentes (CiTIUS), Universidade de Santiago de Compostela, Spain
javier.gallego.fernandez@rai.usc.es, {andrea.cascallar.fuentes, alejandro.ramos, alberto.bugarin.diz}@usc.es

†MeteoGalicia, Xunta de Galicia, Santiago de Compostela, Spain
{calidadedoaire.cma}@xunta.es

Resumen—En este trabajo presentamos un modelo basado en la generación de lenguaje natural y en la lógica borrosa para la generación automática de descripciones lingüísticas a partir de datos numéricos y su aplicación real en el ámbito de la información ambiental. Basado en dicho modelo, describimos el sistema ICA2Text, que genera automáticamente descripciones en lenguaje natural sobre el índice de calidad del aire (ICA), que es un indicador estándar utilizado por todas las agencias meteorológicas a nivel mundial. ICA2Text es una aplicación real que opera sobre los datos del ICA proporcionados por la Red de Calidad del Aire de la Agencia de Meteorología de Galicia (MeteoGalicia). Siguiendo la metodología estándar de evaluación en el campo de la generación de lenguaje natural, presentamos los resultados de evaluación manual del sistema por parte de tres expertos meteorólogos. Los resultados de dicha evaluación fueron muy satisfactorios, confirmando empíricamente que las descripciones en lenguaje natural que se generan a partir de los datos resultaron muy adecuadas, tanto en su contenido como en su calidad lingüística. Por ello, el sistema estará operativo en breve como servicio público para los usuarios de la web de MeteoGalicia.

Index Terms—descripciones lingüísticas de datos, sistemas data-to-text, generación de lenguaje natural.

I. INTRODUCCIÓN

Obtener información relevante a partir de grandes cantidades de datos plantea varios retos que no pueden abordarse únicamente con las técnicas tradicionales basadas en la estadística y las visualizaciones gráficas. En general, estos enfoques son muy útiles para obtener información básica de los datos, pero para que los usuarios comprendan la información realmente relevante que hay detrás de los datos, es necesario emplear técnicas que se adapten mejor a las necesidades específicas de cada dominio y que puedan escalar a medida que aumenta la cantidad de datos. En este sentido, la Inteligencia Artificial proporciona a los usuarios herramientas de análisis de conjuntos de datos para extraer información útil, así como herramientas de procesamiento del lenguaje que permiten una comunicación más fluida entre humanos y máquinas.

Una técnica prometedora para este propósito es la Generación de Lenguaje Natural (NLG, por sus siglas en inglés), que permite generar texto a partir de varias fuentes de datos (principalmente numéricos y textuales).

Dentro de este campo, la arquitectura [1] más empleada en la literatura es la siguiente (Figura 1):

- Determinación de contenido: esta fase se centra en decidir qué información será comunicada en el texto.
- Planificación del discurso: en esta etapa se decide el conjunto de mensajes que serán verbalizados y se les asigna un orden y una estructura.
- Planificación de la sentencia: en esta fase se agrupan los mensajes según sea necesario y se eligen las palabras y expresiones que deben ser utilizadas.
- Realización lingüística: en esta etapa se lleva a cabo el proceso de generar el texto resultante, que debe ser morfológica y ortográficamente correcto.

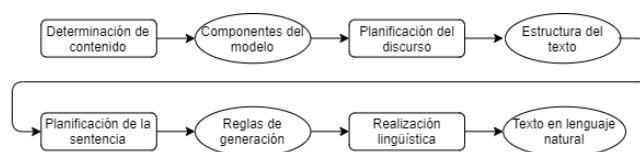


Figura 1. Arquitectura NLG propuesta en [1]. Los rectángulos muestran las fases más importantes mientras que las elipses muestran la salida de cada una de ellas.

Dentro del NLG, los sistemas data-to-text (D2T) [2] generan automáticamente textos a partir de grandes conjuntos de datos numéricos o simbólicos, proporcionando una información comprensible que no podría producirse de otro modo. Los sistemas D2T incluyen: *i*) una etapa de análisis de datos en la que se extrae la información relevante de los mismos y *ii*) una etapa de generación en la que se transmite la información en lenguaje natural.

Además, también relacionado con el NLG, desde el campo de la lógica difusa se propusieron varios enfoques para generar

descripciones lingüísticas de datos (LDD, por sus siglas en inglés), basadas en el uso de términos lingüísticos modelados como conjuntos difusos. Siguiendo este enfoque, algunas aproximaciones [3] resumen en una forma lingüística una o más variables numéricas y sus valores, utilizando la noción general de protoforma [4] y de forma más específica, sentencias cuantificadas borrosas que pueden seguir varios tipos de estructura (por ejemplo, “En algunos puntos la temperatura es alta”).

Las LDD carecen, en general, de la expresividad de los textos reales, pero no por ello dejan de ser elementos de información útiles que pueden utilizarse como entrada de alto nivel para los sistemas NLG en general y D2T en particular [5], [6]. Se han propuesto aplicaciones de descripciones lingüísticas en diversos campos como en el de la salud [7], negocios [8], ahorro de energía [9] o meteorología [10].

En este trabajo proponemos el sistema ICA2Text, que tiene como objetivo la generación de descripciones lingüísticas a partir de datos sobre el Índice de Calidad del Aire (ICA), que es un indicador estándar de la calidad del aire utilizado por todas las agencias meteorológicas y medioambientales del mundo [11]. Nuestro sistema genera descripciones en lenguaje natural a partir de datos de distribución del ICA en las estaciones de la red de observación meteorológica proporcionadas por la Agencia Meteorológica de Galicia, MeteoGalicia [12].

En el campo de la meteorología existen varios enfoques a lo largo de los años para generar descripciones a partir de los datos de calidad del aire. [13] es un prototipo del sistema TEMSIS centrado en describir si se han superado o no los umbrales (diarios, mensuales...) de alerta temprana para cada contaminante. También el sistema MARQUIS [14] de series temporales de índice de calidad del aire numéricas incluye referencias temporales simples (horas o intervalos específicos) para describir el último valor del ICA y los contaminantes (concentración, información de archivo y previsión). En [15] se propone una solución centrada en la generación de descripciones del índice de calidad del aire para una ventana temporal que incluye tres valores diarios.

Este trabajo está estructurado de la siguiente manera: en primer lugar, en la sección II presentamos el contexto del problema gestionado en este trabajo y la arquitectura de nuestro enfoque. En la Sección III describimos en profundidad los detalles del modelo para la descripción en lenguaje natural de la distribución del Índice de Calidad del Aire. En la sección IV se presenta la validación por expertos humanos y sus resultados. Por último, en la sección V ofrecemos algunas observaciones finales y una discusión sobre trabajo futuro.

II. CONTEXTO DEL PROBLEMA

La presencia de contaminantes en el aire y, por lo tanto, el deterioro de la calidad del aire, pueden tener efectos nocivos para la salud de las personas.

El Índice de Calidad del Aire (ICA) es una variable simbólica del tiempo que representa la calidad del aire en cada momento midiendo la presencia y densidad de diversos tipos de partículas contaminantes en la atmósfera, utilizada por las

agencias meteorológicas y los gobiernos para informar a la población sobre la calidad del aire.

En concreto, los datos con los que trabajamos describen el Índice de Calidad del Aire en la red de 50 estaciones meteorológicas (Figura 2) que envían datos actualizados cada hora en tiempo real en Galicia. Se trata de datos oficiales proporcionados por Meteogalicia, que es la agencia de meteorología oficial de la Xunta de Galicia [12].

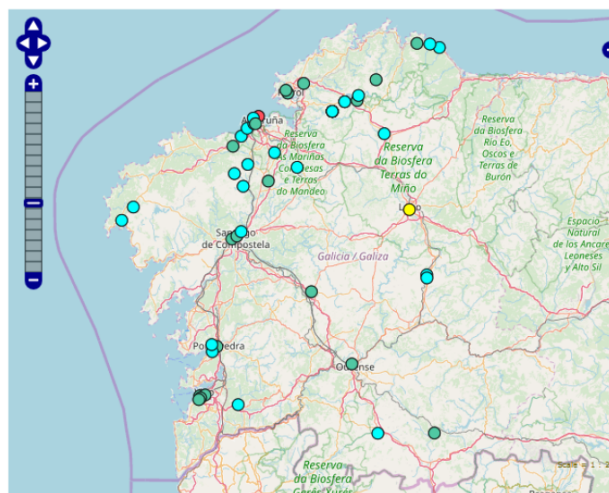


Figura 2. Mapa de las estaciones meteorológicas de Meteogalicia.

Este servicio se ha actualizado recientemente debido a un proceso de unificación del Índice de Calidad del Aire por parte de diferentes agencias meteorológicas, basado en los criterios de la Agencia Europea de Medio Ambiente [11]. Actualmente, cuenta con seis etiquetas con una percepción positiva, neutra o negativa (Tabla I).

Estas etiquetas están representadas por un código de colores en el que, por ejemplo, el color morado significa “pésima” mientras que el amarillo significa “buena”.

Para determinar el valor de calidad del aire adecuado para una situación, este servicio mide cinco contaminantes diferentes: dióxido de azufre (SO_2), dióxido de nitrógeno (NO_2), partículas en suspensión con un diámetro menor o igual a 2.5 micras ($PM_{2.5}$), partículas en suspensión con un diámetro entre 2.5 y 10 micras (PM_{10}) y ozono (O_3).

A partir de los datos obtenidos de las estaciones meteorológicas, Meteogalicia proporciona representaciones gráficas de la distribución de cada valor de calidad del aire en tiempo real para todas las estaciones gallegas. Meteogalicia representa estos valores a través de un gráfico de sectores que incluye una leyenda con la lista de estaciones con índice de calidad del aire “malo” o “muy malo” junto con el contaminante causante de esa situación. En la Figura 3 se muestra un ejemplo de la distribución de las etiquetas del ICA, donde el 58 % de las estaciones tiene una calidad del aire “muy buena” y el 38 % de las estaciones tienen una calidad “buena”. Por otro lado, una estación tiene una calidad del aire “moderada” y

Tabla I
ETIQUETAS DEL ÍNDICE DE CALIDAD DEL AIRE CON SUS PERCEPCIONES E ÍNDICES NUMÉRICOS ASOCIADOS.

Percepción	Positivo		Neutro	Negativo		
Etiqueta	Buena	Favorable	Regular	Mala	Muy mala	Pésima
Índice	0	1	2	3	4	5

también una única estación tiene una situación “mala” debido al contaminante PM10.

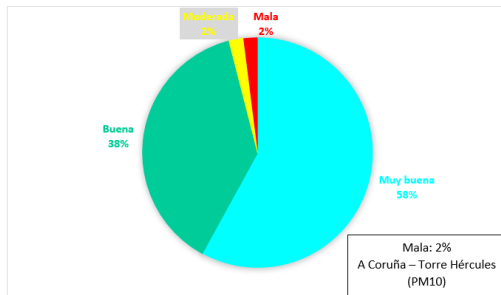


Figura 3. Distribución en tiempo real de los valores de calidad del aire.

III. DESCRIPCIONES LINGÜÍSTICAS DE LA DISTRIBUCIÓN DEL ÍNDICE DE CALIDAD DEL AIRE

Los cuantificadores lingüísticos en el lenguaje son una herramienta muy potente para representar y describir el conocimiento sobre la cantidad de elementos que cumplen determinadas propiedades, cuyo número, al menos desde el punto de vista pragmático, suele encontrarse entre una (cuantificadores unarios) y cuatro (cuantificadores cuaternarios). Las sentencias cuantificadas unarias tienen la siguiente estructura: “ Q X son S ” donde Q es un cuantificador (por ejemplo, “la mayoría”), X es un conjunto referencial (por ejemplo, “días”), y S es un valor lingüístico (por ejemplo, “lluvia”). Así, un ejemplo de enunciado cuantificado unario es “La mayoría de los días son lluviosos”.

En el caso de MeteoGalicia, se recogen en tiempo real datos de valores del índice de calidad del aire (ICA) de las 50 estaciones meteorológicas que componen su la Red [16]. Debido a la importancia que tiene esta información ya que puede afectar a la salud de las personas, surge el interés en proporcionar descripciones en lenguaje natural sobre los datos del ICA, que normalmente se presentan de forma gráfica. Por lo tanto, hemos definido en colaboración con los expertos de MeteoGalicia los requisitos de las descripciones lingüísticas que luego se proyectaron en las diferentes etapas de la arquitectura NLG.

III-A. Determinación de contenido

A partir de los datos del ICA de los que disponemos, surge el interés de generar descripciones en cuanto a la distribución de sus etiquetas a lo largo de la red de estaciones utilizando los porcentajes. Sin embargo, generar descripciones incluyendo porcentajes puede no ser atractivo para los usuarios, por lo tanto, para describir la distribución de las etiquetas del índice de

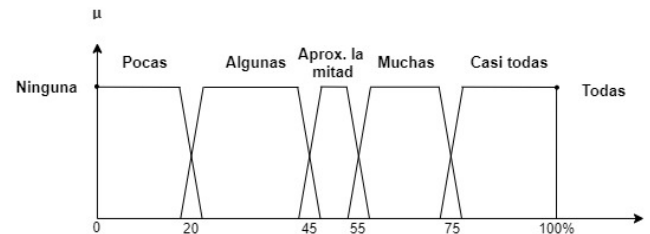


Figura 4. Representación gráfica de la definición de los cuantificadores para la descripción de la distribución del índice de calidad del aire (porcentaje de estaciones en la red meteorológica).

calidad del aire, utilizamos la partición de siete cuantificadores “Ninguna”, “Pocas”, “Algunas”, “Aproximadamente la mitad”, “Muchas”, “Casi todas”, “Todas” representados en la Figura 4, lo que permite verbalizar de forma imprecisa el porcentaje de estaciones que cumplen una etiqueta determinada. Esto resulta más amigable para los usuarios que un porcentaje, puesto que para este sistema no se demanda una descripción precisa.

Siendo Q los cuantificadores previamente definidos, X las estaciones de la red y S las etiquetas del ICA, generamos descripciones unarias, como, por ejemplo, “Muchas estaciones tienen una calidad del aire mala”.

A modo de ejemplo, partiendo de un conjunto de datos formado por la dupla ID de estación y valor del ICA para las 50 estaciones (p.e., 2, 3, 52, 1, 85, 1, 107, 2, 164, 2, ...), en esta fase se determinan las piezas de información que se van a incluir en la descripción textual:

- Muchas estaciones tienen calidad del aire muy mala.
- Muchas estaciones tienen calidad del aire mala.
- Rodís tiene calidad pésima debido al NO₂.
- Marraxón tiene calidad pésima debido a SO₂.
- Pastoriza tiene calidad pésima debido a PM10.
- Magdalena tiene calidad pésima debido a O₃.
- Río Cobo tiene calidad muy buena.
- Laza tiene calidad muy buena.
- Mourence tiene calidad muy buena.

III-B. Planificación del documento

Las descripciones en lenguaje natural están formadas por las siguientes partes:

- Descripción general que resume globalmente la situación, con una proposición cuantificada como “Muchas de las estaciones tienen una calidad del aire muy mala o mala”.
- Intensificación, que destaca las estaciones que presentan valores de calidad del aire que cumplan en mayor medida con la misma percepción que la descripción general. Esta

parte solo se incluye si hay alguna estación que cumpla esta condición. Por ejemplo, “Destaca Rodís, con una calidad del aire pésima debido al NO₂”.

- Excepción, que destaca las estaciones que tienen valores de calidad del aire cuya percepción es contraria a la de la descripción general. Esta parte solo se incluye si existe alguna estación con un valor destacable. Por ejemplo, “Por el contrario, Laza tiene una calidad muy buena”.

En el ejemplo anterior:

- Muchas estaciones tienen calidad del aire muy mala o mala.
- Rodís tiene calidad pésima debido al NO₂.
- Marraxón tiene calidad pésima debido a SO₂.
- Pastoriza tiene calidad pésima debido a PM₁₀.
- Magdalena tiene calidad pésima debido a O₃.
- Río Cobo tiene calidad muy buena.
- Laza tiene calidad muy buena.
- Mourence tiene calidad muy buena.

III-C. Planificación de la sentencia

Estas reglas de planificación están basadas en las máximas de Grice [17], que pueden resumirse como sigue:

- Calidad: las normas deben estar descritas con precisión.
- Cantidad: las reglas deben contener suficiente información para ser comprensibles sin proporcionar más información de la necesaria.
- Relevancia: las reglas deben contener las reglas importantes para este modelo.
- Forma: las reglas deben estar claramente definidas evitando descripciones ambiguas y deben estar ordenadas en función de la estructura de la descripción lingüística.

Las reglas para la descripción general consisten en:

- Resaltar si hay un valor del ICA mayoritario cuando cubre un porcentaje de estaciones superior a un umbral definido, por ejemplo “Muchas de las estaciones tienen una calidad del aire muy mala”¹.
- Destacar si hay dos etiquetas del ICA con la misma percepción que cubren un porcentaje de estaciones por encima de un umbral fijo, por ejemplo “‘Muchas de las estaciones tienen una calidad del aire muy mala’”.
- Si no hay un ICA predominante, se resaltan los peores valores, por ejemplo “Ourense tiene una calidad del aire extremadamente mala debido al contaminante PM₁₀”.

Para la intensificación y la excepción, la descripción debe seguir estas reglas:

- Una vez que se ha descrito el ICA predominante en la parte general de la descripción, el resto de valores se consideran intensificaciones si son etiquetas con la misma percepción o excepciones en caso contrario. Además, los valores con percepciones negativas incluyen los contaminantes que las producen. Por ejemplo “Todas las estaciones tienen una calidad del aire buena, especialmente

Ourense con una calidad del aire muy buena. Por el contrario, Paiosaco tiene una calidad del aire muy mala debido al contaminante O₃”.

- Si dos o más estaciones tienen el mismo ICA y causante se agrupan, por ejemplo “Laza y Santiago-Campus tienen una calidad del aire muy mala debido al contaminante PM₁₀”.

Siguiendo el ejemplo, en esta fase obtenemos:

- Muchas estaciones tienen calidad del aire muy mala o mala.
- Rodís tiene calidad pésima debido al NO₂.
- Marraxón tiene calidad pésima debido a SO₂.
- Pastoriza tiene calidad pésima debido a PM₁₀.
- Magdalena tiene calidad pésima debido a O₃.
- Río Cobo, Laza y Mourence tienen calidad muy buena.

III-D. Realización

En una descripción es tan importante el contenido como la forma de presentarlo, por lo que, con el objetivo de facilitar su lectura, se incluyen reglas de realización relacionadas con la estructura tanto para los casos de intensificación como para los de excepción: si el número de casos destacados es superior a 2, se dispondrán como una lista, en caso contrario se incluirán ambos como texto plano.

Debido a que la web de MeteoGalicia ofrece información tanto en castellano como en gallego, hemos elaborado las descripciones en ambos idiomas utilizando SimpleNLG-ES [18] y SimpleNLG-GL [19], que son versiones extendidas de SimpleNLG [20] para el castellano y el gallego, respectivamente.

La descripción textual completa resultante en el ejemplo anterior es:

Muchas estaciones tienen calidad del aire muy mala o mala. Destacan 4 estaciones:

- Rodís tiene calidad pésima debido al NO₂.
- Marraxón tiene calidad pésima debido a SO₂.
- Pastoriza tiene calidad pésima debido a PM₁₀.
- Magdalena tiene calidad pésima debido a O₃.

Por el contrario, Río Cobo, Laza y Mourence tienen una calidad muy buena.

IV. VALIDACIÓN

Para la validación de los sistemas de Generación de Lenguaje Natural se ha definido una metodología estándar que considera diferentes dimensiones [21]. En cuanto al objeto de la validación, tenemos: *i)* la validación intrínseca, que mide el rendimiento de un sistema en términos de su efectividad con respecto a los usuarios; y *ii)* la validación extrínseca, enfocada en medir la efectividad de un sistema a la hora de conseguir un determinado fin u objetivo. En cuanto a la forma de realizar la validación, tenemos: *i)* la validación manual, realizada a cargo de evaluadores humanos, que pueden ser tanto expertos como no expertos, en función del objetivo final del sistema, y *ii)* la validación automática, utilizando métricas de rendimiento, que son de uso unánime en sistemas de generación neuronal profundo.

¹El sistema ICA2Text genera descripciones en lenguaje natural bilingües en español y gallego, pero en este trabajo mostramos los ejemplos solo en español para facilitar la lectura.

Teniendo en cuenta estas consideraciones, y dadas las características del sistema ICA2Text optamos por una validación intrínseca manual realizada por expertos [22]. Así, tres meteorólogos expertos de la Red de Calidad del Aire de la Agencia de Meteorología de Galicia [12] (MeteoGalicia) participaron en la evaluación de la calidad de los descripciones en lenguaje natural en este ámbito y su adecuación para la descripción de la distribución del ICA. En la validación se les presentó al grupo de tres expertos un cuestionario compuesto por 25 ejemplos de casos reales de datos proporcionados por 20 estaciones meteorológicas diferentes, que evaluaron de forma conjunta, consensuando una única respuesta para cada ítem de cada caso. En la Figura 5 mostramos uno de los casos, donde se da una situación de ICA mala. Cada caso está compuesto por un diagrama de sectores y la descripción textual correspondiente generada por ICA2Text, acompañado de cinco ítems (mostradas en la Tabla II), que debían responder utilizando una escala Likert [23] de 5 puntos en el rango [1, 5] donde 1 indica acuerdo total del grupo de expertos con el ítem del cuestionario y 5 indica desacuerdo total. Los ítems hacen referencia a dos dimensiones principales: contenido de la descripción (Q1, Q2) y su forma (Q3, Q4, Q5). El equipo de expertos consensuó en todos los casos la evaluación numérica conjuntamente de la correspondencia entre las descripciones en lenguaje natural y la información proporcionada por el gráfico.

Tabla II
ÍTEMS QUE COMPONEN EL CUESTIONARIO PARA LA EVALUACIÓN POR PARTE DE LOS EXPERTOS.

Código	Cuestión
Q1	Indicar el grado de concordancia entre la descripción lingüística proporcionado y los datos representados en la figura
Q2	Indicar el grado de concordancia entre la descripción lingüística proporcionado y cómo describirías los datos
Q3	Indicar el grado de conformidad con el uso correcto del vocabulario
Q4	Indicar el grado de acuerdo con la organización de la descripción lingüística para facilitar su comprensión
Q5	Indique el grado de acuerdo con la ortografía, la puntuación y la estructura

Ninguno de estos expertos había participado en la definición del sistema ICA2Text ni de su modelo subyacente, puesto que todos los requisitos fueron definidos con la ayuda de un cuarto experto diferente, perteneciente también a la Red de Calidad del Aire de MeteoGalicia. Por lo tanto, se realizó una evaluación totalmente ciega, puesto que no se proporcionaron detalles de ningún tipo acerca de cómo se generaron las descripciones en lenguaje natural.

En la Tabla III se presenta un resumen de los resultados tras la evaluación por parte de los expertos de las dimensiones de contenido (Q1, Q2) y diseño (Q3, Q4, Q5) y el resultado global.

Los resultados muestran que los expertos están de acuerdo con las descripciones en lenguaje natural generadas, ya que la media de las puntuaciones es de 4.80 con una desviación típica de 0.51. Además, en todos los ítems la moda la máxima puntuación posible.

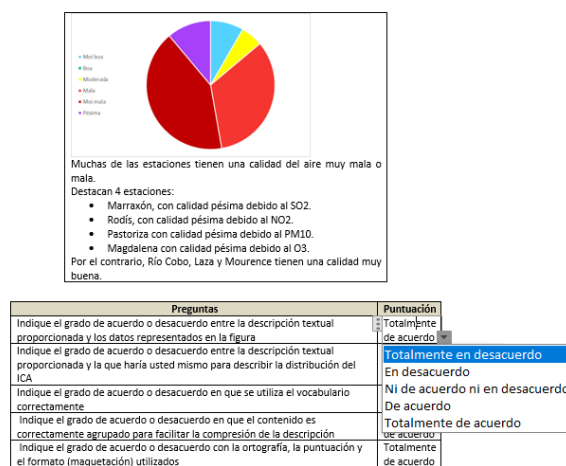


Figura 5. Ejemplo extraído del cuestionario diseñado para la validación de expertos.

Tabla III
RESULTADOS DE LA EVALUACIÓN POR EXPERTOS.

	Media	Desviación típica	Moda
Q1	4.76	0.52	5
Q2	4.68	0.69	5
Q3	4.88	0.33	5
Q4	4.76	0.60	5
Q5	4.92	0.28	5
Contenido	4.72	0.61	5
Diseño	4.85	0.43	5
General	4.80	0.51	5

En general, podemos concluir que estas descripciones lingüísticas generadas son muy adecuadas tanto en contenido como en forma, con medias de 4.72 y 4.85 respectivamente, para describir la distribución del índice de calidad del aire en las 20 estaciones.

A partir de los resultados tan positivos obtenidos en la validación por expertos, este sistema está actualmente en producción y será próximamente desplegado en la página web oficial de MeteoGalicia como un nuevo servicio público para los usuarios que consulten la información sobre la calidad del aire.

V. CONCLUSIONES

En este trabajo presentamos un modelo de generación de lenguaje natural empleando cuantificadores borrosos para la generación automática de descripciones lingüísticas a partir de datos numéricos obtenidos de un caso real de aplicación en el campo de la información medioambiental, proporcionando explicaciones textuales sobre la distribución de los valores del índice de calidad del aire (ICA), que es un indicador muy conocido proporcionado por todas las agencias meteorológicas del mundo. Basado en este modelo, describimos ICA2Text, un sistema D2T para la generación automática de descripciones en lenguaje natural sobre los datos del ICA proporcionados por la Agencia de Meteorología de Galicia.

Los resultados de la validación manual por parte de los expertos meteorólogos muestran que las descripciones en lenguaje natural generadas son muy adecuadas, puesto que en media calificaron las descripciones en lenguaje natural generadas por el sistema ICA2Text con un 4.72 sobre 5 en una escala Likert en cuanto a calidad del contenido y un 4.85 sobre 5 en cuanto a calidad lingüística (diseño).

Actualmente, ICA2Text se ha integrado en la web de producción de MeteoGalicia, de modo que tras un periodo de pruebas, como trabajo futuro se plantea el despliegue como servicio público en la sección de calidad del aire de la página web de MeteoGalicia.

AGRADECIMIENTOS

Este trabajo fue financiado por el Ministerio de Ciencia, Innovación y Universidades (proyecto TIN2017-84796-C2-1-R) y la Consellería de Educación, Universidad e Formación Profesional de la Xunta de Galicia (proyectos ED431C2018/29 y ED431G2019/04). Todos los proyectos fueron cofinanciados por el programa FEDER.

REFERENCIAS

- [1] E. Reiter, "Has a Consensus NL Generation Architecture Appeared, and is it Psycholinguistically Plausible?" in *Proceedings of the Seventh International Workshop on Natural Language Generation, INLG 1994, Kennebunkport, Maine, USA, June 21-24, 1994*, 1994. [Online]. Available: <https://www.aclweb.org/anthology/W94-0319/>
- [2] —, "An architecture for data-to-text systems," in *Proceedings of the Eleventh European Workshop on Natural Language Generation*. Association for Computational Linguistics, 2007, pp. 97–104. [Online]. Available: <https://doi.org/10.3115%2F1610163.1610180>
- [3] J. Kacprzyk and R. R. Yager, "Linguistic summaries of data using fuzzy logic," *International Journal of General System*, vol. 30, no. 2, pp. 133–154, 2001.
- [4] L. A. Zadeh, "A prototype-centered approach to adding deduction capability to search engines-the concept of protoform," in *Intelligent Systems, 2002. Proceedings. 2002 First International IEEE Symposium*, vol. 1. IEEE, 2002, pp. 2–3. [Online]. Available: <https://doi.org/10.1109%2Fis2002.1044219>
- [5] R. R. Yager, K. M. Ford, and A. J. Cañas, "An Approach to the Linguistic Summarization of Data," in *Uncertainty in Knowledge Bases, 3rd International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems, IPMU '90, Paris, France, July 2-6, 1990, Proceedings*, ser. Lecture Notes in Computer Science, vol. 521. Springer, 1990, pp. 456–468. [Online]. Available: <https://doi.org/10.1007/BFb0028132>
- [6] L. A. Zadeh, "A New Direction in AI - Toward a Computational Theory of Perceptions," in *Computational Intelligence, Theory and Applications, International Conference, 7th Fuzzy Days, Dortmund, Germany, October 1-3, 2001, Proceedings*, ser. Lecture Notes in Computer Science, vol. 2206. Springer, 2001, p. 628. [Online]. Available: https://doi.org/10.1007/3-540-45493-4_63
- [7] A. Wilbik, J. M. Keller, and G. L. Alexander, "Linguistic summarization of sensor data for eldercare," in *Proceedings of the IEEE International Conference on Systems, Man and Cybernetics, Anchorage, Alaska, USA, October 9-12, 2011*. IEEE, 2011, pp. 2595–2599. [Online]. Available: <https://doi.org/10.1109/ICSMC.2011.6084067>
- [8] J. Kacprzyk and A. Wilbik, "A comprehensive comparison of time series described by linguistic summaries and its application to the comparison of performance of a mutual fund and its benchmark," in *FUZZ-IEEE 2010, IEEE International Conference on Fuzzy Systems, Barcelona, Spain, 18-23 July, 2010, Proceedings*. IEEE, 2010, pp. 1–8. [Online]. Available: <https://doi.org/10.1109/FUZZY.2010.5584881>
- [9] P. Conde-Clemente, J. M. Alonso, and G. Triviño, "Toward automatic generation of linguistic advice for saving energy at home," *Soft Comput.*, vol. 22, no. 2, pp. 345–359, 2018. [Online]. Available: <https://doi.org/10.1007/s00500-016-2430-5>
- [10] A. Ramos-Soto, A. Bugarin, S. Barro, and J. Taboada, "Linguistic descriptions for automatic generation of textual short-term weather forecasts on real prediction data," *IEEE Transactions on Fuzzy Systems*, vol. 23, no. 1, pp. 44–57, 2015. [Online]. Available: <https://doi.org/10.1109%2Ftffuzz.2014.2328011>
- [11] A. Europea de Medio Ambiente, "Página web de la Agencia Europea de Medio Ambiente," 2021, [Accedido Mayo 2021]. [Online]. Available: www.eea.europa.eu/themes/air/air-quality-index
- [12] MeteoGalicia, "Página web de la agencia de meteorología gallega meteoGalicia," 2021, [Accedido Mayo 2021]. [Online]. Available: www.meteogalicia.gal
- [13] S. Busemann and H. Horacek, "Generating air quality reports from environmental data," in *Proceedings of the DFKI Workshop on Natural Language Generation*, 1997, pp. 15–21.
- [14] L. Wanner, B. Bohnet, N. Bouayad-Agha, F. Lareau, and D. Nicklaß, "MARQUIS: Generation of User-Tailored Multilingual Air Quality Bulletins," *Appl. Artif. Intell.*, vol. 24, no. 10, pp. 914–952, 2010. [Online]. Available: <https://doi.org/10.1080/08839514.2010.529258>
- [15] A. Ramos-Soto, A. Bugarín, S. Barro, N. Gallego, C. Rodríguez, I. Fraga, and A. Saunders, "Automatic Generation of Air Quality Index Textual Forecasts Using a Data-To-Text Approach," in *Advances in Artificial Intelligence - 16th Conference of the Spanish Association for Artificial Intelligence, CAEPIA 2015, Albacete, Spain, November 9-12, 2015, Proceedings*, ser. Lecture Notes in Computer Science, vol. 9422. Springer, 2015, pp. 164–174. [Online]. Available: https://doi.org/10.1007/978-3-319-24598-0_15
- [16] MeteoGalicia, "Página web de la calidad del aire de la agencia de meteorología gallega meteoGalicia," 2021, [Accedido Mayo 2021]. [Online]. Available: www.meteogalicia.gal/Caire
- [17] H. P. Grice, "Logic and conversation," in *Speech acts*. Brill, 1975, pp. 41–58.
- [18] A. Ramos-Soto, J. J. Gallardo, and A. Bugarín, "Adapting SimpleNLG to Spanish," in *Proceedings of the 10th International Conference on Natural Language Generation, INLG 2017, Santiago de Compostela, Spain, September 4-7, 2017*. Association for Computational Linguistics, 2017, pp. 144–148. [Online]. Available: <https://doi.org/10.18653/v1/w17-3521>
- [19] A. Cascallar-Fuentes, A. Ramos-Soto, and A. Bugarín, "Adapting SimpleNLG to Galician language," in *Proceedings of the 11th International Conference on Natural Language Generation, Tilburg University, The Netherlands, November 5-8, 2018*. Association for Computational Linguistics, 2018, pp. 67–72. [Online]. Available: <https://doi.org/10.18653/v1/w18-6507>
- [20] A. Gatt and E. Reiter, "SimpleNLG: A Realisation Engine for Practical Applications," in *ENLG 2009 - Proceedings of the 12th European Workshop on Natural Language Generation, March 30-31, 2009, Athens, Greece*. The Association for Computational Linguistics, 2009, pp. 90–93. [Online]. Available: <https://www.aclweb.org/anthology/W09-0613/>
- [21] C. Barros, "Proposal of a Hybrid Approach for Natural Language Generation and its Application to Human Language Technologies," 2019.
- [22] D. M. Howcroft, A. Belz, M.-A. Clinciu, D. Gkatzia, S. A. Hasan, S. Mahamood, S. Mille, E. van Miltenburg, S. Santhanam, and V. Rieser, "Twenty Years of Confusion in Human Evaluation: NLG Needs Evaluation Sheets and Standardised Definitions," in *Proceedings of the 13th International Conference on Natural Language Generation*. Dublin, Ireland: Association for Computational Linguistics, Dec. 2020, pp. 169–182. [Online]. Available: <https://www.aclweb.org/anthology/2020.inlg-1.23>
- [23] C. van der Lee, A. Gatt, E. van Miltenburg, S. Wubben, and E. Krahmer, "Best practices for the human evaluation of automatically generated text," in *Proceedings of the 12th International Conference on Natural Language Generation, INLG 2019, Tokyo, Japan, October 29 - November 1, 2019*. Association for Computational Linguistics, 2019, pp. 355–368.

Evaluación empírica de modelos de cuantificación borrosa aplicados a un agente conversacional

Mariña Canabal-Juanatey, Jose M. Alonso-Moral, Alejandro Catala, Alberto Bugarín-Diz

Centro Singular de Investigación en Tecnoloxías Intelixentes (CiTIUS),

Universidade de Santiago de Compostela, 15782, Santiago de Compostela, SPAIN

marina.canabal@rai.usc.es, {josemaria.alonso.moral, alejandro.catala, alberto.bugarin.diz}@usc.es

Resumen—En este trabajo presentamos un estudio empírico orientado a poner en evidencia posibles inconsistencias que pueden producirse cuando utilizamos expresiones en lenguaje natural con cuantificadores imprecisos, y éstas son evaluadas con modelos de cuantificación borrosa que incumplen dos propiedades de interés que se han definido en la bibliografía para dichos modelos (antonimia y efecto acumulativo). La experimentación considera el modelo de cuantificación de Zadeh en el contexto de un agente conversacional que realiza recomendaciones. Los resultados muestran que hay diferencias significativas en la percepción de la consistencia y la utilidad de las conversaciones mantenidas entre una persona usuaria y el agente conversacional, relacionadas con el modelo de cuantificación utilizado.

Index Terms—Cuantificadores borrosos, Generación de lenguaje natural, Evaluación humana, Agentes conversacionales

I. INTRODUCCIÓN

La presencia de la cuantificación en el uso del lenguaje natural humano es continua y juega un papel fundamental tanto en términos de significado como de expresividad. A menudo empleamos expresiones imprecisas como “*algunos de los alumnos sacaron buenas notas*”, “*en la mayor parte de los ayuntamientos gallegos la incidencia de coronavirus es baja*” o “*habrá pocos días con cielos nublados esta semana*”, que nos permiten manejar la vaguedad del lenguaje y resumir la información de forma natural. Debido a esta constante aparición de expresiones cuantificadas en el lenguaje natural, su modelado se ha convertido en un problema de gran interés en múltiples campos relacionados con la Inteligencia Artificial.

El objetivo principal de los cuantificadores es permitirnos caracterizar y describir propiedades cuantitativas sobre un conjunto de elementos (un colectivo, denominado *referencial*), en lugar de tener que hacerlo sobre cada elemento individual. La forma en que los datos son agregados y resumidos en una sentencia cuantificada depende directamente del modelo de cuantificación que se utilice para procesar la información. En caso de no seleccionarse un modelo adecuado, podría presentarse un resultado que no fuera representativo de los datos o, directamente, incorrecto.

Para caracterizar el comportamiento de un modelo de cuantificación, en el sentido de que los resultados obtenidos supongan una interpretación coherente y apropiada de los datos descritos, varios autores han propuesto una serie de propiedades intuitivas y plausibles que, en principio, cualquier modelo debería satisfacer [1]–[3]. Sin embargo, la mayor parte de los modelos de cuantificación más conocidos y utilizados

(por ejemplo, Zadeh [4], Yager [5], etc.) incumplen algunas de estas propiedades, por lo que presentan un comportamiento alejado del adecuado, y no deseado en algunos casos.

A pesar de que los modelos de cuantificación borrosa son bien conocidos desde hace tiempo (principio de los años 80 del siglo pasado), y que en diversos trabajos posteriores se ha llevado a cabo un análisis de las propiedades que cumplen a nivel teórico [1], [2], existen muy pocos estudios empíricos que, desde una perspectiva pragmática, midan el alcance o impacto del incumplimiento de dichas propiedades a nivel práctico en ámbitos concretos [6]. Este aspecto pragmático es de especial relevancia y actualidad, especialmente a partir del uso de cuantificadores en las descripciones lingüísticas de datos [7]–[9] y en los sistemas de generación de lenguaje natural [10], puesto que la pragmática es esencial en el lenguaje humano y el incumplimiento de algunas propiedades relevantes por parte de los cuantificadores puede dar lugar a inconsistencias en la operativa interna de las aplicaciones en estos u otros ámbitos que conviene hacer aflorar. La principal aportación de este trabajo se centra precisamente en aportar evidencias que ilustren el alcance del incumplimiento de estas propiedades y, de alguna manera, hagan transparente hacia los usuarios o desarrolladores su existencia.

Así, en este trabajo presentamos un estudio empírico que evidencia las posibles inconsistencias que pueden producirse en la práctica debido al incumplimiento de ciertas propiedades por parte del modelo de cuantificación utilizado. En concreto, nos centramos en el modelo de cuantificación escalar de Zadeh [4], para el cual estudiamos el alcance del incumplimiento de la propiedad de antonimia y el problema del efecto acumulativo, en un escenario de aplicación concreto. Para ello, hemos creado el asistente virtual *Quanversa*¹ [11], que desarrolla conversaciones breves sobre la predicción meteorológica en lenguaje natural a partir de datos [10], [12] que incluyen expresiones cuantificadas borrosas. Con fragmentos de estas conversaciones hemos realizado un estudio basado en cuestionarios, que nos ha permitido evaluar el impacto que tiene el efecto acumulativo y el incumplimiento de la propiedad de antonimia por parte del modelo de cuantificación de Zadeh en cuanto a la coherencia del diálogo mantenido entre el asistente y una persona usuaria, cuando se utilizan en él expresiones cuantificadas imprecisas.

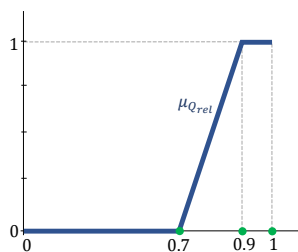
¹<https://demos.citius.usc.es/Quanversa/>

El resto del trabajo está organizado en cuatro secciones. La sección II introduce brevemente los modelos de cuantificación borrosa considerados. La sección III presenta el nuevo modelo de cuantificación propuesto. La sección IV describe los experimentos realizados. Finalmente, la sección V resume las principales conclusiones.

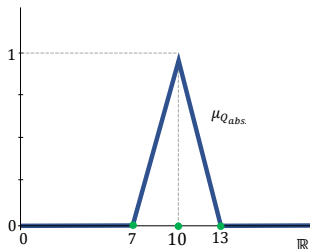
II. MODELOS DE CUANTIFICACIÓN BORROSA

En el lenguaje natural, es común encontrar términos vagos o imprecisos como “alto”, “joven” o “pequeño”, que expresan alguna propiedad sobre uno o más objetos y que pueden ser modelados mediante conjuntos borrosos. Además, la introducción de este tipo de términos dentro de expresiones cuantificadas es habitual (por ejemplo, “muchos trabajadores son jóvenes”, “la mayor parte de los jugadores son altos”, etc.), donde no solo la propiedad descrita presenta incertidumbre, sino que también lo hace el cuantificador lingüístico utilizado.

Así, en el modelado de la cuantificación lingüística, el manejo de la vaguedad presente en el lenguaje es fundamental. Fue Zadeh el primero en modelar el problema de la cuantificación en el lenguaje natural utilizando la teoría de conjuntos borrosos [4]. Para ello, distinguió dos tipos de cuantificadores lingüísticos (ver ejemplos en la figura 1): *cuantificadores borrosos absolutos*, los cuales denotan una cantidad imprecisa absoluta (“aproximadamente cinco”, “un número grande”, “alrededor de diez”, etc.) y *cuantificadores borrosos relativos* (o proporcionales), que referencian cantidades relativas (“la mitad”, “casi todos”, “una pequeña parte”, etc.). Además, propuso identificar ambos tipos de cuantificadores con un conjunto borroso que represente la semántica del cuantificador.



(a) Cuantificador relativo Q_{rel} : casi todos.



(b) Cuantificador absoluto Q_{abs} : alrededor de diez.

Figura 1: Ejemplos de cuantificadores.

Zadeh se centró en el estudio de dos tipos de expresiones: (1) las expresiones de tipo *I*, que siguen el esquema (o

protoforma) “ Q de E son $\tilde{A}$ ”; y (2) las expresiones de tipo *II*, que responden a “ Q de $(\tilde{D}E)$ son $\tilde{A}$ ”, donde E es el conjunto referencial, $\tilde{A}$ y $\tilde{D}$ dos conjuntos borrosos que representan alguna propiedad de E , y Q un cuantificador borroso absoluto o relativo. Un ejemplo de expresión de tipo *I* es “la mayoría de estudiantes de matemáticas son altos”, mientras que un ejemplo de expresión de tipo *II* podría ser “la mayoría de los buenos estudiantes de matemáticas son altos”, identificando en ambos casos E con el conjunto de estudiantes de matemáticas, Q con el cuantificador borroso relativo “la mayoría”, $\tilde{A}$ y $\tilde{D}$ con los conjuntos borrosos que representan las propiedades de “ser alto” y “ser buen estudiante”, respectivamente.

En lo que sigue, si $\tilde{A}$ es un conjunto borroso que representa una propiedad borrosa acerca de los elementos de un referencial $E = \{e_1, \dots, e_n\}$, escribiremos $\mu_{\tilde{A}}(E) = \{a_1, \dots, a_n\}$, con $a_i = \mu_{\tilde{A}}(e_i)$ para $i \in \{1, \dots, n\}$, representando los grados de cumplimiento de la propiedad $\tilde{A}$ sobre los elementos de E .

Siguiendo la aproximación propuesta por Zadeh, un *modelo de cuantificación* es un método que permite combinar el conjunto borroso que representa la semántica del cuantificador Q con los conjuntos borrosos que modelan las propiedades $\tilde{A}$ y $\tilde{D}$, de forma que sea posible obtener una medida de la veracidad o grado de cumplimiento de la expresión cuantificada. El modelo de cuantificación escalar propuesto por Zadeh en [4] es uno de los métodos más conocidos y utilizados en la práctica para la evaluación de expresiones cuantificadas. En el caso de expresiones de tipo *I*, se define, para cuantificadores borrosos absolutos, mediante la expresión:

$$Z_{Q_{abs}}(\tilde{A}) = \mu_{Q_{abs}}\left(\sum_{i=1}^n a_i\right), \quad (1)$$

y, para cuantificadores borrosos relativos, como:

$$Z_{Q_{rel}}(\tilde{A}) = \mu_{Q_{rel}}\left(\sum_{i=1}^n a_i/n\right). \quad (2)$$

En el caso de expresiones de tipo *II* se define, para cuantificadores borrosos absolutos, como:

$$Z_{Q_{abs}}(\tilde{A}/\tilde{D}) = \mu_{Q_{abs}}\left(\sum_{i=1}^n \min\{a_i, d_i\}\right), \quad (3)$$

y, para cuantificadores borrosos relativos, como:

$$Z_{Q_{rel}}(\tilde{A}/\tilde{D}) = \mu_{Q_{rel}}\left(\frac{\sum_{i=1}^n \min\{a_i, d_i\}}{\sum_{i=1}^n d_i}\right). \quad (4)$$

A partir de este primer método introducido por Zadeh, surgieron nuevas propuestas para la evaluación de expresiones cuantificadas [5], [13], [14], de forma que cada uno utiliza un esquema diferente para obtener el grado de verdad de una expresión cuantificada. Con el objetivo de poder comparar y evaluar el comportamiento de cada uno de estos modelos, en relación a la obtención de resultados coherentes y apropiados, varios autores han propuesto una serie de propiedades que, en principio, deberían satisfacer todos aquellos modelos que tengan un comportamiento consistente, intuitivo o plausible [1].

Entre ellas, se encuentran, por ejemplo, la continuidad, la monotonía, las relacionadas con la negación o la antonimia. Sin embargo, muchos de los modelos más utilizados en la práctica no verifican algunas de estas propiedades fundamentales. En concreto, el comportamiento teórico no adecuado del modelo de Zadeh debido al incumplimiento de ciertas propiedades matemáticas ha sido estudiado en profundidad en la bibliografía [1], [2]. En este trabajo, nos centraremos únicamente en dos de sus principales problemas: el incumplimiento de la propiedad de antonimia y el efecto acumulativo.

Debido a su continuo uso en el lenguaje natural, es fundamental que los modelos se comporten correctamente respecto a las negaciones. En este sentido, la propiedad de antonimia establece que todo modelo de cuantificación Γ debe verificar:

$$\Gamma_Q(\tilde{A}) = \Gamma_{Q_{ant}}(\tilde{A}_c) \quad y \quad (5)$$

$$\Gamma_Q(\tilde{A}/\tilde{D}) = \Gamma_{Q_{ant}}(\tilde{A}_c/\tilde{D}), \quad (6)$$

siendo $\tilde{A}_c$ el complementario de $\tilde{A}$, definido por:

$$\mu_{\tilde{A}_c}(e) = 1 - \mu_{\tilde{A}}(e) \text{ para todo } e \in E$$

y Q_{ant} el antónimo de Q , que se define, si Q es un cuantificador borroso absoluto, como:

$$\mu_{Q_{ant}}(x) = \mu_Q(n - x) \text{ para todo } x \in [0, n].$$

y, si Q es un cuantificador borroso relativo, como:

$$\mu_{Q_{ant}}(x) = \mu_Q(1 - x) \text{ para todo } x \in [0, 1].$$

La violación de la propiedad de antonimia puede tener consecuencias importantes en la evaluación de expresiones cuantificadas. Por ejemplo, dos expresiones como “*todos los días del fin de semana habrá temperaturas altas*” y “*no habrá días este fin de semana en los que las temperaturas no sean altas*”, que semánticamente son equivalentes, podrían obtener grados de verdad diferentes al ser evaluadas sobre un mismo referencial. Este comportamiento poco intuitivo no debería ocurrir en ningún caso, puesto que ambas expresiones en lenguaje natural realmente significan lo mismo.

El modelo de Zadeh no verifica la propiedad anterior, como se demuestra con el siguiente contraejemplo; sean E el conjunto de los próximos 4 días, $\tilde{A}$ y $\tilde{D}$ los conjuntos borrosos que representan las propiedades “*tener temperatura alta*” y “*pertenecer al fin de semana*”, de forma que $\mu_{\tilde{A}}(E) = \{0.5, 1, 1, 0.5\}$ y $\mu_{\tilde{D}}(E) = \{0.5, 1, 1, 0.5\}$. Definiendo el cuantificador $\forall :=$ “*todos*” como:

$$\mu_{\forall}(x) = \begin{cases} 1 & \text{si } x = 1, \\ 0 & \text{si } x \neq 1, \end{cases}$$

se tiene que:

$$\mathcal{Z}_{\forall}(\tilde{A}/\tilde{D}) = \mu_{\forall}\left(\frac{0.5 + 1 + 1 + 0.5}{0.5 + 1 + 1 + 0.5}\right) = \mu_{\forall}(1) = 1,$$

mientras que:

$$\begin{aligned} \mathcal{Z}_{\forall_{ant}}(\tilde{A}_c/\tilde{D}) &= \mu_{\forall_{ant}}\left(\frac{0.5 + 0 + 0 + 0.5}{0.5 + 1 + 1 + 0.5}\right) = \\ &= \mu_{\forall_{ant}}(1/3) = \mu_{\forall}(1 - 1/3) = 0. \end{aligned}$$

Por otro lado, el modelo de Zadeh puede provocar un efecto acumulativo poco intuitivo, ya que es posible que varios elementos con un grado de pertenencia bajo alcancen el mismo valor acumulado que un único elemento con un grado de pertenencia alto. Por ejemplo, supongamos que E y E' son los conjuntos de temperaturas para los próximos 5 días en dos ayuntamientos diferentes y $\tilde{A}$ es el conjunto borroso que representa la propiedad de “*ser una temperatura baja*” de forma que $\mu_{\tilde{A}}(E) = \{0.2, 0.2, 0.2, 0.2, 0.2\}$ y $\mu_{\tilde{A}}(E') = \{1, 0, 0, 0, 0\}$. Definiendo el cuantificador absoluto “*al menos uno*” como:

$$\mu_{\text{“al menos uno”}}(x) = \begin{cases} 0 & \text{si } x < 1, \\ 1 & \text{si } x \geq 1, \end{cases}$$

y utilizando la fórmula (1), la evaluación de la expresión “*al menos una temperatura será alta*” mediante el modelo de Zadeh devolverá el mismo resultado para ambos conjuntos referenciales E y E' , pero la situación en cada ayuntamiento es diferente.

A pesar de los problemas evidenciados a nivel teórico, no se conoce el impacto que puede tener el uso de un modelo de cuantificación que presente estas deficiencias durante la interacción entre un agente conversacional y sus usuarios.

III. MODELO DE ZADEH MODIFICADO

Hemos propuesto una modificación del modelo de Zadeh para corregir los potenciales efectos indeseados asociados a la antonimia y el efecto acumulativo [11], y de este modo disponer de dos modelos de cuantificación muy similares que permitiesen evaluar qué impacto tiene el incumplimiento de las propiedades mencionadas en un contexto de aplicación real. Se realizaron las modificaciones imprescindibles para abordar directamente las dos propiedades citadas.

Por un lado, a la hora de realizar una agregación, el modelo modificado solo considera los elementos que tengan un grado de pertenencia suficientemente alto. Esto significa que, en las definiciones (1)–(4), son excluidos aquellos términos de los sumatorios en los que los valores de a_i o d_i no alcancen el 0.5, evitando la acumulación de pequeños grados de pertenencia en la agregación. Se ha escogido el valor 0.5 por ser el umbral que determina la pertenencia al conjunto cuando la variable es desborrosificada. Así, los datos incluidos en la agregación serán aquellos considerados en todo momento como pertenecientes al conjunto.

Por otro lado, el incumplimiento de la propiedad de antonimia se debe a la incorrecta evaluación de una expresión cuantificada cuando ésta se formula en términos de $\tilde{A}_c$. Dado que semánticamente las expresiones “ Q de $(\tilde{D}E)$ son $\tilde{A}_c$ ” y “ Q_{ant} de $(\tilde{D}E)$ son $\tilde{A}$ ” son equivalentes, el modelo modificado realiza un procesamiento previo a la evaluación para transformar la primera en la segunda, cuando es necesario. Así, las expresiones de (5) y (6) se cumplen trivialmente y el modelo verifica la propiedad de antonimia.

En el estudio empírico que describimos en la siguiente sección utilizaremos este modelo de Zadeh modificado junto con el original para la evaluación de proposiciones cuantificadas borrosas sobre un mismo referencial. Se plantean escenarios de

conversación entre el agente conversacional *Quanversa* [11] y una persona usuaria en los cuales las propiedades de antonimia y efecto acumulativo juegan un papel relevante. Ello nos permitirá comparar, desde el punto de vista pragmático, el comportamiento del modelo modificado frente al original, y valorar experimentalmente la importancia en un escenario concreto del incumplimiento de las propiedades mencionadas.

IV. ESTUDIO EMPÍRICO

De un modo más formal, el estudio plantea un experimento que permita verificar la siguiente hipótesis de investigación: “*La interacción con un agente conversacional resulta más consistente y de mayor utilidad para los usuarios cuando los datos son procesados mediante un modelo de cuantificación que supera los problemas de la antonimia y el efecto acumulativo propios del modelo de cuantificación escalar propuesto originalmente por Zadeh [4]*”.

Con el fin de validar esta hipótesis, se construyó un cuestionario que permitió que los participantes en el estudio pudieran valorar la consistencia y la utilidad de una serie de conversaciones reales entre el agente conversacional *Quanversa* y una usuaria ficticia, María, mostradas mediante una captura de pantalla (véase una muestra de un escenario en la figura 2). El cuestionario incluye 6 escenarios distintos, que nos permiten considerar las diferentes situaciones o contextos de estudio en base al modelo considerado, tamaño de la muestra, y la propiedad a analizar (ver la figura 3). Se somete a los participantes a dichos escenarios (o estímulos), sin permitir interacción en tiempo real con el asistente, para evitar efectos de otros factores no controlados que no son objeto de estudio, y así dar respuesta a nuestra hipótesis de partida. El cuestionario se distribuyó a través de listas de distribución y redes sociales, y se mantuvo abierto desde el 6 hasta el 26 de abril de 2021. Durante este tiempo, fue completado por 132 personas, de forma totalmente voluntaria y anónima.

La figura 4 muestra la relación entre el modelo de cuantificación utilizado por *Quanversa* durante la conversación y la consistencia en sus respuestas apreciada por los participantes. Las conversaciones asociadas al nuevo modelo se consideraron más consistentes que en el caso del modelo original.

La figura 5 representa la relación entre el modelo de cuantificación y la utilidad de la conversación para la resolución de la tarea propuesta en el escenario, según los participantes en el estudio. De nuevo, se observan diferencias entre ambos modelos, aunque menos acentuadas que para el caso de la consistencia.

Para contrastar la hipótesis de investigación planteada (con un nivel de significación de 0.01), basta con comprobar la dependencia o independencia entre las variables respuesta y el modelo de cuantificación. Las hipótesis nula y alternativa para los contrastes de independencia entre el modelo de cuantificación y la consistencia/utilidad en cada contexto de estudio (véase C_i , con $i \in \{1, 2, 3\}$, en la figura 3) son:

H_0 : *La variable respuesta consistencia/utilidad es independiente del factor modelo de cuantificación en C_i .*

H_a : *El factor modelo de cuantificación influye en la variable respuesta consistencia/utilidad en C_i .*

Teniendo en cuenta lo anterior, se ha seleccionado el test de McNemar [15], un test de independencia no paramétrico utilizado como la alternativa a los test χ^2 de Pearson cuando los datos son pareados. El test de McNemar estudia si la probabilidad de evento positivo para una variable (en nuestro caso, conversación consistente o conversación útil) es igual en los dos niveles de otra variable (en nuestro caso, para los dos modelos de cuantificación). Su estadístico sigue una distribución χ^2 con 1 grado de libertad.

El test de McNemar aplicado a los tres contextos de estudio revela que existen diferencias significativas, que soportan rechazar H_0 en favor de que sí existe relación entre el modelo de cuantificación y la consistencia de las respuestas dadas por *Quanversa* (C1: $\chi^2 = 73.11$, p -valor < 0.01; C2: $\chi^2 = 81.01$, p -valor < 0.01; C3: $\chi^2 = 100.08$, p -valor < 0.01).

Análogamente, comprobamos la independencia de la variable utilidad respecto al modelo de cuantificación. De nuevo, para los tres contextos estudiados, el test de McNemar arroja diferencias significativas, rechazando H_0 en favor de que sí existe relación entre el modelo de cuantificación y la utilidad de la conversación mostrada entre *Quanversa* y la usuaria *María* (C1: $\chi^2 = 24.32$, p -valor < 0.01; C2: $\chi^2 = 37.21$, p -valor < 0.01; C3: $\chi^2 = 39.2$, p -valor < 0.01).

V. CONCLUSIONES

La aplicación del test no paramétrico de McNemar revela que las diferencias encontradas son significativas tanto en la percepción de la consistencia como en la percepción de la utilidad respecto al modelo de cuantificación considerado. Junto con la información recogida en las figuras 4 y 5, se puede afirmar que un porcentaje significativo de participantes encontró las interacciones con el agente conversacional más consistentes y de mayor utilidad cuando este procesó los datos mediante el modelo de cuantificación de Zadeh modificado (que cumple la propiedad de antonimia y supera el problema del efecto acumulativo). Por lo tanto, el experimento realizado valida la hipótesis de investigación planteada.

Como trabajo futuro nos planteamos ampliar el estudio a contextos de tipo conversacional en otros ámbitos y a otros modelos de cuantificación, de modo que se pueda extender la base experimental del presente estudio y generalizar las conclusiones que hemos presentado en este trabajo. Asimismo, estudiaremos con mayor profundidad dos aspectos pragmáticos que hemos observado en el análisis de resultados, y cuyo alcance merece atención y requiere investigación adicional. Por un lado, el hecho de que, con cierta frecuencia, los usuarios fueron capaces de tomar una decisión en relación a la tarea propuesta incluso en escenarios donde habían detectado inconsistencias en la conversación. Por otro, el hecho de que un buen número de participantes tardaron más en responder a los escenarios que involucraban al modelo de Zadeh, lo cual puede ser debido a la inconsistencia en la conversación que producía dicho modelo.

Escenario 6 de 6

María y sus amigos planean pasar el fin de semana en la localidad gallega de Noia y, **si hubiera algún día en el que las temperaturas máximas fueran altas, les gustaría disfrutar de la playa.**

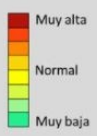
Con el objetivo de llevar el equipaje adecuado, antes de realizar los preparativos, María decide consultar el tiempo mediante el asistente virtual Quanversa. La conversación entre María y Quanversa se muestra en la siguiente imagen.

Quanversa:
Hola, ¿en qué puedo ayudarte?

María:
Pasaré el fin de semana en Noia, ¿cómo van a ser las temperaturas máximas?

Quanversa:
En general, las temperaturas máximas en Noia el fin de semana serán normales respecto a lo esperado para esta época del año.

	Viernes	Sábado	Domingo	Lunes
Temperatura máxima (°C)	18	18	18	18



María:
Entonces, ¿no va a haber temperaturas máximas altas?

Quanversa:
No, las temperaturas máximas serán normales todos los días del fin de semana en Noia.

Pregunta 1

¿Crees que las respuestas de Quanversa son consistentes entre sí? *

☐ Sí.

☒ No.

Explica, de forma breve, qué te ha llevado a dar la valoración anterior: *

Tu respuesta

Pregunta 2

Atendiendo únicamente a la información contenida en la conversación anterior, si fueras María, ¿crees que podrías tomar la decisión de ir o no a la playa? *

☐ Sí.

☒ No.

¿Qué te impide poder tomar una decisión? *

Tu respuesta

Figura 2: Ejemplo de escenario a evaluar en el cuestionario.

AGRADECIMIENTOS

Jose M. Alonso-Moral es investigador *Ramón y Cajal* (RYC-2016-19802). Alejandro Catalá es investigador *Juan de la Cierva* (IJC2018-037522-I). Este trabajo está financiado por el Ministerio de Ciencia, Innovación y Universidades (proyectos TIN2017-84796-C2-1-R y RTI2018-099646-B-I00) y por la Consellería de Educación, Universidade e Formación Profesional de la Xunta de Galicia (proyectos ED431F2018/02, ED431C2018/29, y ED431G2019/04). Todos los proyectos

han sido con financiados por el Fondo Europeo de Desarrollo Regional (FEDER).

REFERENCIAS

- [1] S. Barro, A. Bugarín, P. Cariñena, and F. Díaz-Hermida, "A framework for fuzzy quantification model analysis," *IEEE Transactions on Fuzzy Systems*, vol. 11, no. 1, pp. 89–99, 2003.
- [2] M. Delgado, M. D. Ruiz, D. Sánchez, and M. A. Vila, "Fuzzy quantification: a state of the art," *Fuzzy Sets Syst.*, vol. 242, pp. 1–30, 2014.
- [3] I. Glöckner, *Fuzzy Quantifiers: A computational Theory*. Springer, 2006.

	Modelo de Zadeh		Modelo modificado	
	Efecto acumulativo	Antonimia	Efecto acumulativo	Antonimia
Muestra pequeña	Escenario 1 (C1)	Escenario 3 (C3)	Escenario 4 (C1)	Escenario 6 (C3)
Muestra grande	Escenario 2 (C2)		Escenario 5 (C2)	

Figura 3: Escenarios seleccionados para el estudio.

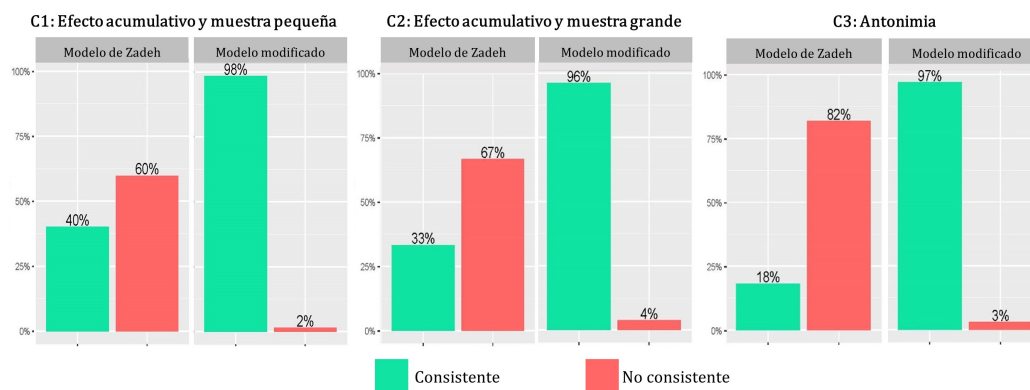


Figura 4: Relación entre la consistencia de la conversación y el modelo de cuantificación.

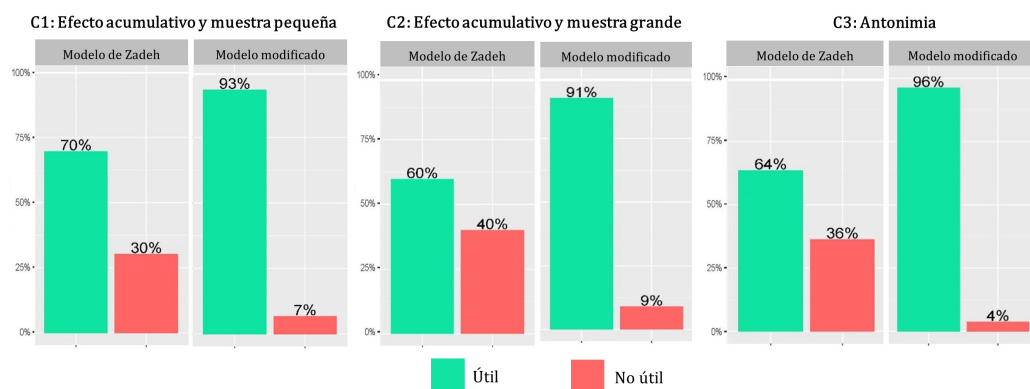


Figura 5: Relación entre la utilidad de la conversación y el modelo de cuantificación utilizado.

- [4] L. A. Zadeh, "A computational approach to fuzzy quantifiers in natural languages," *Computers & Mathematics with Applications*, vol. 9, no. 1, pp. 149–184, 1983.
- [5] R. R. Yager, "On ordered weighted averaging aggregation operators in multicriteria decisionmaking," *IEEE Transactions on Systems, Man, and Cybernetics*, vol. 19, no. 1, pp. 183–191, 1988.
- [6] A. Cascallar-Fuentes, A. Ramos-Soto, and A. Bugarín-Diz, "An experimental study on the use of fuzzy quantification models for linguistic descriptions of data," in *Proc. 24th European Conference on Artificial Intelligence*. IOS Press, 2020, pp. 267–274.
- [7] P. Conde-Clemente, J. M. Alonso, and G. Triviño, "Toward automatic generation of linguistic advice for saving energy at home," *Soft Comput.*, vol. 22, no. 2, pp. 345–359, 2018.
- [8] J. Kacprzyk and A. Wilbik, "A comprehensive comparison of time series described by linguistic summaries and its application to the comparison of performance of a mutual fund and its benchmark," in *Proc. IEEE International Conference on Fuzzy Systems*, 2010, pp. 1–8.
- [9] N. Marín and D. Sánchez, "On generating linguistic descriptions of time series," *Fuzzy Sets and Systems*, vol. 285, pp. 6–30, 2016, special Issue on Linguistic Description of Time Series.
- [10] A. Ramos-Soto, A. Bugarín, S. Barro, and J. Taboada, "Linguistic descriptions for automatic generation of textual short-term weather forecasts on real prediction data," *IEEE Transactions on Fuzzy Systems*, vol. 1, no. 23, pp. 44–57, 2015.
- [11] M. Canabal-Juanatey, *Quanversa: Modelos de cuantificación borrosa aplicados a un agente conversacional*. Universidad de Santiago de Compostela, Trabajo Fin de Grado, Ingeniería Informática, 2021.
- [12] A. Ramos-Soto, J. M. Alonso, E. Reiter, K. van Deemter, and A. Gatt, "Fuzzy-based language grounding of geographical references: From writers to readers," *International Journal of Computational Intelligence Systems*, vol. 12, no. 2, pp. 970–983, 2019.
- [13] P. Bosc and L. Lietard, "Monotonic quantified statements and fuzzy integrals," in *NAFIPS/IFIS/NASA Conference*, 1994, pp. 8–12.
- [14] M. Delgado, D. Sánchez, and M. A. Vila, "Fuzzy cardinality based evaluation of quantified sentences," *International Journal of Approximate Reasoning*, vol. 23, pp. 23–6, 2000.
- [15] Q. McNemar, "Note on the sampling error of the difference between correlated proportions or percentages," *Psychometrika*, vol. 12, no. 2, pp. 153–157, 1947.

Hypothesis Scoring and Model Refinement Strategies for FM-based RANSAC*

*Note: The full contents of this paper have been published in the volume Lecture Notes in Artificial Intelligence 12882 (LNAI 12882)

Alberto Ortiz, Esaú Ortiz, Juan José Miñana, Óscar Valero

Dep. Mathematics and Computer Science

University of the Balearic Islands, and IDISBA (Institut d'Investigació Sanitària de les Illes Balears)

Palma de Mallorca, Spain

{alberto.ortiz,esau.ortiz,jj.minana,o.valero}@uib.es

Abstract—Robust model estimation is a recurring problem in application areas such as robotics and computer vision. Taking inspiration from a notion of distance that arises in a natural way in fuzzy logic, this paper modifies the well-known robust estimator RANSAC making use of a Fuzzy Metric (FM) within the estimator main loop to encode the compatibility of each sample to the current model/hypothesis. Further, once a number of hypotheses have been explored, this FM-based RANSAC makes use of the same fuzzy metric to refine the winning model. The incorporation of this fuzzy metric permits us to express the distance between two points as a kind of degree of nearness measured with respect to a parameter, which is very appropriate in the presence of the vagueness or imprecision inherent to noisy data. By way of illustration of the performance of the approach, we report on the estimation accuracy achieved by FM-based RANSAC and other RANSAC variants for a benchmark comprising a large number of noisy datasets with varying proportion of outliers and different levels of noise. As it will be shown, FM-based RANSAC outperforms the classical counterparts considered.

Index Terms—Model estimation, RANSAC, Fuzzy metric, 2D straight line estimation

DOI: https://doi.org/10.1007/978-3-030-85713-4_10

Análisis de distintos tipos de tendencias en sucesiones de contextos L-fuzzy

Cristina Alcalde

Dep. Matemática Aplicada
Universidad del País Vasco - UPV/EHU
San Sebastián
c.alcalde@ehu.eus

Ana Burusco

Dep. Estadística, Informática y Matemáticas
Universidad Pública de Navarra
Instituto de Smart Cities
Pamplona
burusco@unavarra.es

Resumen—En este trabajo estudiamos el crecimiento, decrecimiento o ausencia de tendencias en las relaciones entre objetos y atributos en una sucesión de contextos *L-fuzzy*. Estas tendencias nos permitirán ordenar los objetos, atributos y conceptos atendiendo a su crecimiento o decrecimiento a lo largo del tiempo. Finalmente, ilustraremos nuestros resultados con una aplicación práctica.

Palabras Clave—Conceptos *L-Fuzzy*, sucesiones de contextos *L-Fuzzy*, análisis de tendencias.

I. INTRODUCCIÓN

El Análisis de Conceptos *L-Fuzzy* [1], [2] es una herramienta matemática para el análisis y la representación del conocimiento conceptual. Esta teoría utiliza los conceptos *L-fuzzy* para extraer información de un contexto *L-fuzzy*. Recordemos que un contexto *L-fuzzy* es una tupla (L, X, Y, R) , donde L es un retículo completo, X e Y son conjuntos de objetos y atributos, y $R \in L^{X \times Y}$ es una relación *L-fuzzy* entre los objetos y los atributos. Podemos entender el Análisis de Conceptos *L-fuzzy* como una extensión del Análisis de Conceptos Formales de Wille [3], [4] que permite trabajar con relaciones entre los objetos y atributos que toman valores en un retículo L , en lugar de valores binarios.

Para trabajar con estos contextos *L-fuzzy*, hemos definido los operadores derivación 1 y 2 por medio de las expresiones:

Para todo $A \in L^X$, para todo $B \in L^Y$

$$A_1(y) = \inf_{x \in X} \{I(A(x), R(x, y))\}, \forall y \in Y$$

$$B_2(x) = \inf_{y \in Y} \{I(B(y), R(x, y))\}, \forall x \in X$$

con I un operador de implicación fuzzy definido en el retículo $(L, \leq)$.

La información almacenada en el contexto se visualiza por medio de los conceptos *L-fuzzy*, que representan a un grupo de objetos que comparten, ellos y sólo ellos, un grupo de atributos. Estos conceptos son pares $(M, M_1) \in L^X \times L^Y$, donde el conjunto $M \in \text{fix}(\varphi)$ es un punto fijo del operador constructor φ , que se define a partir de los operadores derivación 1 y 2 como $\varphi(M) = (M_1)_2 = M_{12}$.

Trabajo parcialmente subvencionado por el Gobierno Vasco (Proyecto IT1256-19), y por el Grupo de Investigación "Inteligencia Artificial y Razonamiento Aproximado" de la Universidad Pública de Navarra.

Además, fijado un conjunto de partida $A \in L^X$ (o $B \in L^Y$), podemos obtener el concepto *L-fuzzy* asociados aplicando sucesivamente los operadores derivación hasta encontrar un punto fijo.

Para el caso de utilizar implicaciones residuadas, tal y como haremos en este trabajo, obtendremos un punto fijo en la segunda aplicación de los operadores derivación, con lo que el cálculo del concepto *L-fuzzy* asociado a un conjunto de partida se simplifica enormemente. Así, dado un conjunto de objetos $A \in L^X$ (o un conjunto de atributos $B \in L^Y$), el concepto *L-fuzzy* asociado será (A_{12}, A_1) (o (B_2, B_{21})).

En este trabajo trataremos de estudiar en qué medida las relaciones entre los objetos y los atributos mejoran o empeoran con el paso del tiempo. Representaremos estas situaciones mediante una secuencia de contextos *L-fuzzy* en la que profundizaremos en el estudio de las tendencias tanto crecientes como decrecientes. También estudiaremos los casos en los que no haya tendencia.

Existen trabajos en la literatura que analizan la evolución temporal en un contexto formal, por ejemplo, [5], [6]. En particular, en [6], Wolff introduce un Sistema de Tiempo Conceptual para definir el Análisis de Conceptos Temporales. En este Sistema de Tiempo Conceptual, el estado y la fase se definen como retículos de conceptos que representan el significado de los estados con respecto a la descripción elegida del tiempo. Además, otros autores definen tendencias de evolución en [5], utilizando temporal matching en el caso del Análisis de Conceptos Formales.

II. TENDENCIAS TEMPORALES

En [7] presentamos un primer estudio de sucesiones de contextos *L-fuzzy* considerando el retículo $L = [0, 1]$. Comenzaremos este apartado recordando cómo definíamos estas sucesiones:

Definición 1: Una sucesión de contextos *L-fuzzy* es una sucesión de tuplas (L, X, Y, R_i) , $i \in \{1, 2, \dots, n\}$, siendo L un retículo completo, X e Y dos conjuntos de objetos y atributos respectivamente y para cada valor $i \in \{1, \dots, n\}$ $R_i \in L^{X \times Y}$ representa la relación entre los conjuntos X e Y en el instante i , la cual toma valores en el retículo L .

Posteriormente, con el objeto de estudiar la evolución de la relación entre los objetos (o atributos) respecto de uno o

varios atributos (u objetos), en [8], [9] realizamos un análisis de tendencias temporales para identificar la evolución con el tiempo de la sucesión de contextos L -fuzzy $(L, X, Y, R_i), i \in \{1, \dots, n\}$.

Una buena herramienta para analizar la evolución en el tiempo de un objeto o un atributo es el estudio de sus conceptos L -fuzzy asociados en los diferentes contextos L -fuzzy de la sucesión. Ésta es la idea en la que se basan las siguientes definiciones:

Definición 2: Consideremos el objeto $x_0 \in X$. Sean $(A_{i\{x_0\}}, B_{i\{x_0\}})$ los conceptos L -fuzzy asociados a $\{x_0\}$ en la sucesión de contextos L -fuzzy (L, X, Y, R_i) con $i \leq n$. Denotamos por $ITrend(x_0)$ al conjunto de atributos cuya relación con el objeto x_0 se va fortaleciendo en la sucesión de contextos, y lo definimos como:

$$ITrend(x_0) = \{y \in Y \mid B_{1\{x_0\}}(y) \neq B_{n\{x_0\}}(y) \wedge B_{i\{x_0\}}(y) \leq B_{i+1\{x_0\}}(y), \forall i < n\}$$

Análogamente, dado el atributo $y_0 \in Y$, y los conceptos L -fuzzy $(A_{i\{y_0\}}, B_{i\{y_0\}})$ asociados a $\{y_0\}$ en la sucesión de contextos L -fuzzy, definiremos el conjunto $ITrend(y_0)$ que representa el conjunto de objetos que están cada vez más relacionados con el atributo y_0 de la siguiente manera:

$$ITrend(y_0) = \{x \in X \mid A_{1\{y_0\}}(x) \neq A_{n\{y_0\}}(x) \wedge A_{i\{y_0\}}(x) \leq A_{i+1\{y_0\}}(x), \forall i < n\}$$

En el caso de tendencias decrecientes, podemos definir los conjuntos de objetos o atributos cuyos grados de relación con un determinado atributo u objeto sea cada vez menor.

Definición 3: Fijado el objeto $x_0 \in X$ definimos el conjunto

$$DTrend(x_0) = \{y \in Y \mid B_{1\{x_0\}}(y) \neq B_{n\{x_0\}}(y) \wedge B_{i\{x_0\}}(y) \geq B_{i+1\{x_0\}}(y), \forall i < n\}$$

formado por los atributos cuyo grado de relación con dicho objeto va disminuyendo con el paso del tiempo.

Del mismo modo, para el atributo $y_0 \in Y$ obtenemos el conjunto de objetos cada vez menos relacionados con él:

$$DTrend(y_0) = \{x \in X \mid A_{1\{y_0\}}(x) \neq A_{n\{y_0\}}(x) \wedge A_{i\{y_0\}}(x) \geq A_{i+1\{y_0\}}(x), \forall i < n\}$$

Finalmente, tendremos que tener en cuenta que podría haber situaciones en las cuales no se observe una tendencia creciente ni decreciente. Para representar estos casos definimos los siguientes conjuntos:

Definición 4: El conjunto de objetos que no muestran ninguna tendencia en relación con el objeto $x_0 \in X$ es el siguiente:

$$NTrend(x_0) = \{y \in Y \mid y \notin ITrend(x_0) \cup DTrend(x_0)\}$$

Para el atributo $y_0 \in Y$, el conjunto de objetos sin ningún tipo de tendencia se define como:

$$NTrend(y_0) = \{x \in X \mid x \notin ITrend(y_0) \cup DTrend(y_0)\}$$

Es sencillo demostrar que, dados un objeto y un atributo cualesquiera del contexto L -fuzzy, se cumple la siguiente propiedad.

Proposición 1: Para todo par de elementos $(x, y) \in X \times Y$ se verifica que:

1. $x \in ITrend(y) \iff y \in ITrend(x)$
2. $x \in DTrend(y) \iff y \in DTrend(x)$

Estas definiciones de conjuntos de objetos (o de atributos) que presentan algún tipo de tendencia en su variación en el tiempo respecto de un atributo (o de un objeto) dado, permitirán establecer pares de objetos y atributos que pueden ser usados para un análisis más completo de la evolución de la sucesión de contextos L -fuzzy $(L, X, Y, R_i), i \in \{1, \dots, n\}$.

A partir de esta idea, se construyen las *matrices de tendencias* en el contexto L -fuzzy que indican la relación entre un objeto y un atributo cuando se observa algún tipo de tendencia relacionada con ellos.

Definición 5: La matriz de tendencia creciente $ITM \subseteq X \times Y$ se define como (ver [10]):

$$ITM(x, y) = \begin{cases} 1 & \text{si } y \in ITrend(x) \text{ (o equival. } x \in ITrend(y)) \\ 0 & \text{en otro caso} \end{cases}$$

De forma análoga podemos definir las matrices de tendencia decreciente y de ausencia de tendencia de la siguiente manera:

$$DTM(x, y) = \begin{cases} 1 & \text{si } x \in DTrend(y) \\ 0 & \text{en otro caso} \end{cases}$$

$$NTM(x, y) = \begin{cases} 1 & \text{si } y \in NTrend(x) \\ 0 & \text{en otro caso} \end{cases}$$

Estas matrices de tendencias nos permiten ahora definir los contextos formales (X, Y, ITM) , (X, Y, DTM) y (X, Y, NTM) y, a partir del análisis de sus respectivos conceptos formales, tener una visión general de las tendencias entre los objetos X y los atributos Y .

Definición 6: Sea el contexto formal (X, Y, ITM) con X conjunto de objetos, Y conjunto de atributos y $ITM \subseteq X \times Y$. Llamaremos a los conceptos de (X, Y, ITM) *conceptos formales ITrend*. Estos conceptos representarán las tendencias crecientes.

De la misma manera, podemos obtener los *conceptos formales DTrend* para las tendencias decrecientes y *conceptos formales NTrend* para los casos de ausencia de tendencias.

El cálculo de las matrices de tendencia nos va a permitir analizar aquellos objetos y atributos que mejoran o empeoran su relación con el tiempo. Sin embargo, no dispondremos de herramientas para cuantificar lo que supone esa mejora o empeoramiento.

En [10] y con el fin de medir el grado de evolución positiva de las relaciones entre los objetos y los atributos, se definía su nivel de tendencia como un valor en el intervalo $[0, 1]$. Esa manera de definir el nivel de tendencia tiene el inconveniente de que recoge el incremento que ha habido en los valores

de la relación entre un objeto y un atributo, pero no la relevancia absoluta del valor. Por este motivo, presentamos aquí una nueva definición para los niveles de tendencias que nos proporcionará una información más precisa.

Dado que los análisis de las tendencias crecientes y de las decrecientes son en muchos aspectos simétricos, por simplificar, escribiremos $\Lambda Trend$ para referirnos a cualquiera de las tendencias $ITrend$ o $DTrend$:

Definición 7: Para cada $x_0 \in X, y \in Y$, el nivel $\Lambda TrendL$ del objeto x_0 para el atributo y se define como:

$$\Lambda TrendL(x_0)_y = \begin{cases} [\min_{1 \leq i \leq n} B_{i\{x_0\}}(y), \max_{1 \leq i \leq n} B_{i\{x_0\}}(y)] & \text{si } y \in \Lambda Trend(x_0) \\ 0 & \text{en otro caso} \end{cases}$$

Análogamente, podemos definir para cada $y_0 \in Y, x \in X$ el nivel $\Lambda TrendL$ del atributo y_0 para el objeto x :

$$\Lambda TrendL(y_0)_x = \begin{cases} [\min_{1 \leq i \leq n} A_{i\{y_0\}}(x), \max_{1 \leq i \leq n} A_{i\{y_0\}}(x)] & \text{si } x \in \Lambda Trend(y_0) \\ 0 & \text{en otro caso} \end{cases}$$

Proposición 2: Para todo par de elementos $(x, y) \in X \times Y$ se verifica que:

$$\Lambda TrendL(x)_y = \Lambda TrendL(y)_x$$

A partir de la definición de los niveles de tendencias podemos establecer nuevas relaciones entre los objetos y los atributos de la siguiente manera:

Definición 8: Denominamos *relación de Λ -tendencia* a la siguiente relación intervalovalorada definida entre los conjuntos de objetos atributos.

Para cada $(x, y) \in X \times Y$,

$$\Lambda TrendLM(x, y) = \Lambda TrendL(x)_y = \Lambda TrendL(y)_x$$

III. RANKINGS DE OBJETOS, ATRIBUTOS Y CONCEPTOS ATENDIENDO A LAS TENDENCIAS

A partir de las definiciones de $ITrendLM$ y $DTrendLM$ para el caso de crecimiento y decrecimiento, vamos a intentar definir relaciones entre los objetos y los atributos de acuerdo con su evolución en el tiempo.

Denotaremos por $ITrendLM(x, y)$ y $\overline{ITrendLM}(x, y)$, respectivamente, al extremo inferior y superior de intervalo $ITrendLM(x, y)$. Análogamente para $DTrendLM(x, y)$.

Definición 9: Para cada $x \in X, y \in Y$, el porcentaje de crecimiento $ITrendLPer$ y de decrecimiento $DTrendLPer$ del objeto x para el atributo y se definen como:

$$ITrendLPer(x, y) = \begin{cases} \frac{ITrendLM(x, y) - \overline{ITrendLM}(x, y)}{1 - \overline{ITrendLM}(x, y)} & \text{si } x \in ITrend(y) \\ 0 & \text{en otro caso} \end{cases}$$

Del mismo modo,

$$DTrendLPer(x, y) = \begin{cases} \frac{DTrendLM(x, y) - \overline{DTrendLM}(x, y)}{DTrendLM(x, y)} & \text{si } x \in DTrend(y) \\ 0 & \text{en otro caso} \end{cases}$$

Agregando los porcentajes de crecimiento y, análogamente, de decrecimiento de cada objeto o atributo podremos analizar sus niveles porcentuales de tendencia.

Definición 10: Para cada $x_0 \in X, y_0 \in Y$ definimos su *nivel porcentual de tendencia creciente (o decreciente)* de la siguiente manera:

$$\begin{aligned} \Lambda TLP(x_0) &= Agr_{y \in Y}(\Lambda TrendLPer(x_0, y)) \\ \Lambda TLP(y_0) &= Agr_{x \in X}(\Lambda TrendLPer(x, y_0)) \end{aligned}$$

donde Agr es un operador de agregación.

Esta definición nos permite establecer relaciones de orden en los conjuntos de objetos y de atributos en función de sus niveles de tendencia:

Definición 11: Dados $z_i, z_j \in X$ o $z_i, z_j \in Y$, definimos las relaciones $<_{\Lambda TL}$ como $z_i <_{\Lambda TL} z_j$ si $\Lambda TLP(z_i) < \Lambda TLP(z_j)$.

De este modo, obtendremos rankings de crecimiento si trabajamos con $ITLP$. Aquellos objetos o atributos con mayores valores de $ITLP$ serán los que hayan evolucionado de forma más positiva a lo largo del tiempo dentro de su posibilidad de mejora. También los habrá de decrecimiento, si lo hacemos con $DTLP$.

Para obtener una información más completa sobre las tendencias que se pueden observar en la secuencia de contextos L -fuzzy, el siguiente paso consistirá en analizar la evolución conjunta de objetos y atributos.

Así, utilizando las matrices de tendencias ΛTM , estableceremos relaciones en los conceptos $Trend$ definidos a partir de los contextos formales $(X, Y, \Lambda TM)$ que nos permitan establecer qué conceptos son los que presentan unas tendencias más acentuadas.

Asignaremos, en primer lugar, un nivel de tendencia porcentual a cada uno de los conceptos formales de la siguiente manera:

Definición 12: Para cada (A, B) concepto formal del contexto formal $(X, Y, \Lambda TM)$ se define el nivel porcentual de Λ -tendencia del concepto como:

$$\Lambda TLPC(A, B) = \begin{cases} Agr_{(x, y) \in A \times B}(\Lambda TrendLPer(x, y)) & \text{si } A, B \neq \emptyset \\ 0 & \text{en otro caso} \end{cases}$$

donde Agr es un operador de agregación.

Esto nos permite obtener relaciones entre los conceptos formales de los contextos definidos a partir de las matrices de tendencias y establecer qué conceptos son los que presentan tendencias crecientes o decrecientes más acentuadas. Para ello utilizaremos la siguiente relación de orden:

Definición 13: Dados $(A, B), (C, D) \in X \times Y$ definimos la relación $<_{\Delta TLC}$ de modo que $(A, B) <_{\Delta TLC} (C, D)$ si $\Delta TLC(A, B) < \Delta TLC(C, D)$.

IV. APLICACIÓN PRÁCTICA

Para ilustrar la aplicación de los resultados vamos a considerar una sucesión de contextos L -fuzzy $(L, X, Y, R_i), i \in \{1, \dots, 5\}$ que recoge la estancia media de los clientes en distintos tipos de alojamientos turísticos en algunas regiones de España a lo largo de un periodo de cinco años.

El objetivo del estudio será analizar si ha habido algún tipo de tendencia creciente o decreciente en dichas estancias, teniendo en cuenta tanto los tipos de alojamientos como las regiones consideradas en el estudio.

En estos contextos L -fuzzy hemos considerado el conjunto de objetos formado por los distintos tipos de alojamientos $X = \{x_1=\text{Hotel}, x_2=\text{Camping}, x_3=\text{Apartamento Turístico}, x_4=\text{Casa Rural}\}$. Los atributos considerados son las distintas regiones en las que se ha realizado el estudio $Y = \{y_1=\text{Andalucía}, y_2=\text{Cataluña}, y_3=\text{Navarra}, y_4=\text{País Vasco}\}$ y los valores de las relaciones se corresponden con las estancias medias registradas desde 2016 hasta 2019. Los valores se han normalizado para trabajar con el retículo formado por la cadena $L = \{0, 0.01, \dots, 0.99, 1\}$ (ver Tabla I).

Analizando los conjuntos derivados obtenidos a partir de los distintos objetos o atributos, se obtienen las matrices de tendencia creciente ITM y de tendencia decreciente DTM que se muestran, respectivamente, en la Tabla II y Tabla III. Estas matrices nos permiten observar para qué objetos hemos tenido algún tipo de tendencia y en cuáles de los atributos se ha mantenido dicha tendencia.

Podemos observar, por ejemplo, que mientras la estancia media en hoteles (x_1) tuvo una tendencia creciente en Andalucía (y_1), en Navarra y País Vasco (y_3 e y_4) la tendencia fue decreciente.

Tras calcular las distintas matrices de tendencias, con el fin de conocer en qué regiones y en qué tipos de alojamientos se ha dado una tendencia más acentuada, vamos a establecer rankings entre los objetos y los atributos. Calculamos para ello los niveles porcentuales de tendencia tanto creciente como decreciente. Los valores obtenidos se muestran en la Tabla IV.

Estos niveles de tendencia nos permiten ordenar los objetos y atributos atendiendo a su mayor crecimiento o decrecimiento. La Figura 1 muestra la ordenación obtenida desde el objeto (o atributo) con una mayor tendencia creciente.

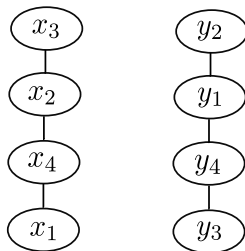


Figura 1. Rankings de objetos y atributos con tendencia creciente.

Podemos observar que el tipo de alojamiento en el que más ha crecido la estancia media ha sido en los apartamentos turísticos (x_3), mientras que la región que ha tenido una mayor subida en sus estancias vacacionales ha sido Cataluña (y_2).

Análogamente, observando el ranking obtenido a partir de los niveles de tendencia decreciente (ver Figura 2), vemos que el tipo de establecimiento que ha acusado una mayor bajada ha sido el camping (x_2). En cuanto a las regiones, la que ha sufrido un mayor descenso ha sido Navarra (y_3).

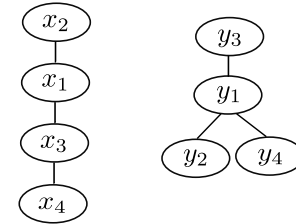


Figura 2. Rankings de objetos y atributos con tendencia decreciente.

Con el fin de tener una visión conjunta de alojamientos y regiones, a partir de las matrices de tendencia mostradas en la Tabla II, definimos los contextos formales (X, Y, ITM) e (X, Y, DTM) y obtenemos los niveles de tendencia asociados a sus distintos conceptos formales.

Para los conceptos formales del contexto (X, Y, ITM) se obtienen los niveles porcentuales de tendencia creciente:

$$\begin{aligned} ITLPC(X, \emptyset) &= 0 \\ ITLPC(\{x_1, x_3, x_4\}, \{y_1\}) &= 0,25 \\ ITLPC(\{x_2, x_3\}, \{y_2\}) &= 0,93 \\ ITLPC(\{x_4\}, \{y_1, y_3\}) &= 0,42 \\ ITLPC(\{x_3\}, \{y_1, y_2, y_4\}) &= 0,54 \\ ITLPC(\emptyset, Y) &= 0 \end{aligned}$$

A partir de estos niveles de tendencia obtenemos la prelación en el conjunto de conceptos, ordenados de mayor a menor crecimiento, representada en la Figura 3. Tal y como podemos

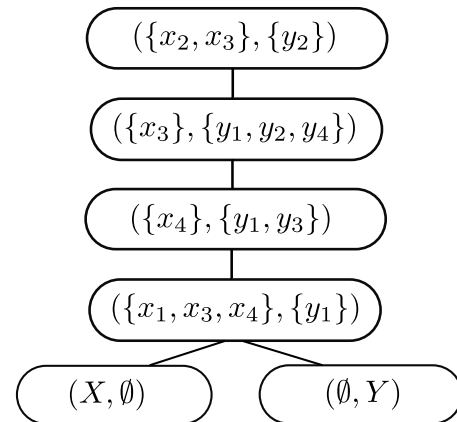


Figura 3. Ranking de conceptos con tendencia creciente.

Tabla I
RELACIONES DE LA SUCESIÓN DE CONTEXTOS L -FUZZY.

2015					2016					2017				
R_1	y_1	y_2	y_3	y_4	R_2	y_1	y_2	y_3	y_4	R_3	y_1	y_2	y_3	y_4
x_1	0.38	0.40	0.31	0.35	x_1	0.39	0.41	0.29	0.33	x_1	0.39	0.46	0.29	0.30
x_2	0.79	0.52	0.80	0.41	x_2	0.63	0.64	0.58	0.37	x_2	0.58	0.65	0.46	0.36
x_3	0.47	0.80	0.56	0.36	x_3	0.49	0.80	0.54	0.42	x_3	0.56	0.83	0.54	0.55
x_4	0.51	0.47	0.40	0.41	x_4	0.52	0.40	0.46	0.39	x_4	0.56	0.39	0.49	0.38

2018					2019				
R_4	y_1	y_2	y_3	y_4	R_5	y_1	y_2	y_3	y_4
x_1	0.43	0.44	0.22	0.30	x_1	0.47	0.42	0.18	0.28
x_2	0.52	0.83	0.45	0.42	x_2	0.50	1	0.45	0.39
x_3	0.56	0.87	0.47	0.59	x_3	0.58	0.97	0.35	0.72
x_4	0.60	0.37	0.52	0.42	x_4	0.71	0.37	0.66	0.43

Tabla II
MATRIZ DE TENDENCIA CRECIENTE.

ITM	y_1	y_2	y_3	y_4
x_1	1	0	0	0
x_2	0	1	0	0
x_3	1	1	0	1
x_4	1	0	1	0

Tabla III
MATRIZ DE TENDENCIA DECRECIENTE.

DTM	y_1	y_2	y_3	y_4
x_1	0	0	1	1
x_2	1	0	1	0
x_3	0	0	1	0
x_4	0	1	0	0

observar, las estancias que han mantenido una mayor tendencia creciente se han dado en campings y apartamentos turísticos (x_2 y x_3) de Cataluña (y_2), seguidas de las estancias en apartamentos turísticos (x_3) de Andalucía, Cataluña y País Vasco (y_1 , y_2 e y_4).

De manera análoga, podemos calcular los niveles porcentuales de tendencia para los conceptos formales del contexto

Tabla IV
NIVELES PORCENTUALES DE TENDENCIA.

$ITLP$			
x_1	0.04	y_1	0.19
x_2	0.25	y_2	0.46
x_3	0.41	y_3	0.11
x_4	0.21	y_4	0.14

$DTLP$			
x_1	0.15	y_1	0.09
x_2	0.20	y_2	0.05
x_3	0.09	y_3	0.31
x_4	0.05	y_4	0.05

(X, Y, DTM) definido a partir de la matriz de tendencia decreciente. Los valores obtenidos son los siguientes:

$$\begin{aligned}
 DTLPC(X, \emptyset) &= 0 \\
 DTLPC(\{x_4\}, \{y_2\}) &= 0,21 \\
 DTLPC(\{x_1, x_2, x_3\}, \{y_3\}) &= 0,41 \\
 DTLPC(\{x_2\}, \{y_1, y_3\}) &= 0,40 \\
 DTLPC(\{x_1\}, \{y_3, y_4\}) &= 0,31 \\
 DTLPC(\emptyset, Y) &= 0
 \end{aligned}$$

La ordenación de los conceptos de mayor a menor decrecimiento es la representada en la Figura 4.

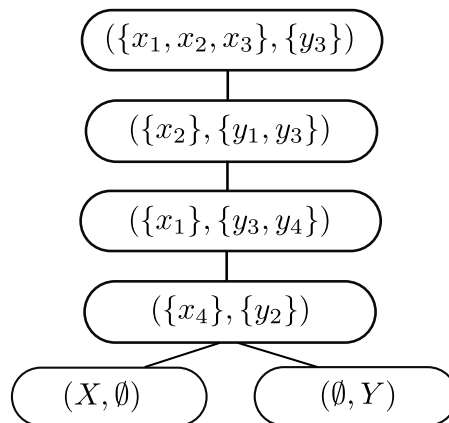


Figura 4. Ranking de conceptos con tendencia decreciente.

Observamos aquí que el mayor decrecimiento se ha dado en las estancias medias en hoteles, campings y apartamentos turísticos (x_1 , x_2 y x_3) de Navarra (y_3).

V. CONCLUSIONES Y TRABAJO FUTURO

Hemos continuado con el estudio de tendencias en una sucesión de contextos L -fuzzy que representa la evolución en el tiempo de un contexto L -fuzzy analizando las tendencias

crecientes, decrecientes y la ausencia de tendencias. Se ha analizado la relevancia de dichas tendencias temporales lo que ha permitido establecer rankings de objetos, atributos y conceptos formales. Como línea futura se estudiarán clasificaciones de objetos y de atributos a partir de conjuntos de partida y según su crecimiento o decrecimiento.

REFERENCIAS

- [1] A. Burusco, R. Fuentes-González: The Study of the L-fuzzy Concept Lattice, *Mathware and Soft Computing* 1 (3), pp. 209–218, 1994.
- [2] A. Burusco, R. Fuentes-González: Construction of the L-fuzzy Concept Lattice, *Fuzzy Sets and Systems* 97 (1), pp. 109–114, 1998.
- [3] B. Ganter, R. Wille: Formal concept analysis: Mathematical foundations, Springer, Berlin - New York, 1999.
- [4] R. Wille: Restructuring lattice theory: an approach based on hierarchies of concepts, in: *Rival I. (Ed.), Ordered Sets, Reidel, Dordrecht-Boston*, pp. 445–470, 1982.
- [5] R. Neouchi, A. Y. Tawfik, R. A. Frost: Towards a Temporal Extension of Formal Concept Analysis. In *Proceedings of Canadian Conference on Artificial Intelligence*, pp. 335–344, 2001.
- [6] K. E. Wolff: Temporal Relational Semantic Systems, *ICCS*, pp. 165–180, 2010.
- [7] C. Alcalde, A. Burusco, R. Fuentes-González: The study of fuzzy context sequences *Int. Journal of Computational Intelligence Systems* 6 (3), pp. 518–529, 2013.
- [8] C. Alcalde, A. Burusco: L-fuzzy context sequences on complete lattices, *IPMU 2014, Lecture Notes in Computer Science Part III, CCIS 444*, pp. 31–40, 2014.
- [9] C. Alcalde, A. Burusco, H. Bustince, A. Jurio, J.A. Sanz: Evolution in time of the L-fuzzy context sequences, *Information Sciences* 326, pp. 202–214, 2016.
- [10] C. Alcalde, A. Burusco: Estudio sobre la evolución de la relación entre objetos y atributos en una sucesión de contextos L-fuzzy, *XVIII Conferencia de la Asociación Española para la Inteligencia Artificial, CAEPIA 2018, Granada*, pp. 273–278, 2018.

On the standard fuzzy metric: generalizations and application to model estimation

Juan José Miñana , Alberto Ortiz , Esaú Ortiz and Óscar Valero 

Department of Mathematics and Computer Science, University of the Balearic Islands and IDISBA, Palma, Spain
 {jj.minana, alberto.ortiz, esau.ortiz, o.valero}@uib.es

Abstract—Different approaches to obtain a notion of metric in the context of fuzzy setting can be found in the literature. In this paper, we deal with the concept due to George and Veeramani, which is defined by means of continuous triangular norms. Different authors have addressed the study of such a concept from a theoretical point of view. In this paper, we provide a new methodology to induce fuzzy metrics which generalize the celebrated standard fuzzy metric. The aforementioned methodology allows us to approach some questions related to the continuous triangular norms from which such fuzzy metrics are defined. Moreover, we show the applicability of the new fuzzy metrics to an engineering problem. More specifically, we address successfully robust model estimation through a variant of the well-known estimator RANSAC. By way of illustration of the performance of the approach, we report on the accuracy achieved by the new estimator and other RANSAC variants for a benchmark involving a specific model estimation problem and a large number of datasets with varying proportion of outliers and different levels of noise. The resulting estimator is shown able to outperform the classical counterparts considered.

Index Terms—Fuzzy metric; continuous t -norm; Dombi t -norm; standard fuzzy metric; model estimation; RANSAC

I. INTRODUCTION AND PRELIMINARIES

In 1965, L. A. Zadeh introduced the notion of fuzzy set in [1]. Since then, such a concept has constituted the grounds of many lines of research in different fields, such as Mathematics, Computer Science, Economics. In Mathematics and, in particular, in Topology, an interesting issue consists in providing a notion of metric, in the fuzzy setting, in accordance with the essence of the classical concept. With this aim, in [2], I. Kramosil and J. Michalek introduced a notion of fuzzy metric space by adapting the concept of statistical metric due to Menger (see [3]) to the fuzzy context. Later on, in [4], A. George and P. Veeramani slightly modified the notion of Kramosil and Michalek with the aim of obtaining a more faithful adaptation to the fuzzy setting of the classical concept of metric. In both cases, the concept of fuzzy metric is defined by means of continuous t -norms (see [5] to find a deep treatment on t -norms). Following [4], a *fuzzy metric space* is a triplet $(X, M, *)$ where X is a non-empty set, $*$ is a continuous t -norm and M is a fuzzy set on $X \times X \times]0, \infty[$ satisfying, for each $x, y, z \in X$ and $t, s \in]0, \infty[$, the following:

- (GV1) $M(x, y, t) > 0$;
- (GV2) $M(x, y, t) = 1$ if and only if $x = y$;
- (GV3) $M(x, y, t) = M(y, x, t)$;
- (GV4) $M(x, z, t + s) \geq M(x, y, t) * M(y, z, s)$;
- (GV5) The assignment $M_{x,y} :]0, \infty[\rightarrow]0, 1]$ is a continuous function.

As usual, we say that $(M, *)$, or simply M if no confusion arises, is a fuzzy metric on X .

On account of the previous definition, the value of $M(x, y, t)$ can be interpreted as a degree of nearness between the point x and y of X with respect to the parameter $t \in]0, \infty[$. Then, the closer to 1 is such a value, the nearer the points x and y with respect to t are. Contrarily, values close to 0 indicate a lower degree of nearness. Thus, in this notion of fuzzy metric, 1 plays a similar role to 0 for the classical case, whereas 0 can be seen as ∞ in classical metrics. So, axiom (GV1) is justified by the fact that the degree of nearness with respect to a parameter never can be zero, just as in the classical case the distance between two points cannot become ∞ .

One can easily identify (GV2), (GV3) and (GV4) as fuzzy versions of the axioms of, respectively, separation, symmetry and transitivity, which altogether define the notion of classical metric. Concretely, (GV2) means that, on the one hand, the degree of nearness between two points with respect to an arbitrary parameter only can be 1 whenever both points are the same. On the other hand, the degree of nearness between a point and itself is 1, with respect to an arbitrary parameter. Finally, (GV5) ensures that no drastic changes arise in the degree of nearness due to slight modifications of the parameter with respect to which it is being measured.

An immediate consequence of (GV4), which was pointed out by M. Grabiec for fuzzy metrics in the sense of Kramosil and Michalek (see [6]), is that the degree of nearness between two points does not decrease when the parameter for which such a degree is relative increases, i.e. for each $x, y \in X$, we have that $M(x, y, t) \geq M(x, y, s)$ for each $t, s \in]0, \infty[$ with $t > s$.

This kind of fuzzy metric spaces has been studied by several authors from the mathematical point of view. Besides, they have been used successfully in engineering problems such as colour image filtering or perceptual colour difference (see [7]–[11]). Indeed, fuzzy metrics show some advantages with respect to the classical ones. On the one hand, the parameter t allows the fuzzy metric to be better adapted to context in which it is to be used. On the other hand, fuzzy

This work is partially supported by EU-H2020 projects BUGWRIGHT2 (GA 871260) and ROBINS (GA 779776), and by projects PGC2018-095709-B-C21 (MCIU/AEI/FEDER, UE), and PROCOE/4/2017 (Govern Balear, 50% P.O. FEDER 2014-2020 Illes Balears). This publication reflects only the authors views and the European Union is not liable for any use that may be made of the information contained therein.

metrics match perfectly with the employment of other fuzzy techniques, since the value given by a fuzzy metric, as pointed out before, can be directly interpreted as a fuzzy degree of nearness. So, providing useful techniques for generating fuzzy metrics becomes an interesting issue in order to provide a wider range of measurement tools in such a way that the fuzzy metric that best fits the problem being studied can be applied to solve it.

A celebrated example of fuzzy metric is the so-called standard fuzzy metric, which is defined from a classical metric (see [4]). Indeed, let (X, d) be a metric space and define the fuzzy set M_d on $X \times X \times]0, \infty[$ as follows:

$$M_d(x, y, t) = \frac{t}{t + d(x, y)}, \text{ for each } x, y \in X, t \in]0, \infty[. \quad (1)$$

The standard fuzzy metric on X deduced from d is the pair $(M_d, *_M)$, where $*_M$ denotes the minimum t -norm (i.e. $a *_M b = \min\{a, b\}$ for each $a, b \in [0, 1]$).

Observe that $(X, M_d, *)$ is also a fuzzy metric space for each continuous t -norm $*$, since $*_M$ is the largest t -norm. Indeed, given a continuous t -norm $*$, the inequality $a *_M b \geq a * b$ is satisfied for each $a, b \in [0, 1]$.

From the topological point of view, the standard fuzzy metric enjoys outstanding properties. The topologies generated from the standard fuzzy metric and from the classical metric, from which it is induced, coincide. Besides, it fulfils some interesting properties which do not make sense in the classical context but they do in the fuzzy context. Among others, it should be stressed the property of being strong (see [12]). Let us recall that a fuzzy metric space $(X, M, *)$ is said to be *strong* if, in addition, M satisfies, for each $x, y, z \in X$ and $t \in]0, \infty[$, the next inequality:

$$M(x, z, t) \geq M(x, y, t) * M(y, z, t). \quad (2)$$

Observe that the preceding inequality is stronger than that given in the axiom (GV4).

It is a well-known fact that, given a metric space (X, d) , then the standard fuzzy metric space $(X, M_d, *_P)$ is strong, where $*_P$ denotes the usual product t -norm, i.e. $a *_P b = a \cdot b$ for each $a, b \in [0, 1]$. Nevertheless, $(X, M_d, *_M)$ is not a strong fuzzy metric space in general, as pointed out in [12]. In view of the preceding fact, an interesting question arises: there exists a continuous t -norm $*$, different from $*_P$, with $* \geq *_P$ and such that $(X, M_d, *)$ is a strong fuzzy metric space for each metric space (X, d) ?

In [13], a generalization of the fuzzy set M_d given by (1) was introduced defining, for each $x, y \in X$ and $t \in]0, \infty[$, the next fuzzy set:

$$M_d^{g,m}(x, y, t) = \frac{g(t)}{g(t) + m \cdot d(x, y)}, \quad (3)$$

where $m \in]0, \infty[$ and $g :]0, \infty[\rightarrow]0, \infty[$ is a non-decreasing continuous function. According to [13], $(X, M_d^{g,m}, *_P)$ is a strong fuzzy metric space. Nevertheless, an extra condition on g is required so that $(X, M_d^{g,m}, *_M)$ is a fuzzy metric space for any arbitrary metric space (X, d) . Indeed, if the function g

is not superadditive, i.e. $g(t + s) \geq g(t) + g(s)$ for each $t, s \in]0, \infty[$, then $(X, M_d^{g,m}, *_M)$ is not, in general, a fuzzy metric space. Again, similar to the case of the standard fuzzy metric, it seems natural to wonder whether there exists a continuous t -norm, different from $*_P$, with $* \geq *_P$ such that $(X, M_d^{g,m}, *)$ is a fuzzy metric space for each metric space (X, d) without requiring any extra condition on g .

Coming back to the applicability of fuzzy metrics, in most problems, we are interested in measuring some kind of difference or similarity between objects. Therefore, fuzzy metrics can be good candidates to evaluate such a measurement. Concretely, the fuzzy set M is used to provide the aforementioned difference or similarity. However, the t -norm that defines M as a fuzzy metric does not play any role in the way in which such a measure is provided and, thus, it does not contribute anything that can make the fuzzy metric better fit for the problem under consideration. Since the fuzzy set $M_d^{g,m}$ given by expression (3) depends on more elements than the standard fuzzy metric M_d , $M_d^{g,m}$ allows to get more flexibility to obtain a measurement tool that fits better to the problem under consideration than M_d . So, providing a fuzzy set that generalizes expression (3) could improve the potential applicability of fuzzy metrics, even though such a generalization does not become a fuzzy metric for the same class of t -norms for which $M_d^{g,m}$ is so.

In the light of the exposed facts, the aim of this paper is twofold. On the one hand, we focus our efforts on obtaining a fuzzy set that generalizes expression (3) and on finding a family of continuous t -norms for which this new fuzzy set becomes a fuzzy metric. Moreover, we are interested in the study of those continuous t -norms for which this new fuzzy metric fulfils the property of being strong. Such a study allows us to approach the two questions posed above. On the other hand, we address a model estimation problem as an example of engineering application to illustrate the applicability of the new fuzzy metric proposed in Section II.

II. THE GENERALIZED STANDARD FUZZY METRIC

In this section, we build a new fuzzy metric which generalizes, in some sense, the standard fuzzy metric and the fuzzy metric given by expression (3). To this end, we recall a well-known family of continuous t -norms introduced by J. Dombi in [14].

Given $\lambda \in]0, \infty[$ the t -norm $*_{Dom}^\lambda$ is defined, for each $a, b \in [0, 1]$, by the following expression:

$$a *_{Dom}^\lambda b = \begin{cases} 0, & \text{if } a = 0 \text{ or } b = 0 \\ \frac{1}{1 + \left(\left(\frac{1-a}{a} \right)^\lambda + \left(\frac{1-b}{b} \right)^\lambda \right)^{\frac{1}{\lambda}}}, & \text{otherwise} \end{cases}. \quad (4)$$

The construction of the promised fuzzy metric can be found in the next result:

Theorem 2.1: Let (X, d) be a metric space, $m, n \in]0, \infty[$ and $g :]0, \infty[\rightarrow]0, \infty[$ be a non-decreasing continuous function. Define the fuzzy set $\tilde{M}_d^{g,m,n}$ on $X \times X \times]0, \infty[$ as:

$$\tilde{M}_d^{g,m,n}(x, y, t) = \frac{g(t)}{g(t) + m \cdot d^n(x, y)}, \quad (5)$$

where $d^n(x, y)$ denotes $(d(x, y))^n$. Then, $(X, \tilde{M}_d^{g,m,n}, *)$ is a fuzzy metric space for each continuous t -norm $*$ satisfying $*$ $\leq *_{Dom}^{\frac{1}{n}}$.

Proof.

Let $*$ be a continuous t -norm such that $*$ $\leq *_{Dom}^{\frac{1}{n}}$. It is not hard to check that $\tilde{M}_d^{g,m,n}$ satisfies axioms (GV1), (GV2), (GV3) and (GV5). It remains to prove that (GV4) also holds.

Let $x, y, z \in X$ and $t, s \in]0, \infty[$. We will see that

$$\tilde{M}_d^{g,m,n}(x, z, t+s) \geq \tilde{M}_d^{g,m,n}(x, y, t) * \tilde{M}_d^{g,m,n}(y, z, s).$$

Set $\alpha = \max\{g(t), g(s)\}$. Observe that

$$\tilde{M}_d^{g,m,n}(x, z, t+s) \geq \frac{\alpha}{\alpha + m \cdot d^n(x, z)},$$

$$\tilde{M}_d^{g,m,n}(x, y, t) \leq \frac{\alpha}{\alpha + m \cdot d^n(x, y)}$$

and

$$\tilde{M}_d^{g,m,n}(y, z, s) \leq \frac{\alpha}{\alpha + m \cdot d^n(y, z)}.$$

So, since $*$ $\leq *_{Dom}^{\frac{1}{n}}$, we have that

$$\begin{aligned} & \tilde{M}_d^{g,m,n}(x, y, t) * \tilde{M}_d^{g,m,n}(y, z, s) \leq \\ & \leq \tilde{M}_d^{g,m,n}(x, y, t) *_{Dom}^{\frac{1}{n}} \tilde{M}_d^{g,m,n}(y, z, s) \leq \\ & \leq \frac{\alpha}{\alpha + m \cdot d^n(x, y)} *_{Dom}^{\frac{1}{n}} \frac{\alpha}{\alpha + m \cdot d^n(y, z)} = \\ & = \frac{1}{1 + \frac{m \cdot (d(x, y) + d(y, z))^n}{\alpha}} = \frac{\alpha}{\alpha + m \cdot (d(x, y) + d(y, z))^n} \leq \\ & \leq \frac{\alpha}{\alpha + m \cdot d^n(x, z)} \leq \tilde{M}_d^{g,m,n}(x, z, t+s). \end{aligned}$$

Therefore, for each $x, y, z \in X$ and $t, s \in]0, \infty[$, $\tilde{M}_d^{g,m,n}$ satisfies (GV4) for $*$ and we conclude that $(X, \tilde{M}_d^{g,m,n}, *)$ is a fuzzy metric space. $\square$

It must be stressed that $(X, \tilde{M}_d^{g,m,n}, *)$ is not a fuzzy metric space, in general, when $*$ does not satisfy the condition $*$ $\leq *_{Dom}^{\frac{1}{n}}$, as the next example shows.

Example 2.2: Let $(\mathbb{R}, d_u)$ be the metric space where d_u is the Euclidean metric on $\mathbb{R}$, i.e. $d_u(x, y) = |x - y|$. Consider the non-decreasing continuous function $g_1 :]0, \infty[\rightarrow]0, \infty[$ given by $g(t) = 1$, for each $t \in]0, \infty[$, and $m = n = 1$. Then, the fuzzy set $\tilde{M}_{d_u}^{g_1,1,1}$ is given by expression (5) as follows:

$$\tilde{M}_{d_u}^{g_1,1,1}(x, y, t) = \frac{1}{1 + d_u(x, y)}, \text{ for each } x, y \in \mathbb{R}, t \in]0, \infty[.$$

Let $*$ be a continuous t -norm such that $*$ $\not\leq *_{Dom}^{\frac{1}{n}}$. Then, there exists $a, b \in]0, 1[$ such that $a * b > a *_{Dom}^{\frac{1}{n}} b$.

Consider $x = \frac{a-1}{a}$, $y = 0$, $z = \frac{1-b}{b}$ and $t, s \in]0, \infty[$. Then,

$$\tilde{M}_{d_u}^{g_1,1,1}(x, z, t+s) = \frac{1}{1 + \frac{1-a}{a} + \frac{1-b}{b}} = a *_{Dom}^{\frac{1}{n}} b,$$

$$\tilde{M}_{d_u}^{g_1,1,1}(x, y, t) = a \text{ and } \tilde{M}_{d_u}^{g_1,1,1}(y, z, s) = b.$$

Therefore, $\tilde{M}_{d_u}^{g_1,1,1}(x, y, t) * \tilde{M}_{d_u}^{g_1,1,1}(y, z, s) = a * b > a *_{Dom}^{\frac{1}{n}} b = \tilde{M}_{d_u}^{g_1,1,1}(x, z, t+s)$, and so $\tilde{M}_{d_u}^{g_1,1,1}$ does not satisfy (GV4).

On account of Theorem 2.1 and the preceding example, we conclude that $*_{Dom}^{\frac{1}{n}}$ is the largest (continuous) t -norm for which $\tilde{M}_d^{g,m,n}$ is a fuzzy metric on X , for each arbitrary metric space (X, d) , each non-decreasing continuous function $g :]0, \infty[\rightarrow]0, \infty[$ and each $m, n \in]0, \infty[$. Such a conclusion allows us to approach the two questions posed in Section I.

On the one hand, Theorem 2.1 introduces a generalization of the fuzzy set given by expression (3). Indeed, such a fuzzy set is obtained by considering $n = 1$ in the fuzzy set defined by expression (5), i.e. $M_d^{g,m} = \tilde{M}_d^{g,m,1}$. Besides, the aforementioned theorem establishes that $M_d^{g,m}$ is a fuzzy metric on X for each t -norm $*$ with $*$ $\leq *_{Dom}^{\frac{1}{n}}$. This fact allows us to answer in affirmative way one of the questions that we wondered in Section I, which is whether there exists a continuous t -norm $*$ $\geq *_P$ such that $(X, M_d, *)$ is a fuzzy metric space for each metric space (X, d) without requiring any extra condition on g .

First, observe that, for each $a, b \in]0, 1[$, we have that

$$a *_{Dom}^{\frac{1}{n}} b = \frac{1}{1 + \frac{1-a}{a} + \frac{1-b}{b}} = \frac{ab}{a + b - ab}.$$

Now,

$$\frac{ab}{a + b - ab} \geq a \cdot b \Leftrightarrow 1 \geq a + b - ab = a(1 - b) + b.$$

Taking into account that $a \leq 1$ we have that $1 = 1 - b + b \geq a(1 - b) + b$. So, $a *_{Dom}^{\frac{1}{n}} b \geq a *_P b$ for each $a, b \in]0, 1[$. Thus $*_{Dom}^{\frac{1}{n}} \geq *_P$ for each $a, b \in [0, 1]$, since $a *_{Dom}^{\frac{1}{n}} b = 0 = a *_P b$ whenever $a = 0$ or $b = 0$.

Hence, we have found a continuous Archimedean t -norm greater than the product t -norm $*_P$ for which $M_d^{g,m}$ is a fuzzy metric on X , for each arbitrary metric space (X, d) , each non-decreasing continuous function $g :]0, \infty[\rightarrow]0, \infty[$ and each $m \in]0, \infty[$. Furthermore, on account of Example 2.2 we conclude that $*_{Dom}^{\frac{1}{n}}$ is the largest t -norm for which $(X, M_d^{g,1}, *_{Dom}^{\frac{1}{n}})$ is a fuzzy metric space, in general.

On the other hand, $\tilde{M}_d^{g,m,n}$ becomes the standard fuzzy metric when we consider the non-decreasing continuous function $g(t) = t$ and $m = n = 1$. Under this remark and, based on the argument exposed in the proof of Theorem 2.1, we prove the next result which will be useful to answer the first question about the standard fuzzy metric set out in Section I.

Theorem 2.3: Let (X, d) be a metric space, $m, n \in]0, \infty[$ and $g :]0, \infty[\rightarrow]0, \infty[$ be a non-decreasing continuous function. Then the fuzzy metric space $(X, \tilde{M}_d^{g,m,n}, *)$, where $\tilde{M}_d^{g,m,n}$ is given by (5), is strong for each continuous t -norm satisfying $*$ $\leq *_{Dom}^{\frac{1}{n}}$.

Proof.

Consider a continuous t -norm $*$ such that $*$ $\leq *_{Dom}^{\frac{1}{n}}$. By Theorem 2.1 we conclude that $(X, \tilde{M}_d^{g,m,n}, *)$ is a fuzzy metric space. It remains to show that inequality (2) holds.

Let $x, y, z \in X$ and $t, s \in]0, \infty[$. Then,

$$\begin{aligned} \tilde{M}_d^{g,m,n}(x, z, t) &= \frac{g(t)}{g(t) + m \cdot d^n(x, z)} \geq \\ &\geq \frac{g(t)}{g(t) + m \cdot (d(x, y) + d(y, z))^n} = \\ &= \tilde{M}_d^{g,m,n}(x, y, t) *_{Dom}^1 \tilde{M}_d^{g,m,n}(y, z, t) \geq \\ &\tilde{M}_d^{g,m,n}(x, y, t) * \tilde{M}_d^{g,m,n}(y, z, t). \end{aligned}$$

Hence, $(X, \tilde{M}_d^{g,m,n}, *)$ is a strong fuzzy metric space. $\square$

As a consequence of the previous theorem, we conclude that $(X, M_d, *_{Dom}^1)$ is a strong fuzzy metric space. This fact answers affirmatively the first question lay out in Section I by providing a continuous t -norm greater than the product t -norm $*_P$ for which the standard fuzzy metric is strong for any arbitrary metric space (X, d) .

Even more, on account of Example 2.2 we conclude that the standard fuzzy metric is just a strong fuzzy metric in general for continuous t -norms less than $*_{Dom}^1$. Notice that the aforementioned example provides that the standard fuzzy metric $(\mathbb{R}, M_{d_u}, *)$ is not strong if $* \not\leq *_{Dom}^1$. Indeed, let $a, b \in]0, 1[$ such that $a * b > a *_{Dom}^1 b$. Then, take $x = \frac{a-1}{a}$, $y = 0$ and $z = \frac{1-b}{b}$. Therefore,

$$\begin{aligned} M_{d_u}(x, z, 1) &= \frac{1}{1 + \frac{1-a}{a} + \frac{1-b}{b}} = a *_{Dom}^1 b < a * b = \\ &= M_{d_u}(x, y, 1) * M_{d_u}(y, z, 1). \end{aligned}$$

III. APPLICATION CASE: ROBUST MODEL ESTIMATION

Solving model estimation problems is a fundamental component of numerous applications involving perception tasks. Nowadays, facing this kind of problem requires to cope with new challenges due to an increased use of poor, low-cost sensors, and the ever growing deployment of robotic devices which may operate in potentially unknown environments. Generally speaking, the underlying algorithms have to be robust against uncertain data that besides may be corrupted by outliers, i.e. data items which are not consistent with the original model due to an arbitrary bias affecting them. A *robust estimator* is able to correctly find the original model that supposedly the input data fits to under the aforementioned conditions [15]. The Random Sample Consensus algorithm (RANSAC) [16] is one of these robust estimation techniques, which is widely used nowadays, so much that it has become common in robotics and computer vision.

Briefly speaking, RANSAC tries to achieve a maximum consensus in the input dataset in order to deduce the inliers by generating random hypotheses on the model parameters through a *hypothesize-and-verify* approach. That is to say, instead of using every sample in the dataset to perform the estimation as in traditional regression techniques, RANSAC tests many random sets of samples and outputs the one leading to the best fitting. Since picking an extra point decreases

exponentially the probability of selecting an outlier-free sample [17], RANSAC takes the Minimum Sample Set size (MSS) to determine a unique candidate model, thus increasing its chances of finding an all-inlier sample set. This model is assigned a score based on the cardinality of its consensus set. Finally, RANSAC returns the hypothesis that has achieved the highest consensus and the set of inliers, which are used next to estimate the ultimate model by regression.

Searching for an all-inlier sample, RANSAC typically runs for N iterations:

$$N = \frac{\log(1 - \rho)}{\log(1 - (1 - \omega)^s)} \quad (6)$$

where ρ is the desired probability of success, i.e. at least one of the considered random sets is outlier-free, s is the size of the MSS for the problem at hand and ω is the ratio of outliers. See [16] for the details on Eq. (6).

Algorithm 1 outlines FM-based RANSAC, a variant of RANSAC described in [18] that avoids discriminating between inliers and outliers by means of the use of a fuzzy metric that encodes as a similarity the compatibility of each sample to the currently hypothesized model. In this work, we particularize FM-based RANSAC for the fuzzy metric $\tilde{M}_d^{g,m,n}$ introduced as Eq. (5) in Section II. $\tilde{M}_d^{g,m,n}$ is also incorporated into the final model refinement step that follows the main hypothesis selection loop. Finally, in Section IV, we report on the accuracy achieved by FM-based RANSAC for a specific model estimation problem when using $\tilde{M}_d^{g,m,n}$ for different values of m and n . The assessment involves a comparison with RANSAC and MSAC [19] for a benchmark comprising a large number of datasets with varying proportion of outliers and different levels of noise.

We detail next the features of FM-based RANSAC:

- 1) **Samples classification.** In the original RANSAC, for every model considered, data samples are classified into inliers and outliers by comparing the fitting error with a threshold τ_I related to data noise. As already mentioned, FM-based RANSAC does not distinguish between inliers and outliers, but makes use of a compatibility value $\phi \in [0, 1]$ between each sample x_j and the current model $\mathcal{M}_{\hat{\Theta}_k}$, given the fitting error $\epsilon(x_j; \mathcal{M}_{\hat{\Theta}_k})$. Such compatibility value derives from the fuzzy metric $\tilde{M}_d^{g,m,n}$ once parameterized by (d, Φ) with $\Phi = (n, m, g)$. Since in the following we contemplate the use of an only, specific distance d , i.e. the Euclidean metric, and g is set to the constant function θ^n as a reference of noise scale¹, we denote the fuzzy metric as $M^{m,n}$ eliminating the allusion to d and g . From now on, the value of $M^{m,n}$ will be denoted by $\phi(\epsilon; \Phi)$.
- 2) **Model scoring.** The individual compatibility values $\phi(\epsilon; \Phi)$ are aggregated by simple summation to obtain the model score (step 6 in Alg. 1) and hence the *so-far-the-best-model* is given by the maximum score found up to the current iteration (steps 7 - 9 of Alg. 1).

¹In this regard, we refer to the form $M_d^{g,m,n}(x, y) = \frac{1}{1 + m \cdot (d(x, y)/\theta)^n}$.

Algorithm 1 FM-based RANSAC

Input: D - dataset comprising samples $\{x_j\}$
 $\phi(\epsilon; \Phi)$ - FM compatibility value for fitting error
 $k_{\max}$ - maximum number of iterations of the main loop
 $t_{\max}$ - maximum number of iterations of the refinement stage

Output: $\mathcal{M}_{\hat{\Theta}}$ - estimated model

```

1:  $k := 0, \varphi_{\max} := -\infty$ 
2: for  $k := 1$  to  $k_{\max}$  do  $\triangleright$  find best consensus model  $\mathcal{M}_{\hat{\Theta}_k}$ 
3:   select randomly a minimal sample set  $S_k$  of size  $s$ 
4:   estimate model  $\mathcal{M}_{\hat{\Theta}_k}$  from  $S_k$ 
5:   calculate fitting errors  $\epsilon(x_j; \mathcal{M}_{\hat{\Theta}_k}), \forall x_j \in D$ 
6:   find model score  $\varphi_k := \sum_{x_j \in D} \phi(\epsilon(x_j; \mathcal{M}_{\hat{\Theta}_k}); \Phi)$ 
7:   if  $\varphi_k > \varphi_{\max}$  then
8:      $\varphi_{\max} := \varphi_k, \mathcal{M}_{\hat{\Theta}}^0 := \mathcal{M}_{\hat{\Theta}_k}$ 
9:   end if
10: end for
11:  $t := 0$ 
12: repeat  $\triangleright$  refine model  $\mathcal{M}_{\hat{\Theta}}$ 
13:   calculate fitting errors  $\epsilon(x_j; \mathcal{M}_{\hat{\Theta}}^t), \forall x_j \in D$ 
14:   estimate model  $\mathcal{M}_{\hat{\Theta}}^{t+1}$  using weights  $\phi(\epsilon(x_j; \mathcal{M}_{\hat{\Theta}}^t); \Phi)$ 
15:    $t := t + 1$ 
16: until convergence or  $t \geq t_{\max}$ 
17: return  $\mathcal{M}_{\hat{\Theta}}^t$ 

```

- 3) **Model refinement.** Once a sufficient number of models have been considered, we re-estimate the winning model using iterative weighted least squares, where the compatibility values $\phi(\epsilon; \Phi)$, calculated for the fitting errors resulting from the current model, are used as weights for the new, refined model (steps 12 - 16 of Alg. 1). The loop iterates until changes in the estimated parameters of the model $\hat{\Theta}$ are negligible (or after $t_{\max}$ iterations).

IV. EXPERIMENTAL RESULTS

A. Experimental setup

For testing purposes, we consider a hyperplane model estimation problem for 2D (straight lines), 3D (planes) and 10D, the latter as a case of higher dimensionality. To this end, we generate synthetic datasets stemming from hyperplanes in random orientations and positions: 500 for 2D/3D hyperplanes and 250 for 10D hyperplanes. Given a 2D/3D/10D random point p belonging to a hyperplane with normal vector $\vec{n}$, an inlier p_I is generated by shifting p along $\vec{n}$ using a zero-mean Gaussian distribution with standard deviation σ , i.e. $p_I = p + \mathcal{N}(0, \sigma) \cdot \vec{n}$. Outliers p_O are uniformly generated within a rectangular area containing part of the hyperplane, ensuring that they lie out of a 3σ stripe at both sides of the hyperplane. Every pair (σ, ω) gives rise to a different dataset.

Regarding hypothesis generation within the main loop, in all experiments, the size of the MSS is always set to the minimum, i.e. $s = 2$, $s = 3$ and $s = 10$ for respectively 2D, 3D and 10D. Besides, the number of iterations $k_{\max}$ is calculated according

to Eq. (6), with $\rho = 99\%$. The parameters of $\phi(\epsilon; \Phi)$, $\Phi = (n, m, g)$, are set as follows: $n, m \in \{1, 2\}$, as indicated for each experiment; and g is the constant function θ^n , where $\theta = \kappa \cdot \sigma$. For RANSAC/MSAC $\tau_I = \kappa \cdot \sigma$. Different values for κ are considered for both θ and τ_I . Finally, to compare properly RANSAC, MSAC and our estimator, we make use of the same sequence of MSS's to avoid the effect of randomness.

B. Results and discussion

In the following, to measure the estimation accuracy for the hyperplane fitting problem, we make use of the average $\mu[\epsilon]$ of the angle ϵ between the true and the estimated normal vector. We also report on the average number of iterations spent during model refinement $\mu[t]$.

Table I and Fig. 1 show performance results for the fuzzy metric $M^{m,n}$ and several outlier ratios ω and Gaussian noise magnitudes σ . In sight of these results, it is worth noting that: (1) the estimation accuracy for $M^{2,2}$ is above that of plain RANSAC and MSAC in all cases, while, for the other configurations of $M^{m,n}$, the accuracy is in general better than the classical counterparts though not in all cases; (2) the value of θ in $M^{m,n}$ does not seem to be critical, since the highest change in $\mu[\epsilon]$ for θ with $\kappa \in \{1, 2, 2.5, 3, 4\}$ is less than 1° ; (3) the distribution of the average error $\mu[\epsilon]$ shows always larger errors for RANSAC/MSAC than for $M^{2,2}$ for all percentiles. As for the number of iterations of the refinement stage t : (4) in general, $\mu[t]$ is similar for every combination of $M^{m,n}$ when varying the noise parameters (σ, ω) and particularly higher for $M^{2,2}$ when κ is low; (5) lower values of κ allow the proposed estimator to perform a better refinement stage in terms of accuracy but at the expense of computational cost since more iterations are required. Regarding the fuzzy metric $M^{m,n}$, both $M^{2,n}$ and $M^{m,2}$ lead in general to higher accuracy, with a slight increase in the number of refining iterations for low κ values or higher noise (σ, ω) .

V. CONCLUSIONS

This work introduces a methodology to induce fuzzy metrics that generalizes the celebrated standard fuzzy metric. Moreover, some questions related to the continuous triangular norms from which such fuzzy metrics are defined have been posed and answered. A concrete new fuzzy metric induced through the aforementioned methodology has been successfully embedded within a revised version of RANSAC. By means of this metric, we avoid discriminating between inliers and outliers, to instead make use of a compatibility value to the current model for each sample. Experimental results show good performance, actually outperforming two classical counterparts, RANSAC and MSAC.

REFERENCES

- [1] L. Zadeh, "Fuzzy sets," *Information and Control*, vol. 8, no. 3, pp. 338–353, 1965.
- [2] I. Kramosil and J. Michalek, "Fuzzy Metrics and Statistical Metric Spaces," *Kybernetika*, vol. 11, no. 5, pp. 336–334, 1975.
- [3] K. Menger, "Statistical metrics," *Proceedings of the National Academy of Sciences of the United States of America*, vol. 28, no. 12, pp. 535–537, 1942.

TABLE I

2D/3D/10D HYPERPLANES ESTIMATION ACCURACY AND NUMBER OF ITERATIONS OF THE REFINEMENT STAGE FOR (TOP) DIFFERENT OUTLIER RATIOS ω , (MIDDLE) DIFFERENT NOISE MAGNITUDES σ AND (BOTTOM) DIFFERENT SETTINGS FOR $\tau_I, \theta = \kappa \cdot \sigma$. WHEN KEPT CONSTANT: $\sigma = 1$, $\omega = 0.4$, $\kappa = 3$. LIGHTER BACKGROUND MEANS HIGHER PERFORMANCE.

2D		$\mu[e] (^{\circ})$					$\mu[t]$				
ω	RANSAC	MSAC	FM-based RANSAC				FM-based RANSAC				
			$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	
0.60	4.43	3.14	4.63	4.10	3.97	3.10	0.60	9.11	10.44	10.79	11.05
0.50	3.03	2.33	2.61	2.33	2.28	1.88	0.50	7.28	7.93	8.57	8.58
0.40	2.13	1.81	1.77	1.59	1.57	1.33	0.40	6.47	6.71	7.68	7.40
0.20	1.58	1.53	0.96	0.88	0.89	0.80	0.20	5.68	5.51	6.82	6.32
σ	RANSAC	MSAC	FM-based RANSAC				FM-based RANSAC				
			$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	
2.00	9.82	6.92	4.08	4.62	3.87	4.46	2.00	7.40	8.38	9.00	9.64
1.00	2.13	1.81	1.77	1.59	1.57	1.33	1.00	6.47	6.71	7.68	7.40
0.50	0.74	0.71	0.90	0.55	0.73	0.45	0.50	6.20	5.63	7.21	6.29
0.25	0.37	0.36	0.48	0.20	0.36	0.17	0.25	6.02	4.93	6.87	5.57
κ	RANSAC	MSAC	FM-based RANSAC				FM-based RANSAC				
			$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	
4.00	2.85	2.09	1.85	1.84	1.65	1.54	4.00	6.05	6.26	7.14	6.81
3.00	2.13	1.81	1.77	1.59	1.57	1.33	3.00	6.47	6.71	7.68	7.40
2.50	2.03	1.88	1.71	1.45	1.52	1.23	2.50	6.79	7.02	8.03	7.94
2.00	2.18	2.18	1.65	1.29	1.47	1.13	2.00	7.14	7.54	8.61	8.78
1.00	3.60	3.58	1.47	1.04	1.35	1.01	1.00	8.61	10.65	10.66	13.56
3D		$\mu[e] (^{\circ})$					$\mu[t]$				
ω	RANSAC	MSAC	FM-based RANSAC				FM-based RANSAC				
			$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	
0.60	6.14	4.58	5.11	4.00	4.24	3.04	0.60	10.02	10.87	11.69	11.47
0.50	4.07	3.48	2.90	2.34	2.46	1.87	0.50	7.55	8.10	9.05	8.83
0.40	3.13	3.01	1.95	1.61	1.69	1.35	0.40	6.69	6.91	8.01	7.67
0.20	2.31	2.29	1.07	0.92	0.98	0.84	0.20	5.73	5.55	6.87	6.41
σ	RANSAC	MSAC	FM-based RANSAC				FM-based RANSAC				
			$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	
2.00	13.32	9.97	5.21	4.92	4.53	4.32	2.00	8.68	9.63	10.45	10.66
1.00	3.13	3.01	1.95	1.61	1.69	1.35	1.00	6.69	6.91	8.01	7.67
0.50	1.11	1.08	0.99	0.57	0.79	0.47	0.50	6.32	5.82	7.43	6.52
0.25	0.64	0.63	0.52	0.21	0.39	0.18	0.25	6.15	5.19	7.14	5.91
κ	RANSAC	MSAC	FM-based RANSAC				FM-based RANSAC				
			$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	
4.00	3.71	2.95	2.07	1.87	1.79	1.56	4.00	6.23	6.47	7.40	7.04
3.00	3.13	3.01	1.95	1.61	1.69	1.35	3.00	6.69	6.91	8.01	7.67
2.50	3.23	3.22	1.88	1.46	1.63	1.25	2.50	7.00	7.28	8.45	8.25
2.00	3.75	3.75	1.79	1.31	1.57	1.15	2.00	7.40	7.87	9.04	9.17
1.00	5.73	5.83	1.57	1.07	1.43	1.06	1.00	9.04	11.25	11.33	14.38
10D		$\mu[e] (^{\circ})$					$\mu[t]$				
ω	RANSAC	MSAC	FM-based RANSAC				FM-based RANSAC				
			$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	
0.60	10.69	10.18	7.46	4.19	5.53	3.00	0.60	14.12	14.53	15.61	14.52
0.50	8.82	8.77	3.56	2.26	2.83	1.78	0.50	8.85	9.50	10.67	10.34
0.40	6.92	6.94	2.29	1.54	1.88	1.29	0.40	7.30	7.60	8.88	8.50
0.20	5.66	5.70	1.17	0.90	1.03	0.84	0.20	6.02	5.98	7.28	6.88
σ	RANSAC	MSAC	FM-based RANSAC				FM-based RANSAC				
			$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	
2.00	26.76	20.83	9.83	7.47	7.49	5.52	2.00	15.13	16.20	16.62	16.38
1.00	6.92	6.94	2.29	1.54	1.88	1.29	1.00	7.30	7.60	8.88	8.50
0.50	3.03	3.04	1.08	0.55	0.84	0.48	0.50	6.74	6.38	8.00	7.22
0.25	1.30	1.31	0.56	0.22	0.42	0.20	0.25	6.45	5.81	7.56	6.63
κ	RANSAC	MSAC	FM-based RANSAC				FM-based RANSAC				
			$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	$M^{1,1}$	$M^{1,2}$	$M^{2,1}$	$M^{2,2}$	
4.00	6.26	6.09	2.50	1.85	2.03	1.49	4.00	6.96	7.16	8.14	7.74
3.00	6.92	6.94	2.29	1.54	1.88	1.29	3.00	7.30	7.60	8.88	8.50
2.50	7.75	7.86	2.16	1.39	1.80	1.20	2.50	7.78	8.03	9.28	9.11
2.00	9.03	9.13	2.03	1.26	1.71	1.14	2.00	8.14	8.79	10.06	10.06
1.00	12.61	12.61	1.71	1.10	1.52	1.12	1.00	10.06	12.41	12.58	15.87

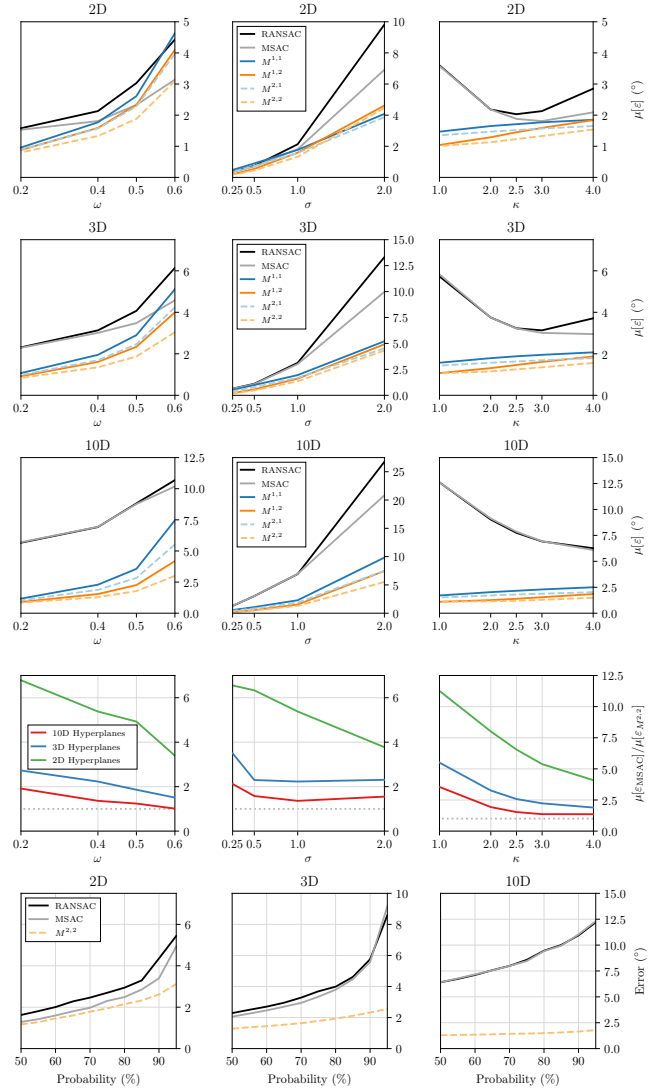


Fig. 1. (1st, 2nd, 3rd rows) Estimation accuracy for resp. 2D, 3D and 10D hyperplanes varying (left) the outlier ratio ω , (middle) the noise magnitude σ and (right) the setting of $\tau_I, \theta = \kappa \cdot \sigma$. (4th row) Ratio between the error of MSAC $\mu[\epsilon_{\text{MSAC}}]$ and the error of FM-based RANSAC for $M^{2,2}$ $\mu[\epsilon_{M^{2,2}}]$. (5th row) Comparison between FM-based RANSAC for $M^{2,2}$ and RANSAC/MSAC for different percentiles of the estimation error.

[4] A. George and P. Veeramani, "On some results in fuzzy metric spaces," *Fuzzy Sets and Systems*, vol. 64, no. 3, pp. 395–399, 1994.

[5] E. Klement, R. Mesiar, and E. Pap, *Triangular norms*. Springer, 2000.

[6] M. Grabiec, "Fixed points in fuzzy metric spaces," *Fuzzy Sets and Systems*, vol. 27, pp. 385–389, 1989.

[7] J. Camarena, V. Gregori, S. Morillas, and A. Sapena, "Fast detection and removal of impulsive noise using peer groups and fuzzy metrics," *Journal of Visual Communication and Image Representation*, vol. 19, no. 1, pp. 20–29, 2008.

[8] S. Morillas, V. Gregori, and G. Peris-Fajarnés, "New adaptive vector filter using fuzzy metrics," *Journal of Electronic Imaging*, vol. 16, no. 3, pp. 033 007:1–15, 2007.

[9] S. Morillas, V. Gregori, G. Peris-Fajarnés, and P. Latorre, "Isolating impulsive noise color images by peer group techniques," *Computer Vision and Image Understanding*, vol. 110, no. 1, pp. 102–116, 2008.

[10] S. Morillas, V. Gregori, G. Peris-Fajarnés, and A. Sapena, "Local self-adaptive fuzzy filter for impulsive noise removal in color image," *Signal Processing*, vol. 8, no. 2, pp. 390–398, 2008.

[11] V. Gregori, J. Miñana, and S. Morillas, "Some questions in fuzzy metric spaces," *Fuzzy Sets and Systems*, vol. 204, pp. 71–85, 2012.

[12] V. Gregori, S. Morillas, and A. Sapena, "On a class of completable fuzzy metric spaces," *Fuzzy Sets and Systems*, vol. 161, pp. 2193–2205, 2010.

[13] —, "Examples of fuzzy metrics and applications," *Fuzzy Sets and Systems*, vol. 170, pp. 95–111, 2011.

[14] J. Dombi, "A general class of fuzzy operators, the de morgan class of fuzzy operators and fuzziness measures induced by fuzzy operators," *Fuzzy Sets and Systems*, vol. 8, pp. 149–163, 1982.

[15] P. J. Huber and E. M. Ronchetti, *Robust statistics*. Wiley, 2011.

[16] M. A. Fischler and R. C. Bolles, "Random sample consensus," *Communications of the ACM*, vol. 24, no. 6, pp. 381–395, 1981.

[17] O. Chum, J. Matas, and J. Kittler, "Locally Optimized RANSAC," in *Pattern Recognition (Proc. DAGM)*, ser. LNCS vol. 2781, B. Michaelis and G. Krell, Eds. Springer, 2003, pp. 236–243.

[18] A. Ortiz, E. Ortiz, J. J. Miñana, and O. Valero, "Hypothesis Scoring and Model Refinement Strategies for FM-based RANSAC," in *Proc. Spanish Conference on Fuzzy Logic and Technologies*, ser. Lecture Notes in Artificial Intelligence. Springer, 2021, in press.

[19] P. Torr and A. Zisserman, "MLESAC: A New Robust Estimator with Application to Estimating Image Geometry," *Computer Vision and Image Understanding*, vol. 78, no. 1, pp. 138–156, 2000.

Uso de t-normas para el estudio de la convexidad en conjuntos difusos intervalo-valuados

1st Pedro Huidobro
Dept. Statistics, O.R & M.E.
University of Oviedo
Oviedo, Spain
huidobropedro@uniovi.es

2nd Pedro Alonso
Dept. Mathematics
University of Oviedo
Oviedo, Spain
palonso@uniovi.es

3rd Humberto Bustince
Dept. of Statistics, Computing and Mathematics
University of Navarra
Navarra, Spain
bustince@unavarra.es

4th Vladimír Janiš
Dept. Mathematics
Matej Bel University
Banská Bystrica, Slovakia
vladimir.janis@umb.sk

5th Susana Montes
Dept. Statistics, O.R & M.E.
University of Oviedo
Oviedo, Spain
montes@uniovi.es

Resumen—En muchos problemas reales no se pueden tomar medidas de forma exacta. Así, los conjuntos difusos surgieron como una forma de intentar tratar con la incertidumbre de la forma más eficiente posible. Por otro lado, debe señalarse que la convexidad es un concepto interesante en varias áreas dentro de las matemáticas. Teniendo esto en cuenta, en este documento proponemos una extensión del concepto de convexidad para conjuntos difusos intervalo-valuados basada en el uso de t-normas para intervalos. Para ello, y teniendo en consideración la literatura científica existente respecto de t-normas, presentamos una definición de t-norma aplicada a intervalos. Por último, comprobamos que nuestra definición de convexidad, utilizando t-normas, preserva la convexidad a través de intersecciones, es decir, que la intersección de dos conjuntos difusos intervalo-valuados convexos es también convexa.

Index Terms—Conjuntos difusos intervalo-valuados, t-normas, convexidad

I. INTRODUCCIÓN

En la vida real no todas las mediciones son exactas, por ello, en 1965, Zadeh [29] introdujo los conjuntos difusos para lidiar con la imprecisión. Desde entonces, muchos autores los han estudiado, así como a sus extensiones. Una de las más conocidas es la que da lugar a los conjuntos difusos intervalo-valuados, que fueron en la década de 1970 presentados independientemente por Zadeh [30], Grattan-Guinness [13], Jahn [16] y Sambuc [22]. Hoy en día, los estudios sobre este tipo de conjuntos han aumentado gracias a su aplicabilidad (ver [5], [7], [12], [22], [26]).

Por otro lado, la convexidad es una herramienta muy útil en muchos campos de las matemáticas (ver por ejemplo [17]–[19], [25]).

Este estudio ha sido parcialmente patrocinado por el programa español MI-NECO (TIN-2017-87600-P: P. Alonso; PGC2018-098623-B-I00: P. Huidobro and S. Montes), MICIN (PID2019-108392GB-I00: H. Bustince), la ayuda no. 1/0150/21 proporcionada por la agencia de subvenciones eslovaca VEGA (V. Janiš) y el programa de ayudas Severo Ochoa PA-20PF-BP19-169 (P. Huidobro).

Cuando Zadeh presentó los conjuntos difusos, estudió también su convexidad. Desde entonces, numerosos autores han trabajado sobre este tema, así como sobre la convexidad de sus extensiones, por ejemplo Ammar y Metz [1], Diaz et al. [8], Gupta y Dabgar [14], Ramik y Vlach [21], Sarkar [23], Syau y Lee [24], Yang [28] o Zhang et al. [31].

Según nuestro conocimiento, la convexidad en conjuntos difusos intervalo-valuados no ha sido estudiada en profundidad, por lo que Huidobro et al. [15] han trabajado recientemente en esa dirección.

Este trabajo está organizado del siguiente modo. Primero, en la sección 2, haremos una presentación de los conceptos necesarios para los desarrollos llevados a cabo. A continuación, en la sección 3 estudiaremos las t-normas para intervalos cerrados contenidos en el $[0, 1]$. Finalmente, analizaremos si la convexidad de conjuntos difusos intervalo-valuados se preserva por intersecciones.

II. CONCEPTOS BÁSICOS

Denotaremos por X al conjunto de referencia. Un conjunto difuso μ_A puede caracterizarse por una función de pertenencia $\mu_A : X \rightarrow [0, 1]$. La colección de todos los conjuntos difusos sobre X se representa como $FS(X)$. Un conjunto difuso intervalo-valuado está caracterizado por una función $A : X \rightarrow L([0, 1])$, donde $A(x) = [\underline{A}(x), \overline{A}(x)]$ y $\underline{A}(x) \leq \overline{A}(x)$ para todo $x \in X$, y $L([0, 1])$ son todos los intervalos cerrados contenidos en el $[0, 1]$. Denotaremos como $IVFS(X)$ a la colección de todos los conjuntos difusos intervalo-valuados en X . Podemos observar que los conjuntos difusos intervalo-valuados son una generalización de los conjuntos difusos, donde la función de pertenencia toma como valor intervalos.

El primer problema que encontramos es cómo ordenar los elementos. Mientras que en $\mathbb{R}$ hay un orden natural ($a \leq b$ si $b - a$ no es un número negativo), en $L([0, 1])$ no ocurre lo mismo. Algo que parece lógico, a la hora de considerar órdenes, es respetar el orden reticular (*lattice order* en inglés)

[11], definido como $a \preceq_{Lo} b$ if $\underline{a} \leq \underline{b}$ y $\bar{a} \leq \bar{b}$ para cualquier $a = [\underline{a}, \bar{a}]$ y $b = [\underline{b}, \bar{b}]$ en $L([0, 1])$. En este trabajo seguiremos las ideas de Bustince et al. [6] usando órdenes admisibles, que son órdenes totales refinando el orden reticular. Para un estudio más detallado del uso de órdenes admisibles entre intervalos ver [15].

Definición 2.1: [6] Sea $(L([0, 1]), \preceq)$ un conjunto parcialmente ordenado. El orden $\preceq$ se dice admisible si verifica

- i) $\preceq$ es un orden total en $L([0, 1])$,
- ii) para todo $[\underline{a}, \bar{a}], [\underline{b}, \bar{b}] \in L([0, 1])$, $[\underline{a}, \bar{a}] \preceq [\underline{b}, \bar{b}]$ cuando $[\underline{a}, \bar{a}] \leq_{Lo} [\underline{b}, \bar{b}]$.

Los órdenes admisibles pueden generarse con funciones de agregación, así que recordemos su definición.

Definición 2.2: [4], [20] Sea $\mathcal{A} : \bigcup_{i=1}^n [0, 1]^n \rightarrow [0, 1]$ cumpliendo

- i) $\mathcal{A}(0, 0, \dots, 0) = 0, \mathcal{A}(1, 1, \dots, 1) = 1$,
- ii) $\mathcal{A}$ es monótona en cada variable,

entonces $\mathcal{A}$ se llama función de agregación.

Hemos restringido el método propuesto por Bustince et al. [6] al intervalo $[0, 1]$, ya que aunque el resultado original trabaja con funciones en $\mathbb{R}$, nosotros estamos centrados en el intervalo $[0, 1]$:

Proposición 2.1: Sean $\mathcal{A}, \mathcal{B} : [0, 1]^2 \rightarrow [0, 1]$ dos funciones de agregación continuas y sean $(u^*, v^*), (u', v') \in \{(u, v) \in [0, 1]^2 \mid u \leq v\}$. Las igualdades $\mathcal{A}(u^*, v^*) = \mathcal{A}(u', v')$ y $\mathcal{B}(u^*, v^*) = \mathcal{B}(u', v')$ sólo pueden darse si $(u^*, v^*) = (u', v')$. Definida la relación $\preceq_{\mathcal{A}, \mathcal{B}}$ en $L([0, 1])$ por $a \preceq_{\mathcal{A}, \mathcal{B}} b$ si y solo si

$$\mathcal{A}(\underline{a}, \bar{a}) < \mathcal{A}(\underline{b}, \bar{b})$$

o

$$\mathcal{A}(\underline{a}, \bar{a}) = \mathcal{A}(\underline{b}, \bar{b}) \text{ y } \mathcal{B}(\underline{a}, \bar{a}) \leq \mathcal{B}(\underline{b}, \bar{b}).$$

Entonces $\preceq_{\mathcal{A}, \mathcal{B}}$ es un orden admisible en $L([0, 1])$.

Con este método, no solo podemos construir nuevos órdenes admisibles, sino que también podemos comprobar que algunos de los órdenes más conocidos para intervalos son, en efecto, órdenes admisibles. Por ejemplo:

- Orden lexicográfico tipo 1 [6]: $a \preceq_{Lex1} b \Leftrightarrow \underline{a} < \underline{b}$ o $\underline{a} = \underline{b}$ y $\bar{a} \leq \bar{b}$, donde $\mathcal{A}(x, y) = x$ and $\mathcal{B}(x, y) = y$.
- Orden lexicográfico tipo 2 [6]: $a \preceq_{Lex2} b \Leftrightarrow \bar{a} < \bar{b}$ o $\bar{a} = \bar{b}$ y $\underline{a} \leq \underline{b}$, donde $\mathcal{A}(x, y) = y$ and $\mathcal{B}(x, y) = x$.
- Orden de Xu y Yager [27]: $a \preceq_{YX} b \Leftrightarrow \underline{a} + \bar{a} < \underline{b} + \bar{b}$ o $\underline{a} + \bar{a} = \underline{b} + \bar{b}$ y $\bar{a} - \underline{a} \leq \bar{b} - \underline{b}$, donde $\mathcal{A}(x, y) = x + y$ and $\mathcal{B}(x, y) = y - x$.

Por otro lado, dados dos conjuntos difusos μ_A y μ_B , se dice que μ_A está contenido en μ_B si $\mu_A(x) \leq \mu_B(x)$ [29]. Basándose en esta idea, si A y B son dos conjuntos difusos intervalo-valuados, decimos que A está o -contenido en B , y lo denotamos como $A \subseteq_o B$, si $A(x) \preceq_o B(x)$ para todo $x \in X$, donde $\preceq_o$ es una relación de orden en $L([0, 1])$. Teniendo en cuenta este punto de vista, Huidobro et al. [15] definieron la intersección de conjuntos difusos intervalo-valuados.

Definición 2.3: [15] Sean A, B dos conjuntos difusos intervalo-valuados en X y sea $\preceq_o$ una relación de orden en

$L([0, 1])$. Se llama o -intersección de A y B , $A \cap_o B$, al mayor conjunto difuso intervalo-valuado o -contenido en A y B .

En [15], Huidobro et al. probaron que esta definición depende del orden considerado en $L([0, 1])$, obteniéndose distintas intersecciones para los distintos órdenes, pero con algo en común si se utilizan órdenes totales.

Proposición 2.2: [15] Sea $\preceq_o$ un orden total en $L([0, 1])$. Para cualquier $A, B \in IVFS(X)$, la o -intersección de A y B es el conjunto difuso intervalo-valuado definido por:

$$A \cap_o B(x) = \begin{cases} A(x) & \text{si } A(x) \preceq_o B(x), \\ B(x) & \text{si } B(x) \preceq_o A(x). \end{cases}$$

En consecuencia, resulta evidente que si A está o -incluido en B , entonces la intersección de A y B es A , ya que es el conjunto más grande contenido en ambos.

III. T-NORMAS PARA INTERVALOS

En esta sección repasaremos algunas de las definiciones que se pueden encontrar en la literatura sobre t-normas para después presentar nuestra propuesta.

Las t-normas en $[0, 1]$ son funciones $t : [0, 1] \times [0, 1] \rightarrow [0, 1]$ asociativas, conmutativas y crecientes en cada argumento verificando que $t(1, u) = u$ para todo $u \in [0, 1]$. Para intervalos, Gehrke et al. [10] consideraron el orden reticular y el contenido usual. Partieron de que dado un punto $c \in L([0, 1])$, tiene un intervalo asociado $[c, c]$ y con varias premisas más llegaron a la siguiente definición:

Definición 3.1: [10] Una operación binaria T en $L([0, 1])$ que es asociativa y conmutativa es una t-norma si para todo $a, b, c \in L([0, 1])$ se cumplen las siguientes propiedades:

- i) $T([u, u], [v, v])$ es un intervalo degenerado (origen y extremo coincidentes), donde $u, v \in [0, 1]$,
- ii) $T(a, b \vee c) = T(a, b) \vee T(a, c)$,
- iii) $T(a, b \wedge c) = T(a, b) \wedge T(a, c)$,
- iv) $T(a, [1, 1]) = a$,
- v) $T([0, 1], [\underline{a}, \bar{a}]) = [\underline{a}, \bar{a}]$,

donde $\vee$ y $\wedge$ se definen como:

$$\begin{aligned} [\underline{a}, \bar{a}] \vee [\underline{b}, \bar{b}] &= [\underline{a} \vee \underline{b}, \bar{a} \vee \bar{b}], \\ [\underline{a}, \bar{a}] \wedge [\underline{b}, \bar{b}] &= [\underline{a} \wedge \underline{b}, \bar{a} \wedge \bar{b}]. \end{aligned}$$

Además, también dieron una caracterización de las t-normas para intervalos.

Teorema 3.1: [10] Toda t-norma en $L([0, 1])$ se puede poner de la forma $T(a, b) = [t(\underline{a}, \underline{b}), t(\bar{a}, \bar{b})]$, donde t es una t-norma en $[0, 1]$.

Además, fueron capaces de relajar la última condición de la definición si el operador binario es convexo, es decir, si $T(a, b) \leq_{Lo} c \leq_{Lo} T(d, e)$, entonces existen $r, s \in L([0, 1])$ con $a \leq_{Lo} r \leq_{Lo} d$ y $b \leq_{Lo} s \leq_{Lo} e$ de manera que $c = T(r, s)$.

Teorema 3.2: [10] Un operador binario convexo, conmutativo y asociativo, T en $L([0, 1])$, es una t-norma si para todo $a, b, c \in L([0, 1])$ se verifican las siguientes propiedades:

- i) $T([u, u], [v, v])$ es un intervalo degenerado (origen y extremo coincidentes), donde $u, v \in [0, 1]$,
- ii) $T(a, b \vee c) = T(a, b) \vee T(a, c)$,

- iii) $T(a, b \wedge c) = T(a, b) \wedge T(a, c)$,
- iv) $T(a, [1, 1]) = a$,
- v) $T([0, 1], [0, 1]) = [0, 1]$.

Bedregal y Takahashi [3] propusieron una extensión de la definición dada por Gehrke et al.:

Definición 3.2: [3] Una aplicación T en $L([0, 1])$ es una t-norma para intervalos si para todo $a, b, c, d \in L([0, 1])$ se verifica:

- i) $T(a, b) = T(b, a)$,
- ii) $T(a, T(b, c)) = T(T(a, b), c)$,
- iii) si $a \leq_{Lo} b$ y $c \leq_{Lo} d$, entonces $T(a, c) \leq_{Lo} T(b, d)$,
- iv) si $a \subseteq b$ y $c \subseteq d$, entonces $T(a, c) \subseteq T(b, d)$,
- v) $T(a, [1, 1]) = a$.

De este modo obtenemos t-normas para intervalos partiendo de t-normas.

Proposición 3.1: [3] Sean t_1 y t_2 dos t-normas en $[0, 1]$. Si $t_1 \leq t_2$, entonces la aplicación $I[t_1, t_2]$, definida como $I[t_1, t_2](a, b) = [t_1(\underline{a}, \underline{b}), t_2(\bar{a}, \bar{b})]$, es una t-norma para intervalos llamada t-norma para intervalos derivada de t_1 y t_2 .

Cuando consideramos un intervalo como un conjunto de números reales, tras hallar la intersección de a y b , $a \cap b$, tiene sentido que si $a_1 \subseteq a$ y $b_1 \subseteq b$, entonces $a_1 \cap b_1 \subseteq a \cap b$ [32]. Sin embargo, nosotros estamos considerando que un intervalo es una descripción imprecisa de la función de pertenencia, ya que estamos usando conjuntos difusos intervalo-valuados. Desde esta perspectiva, la monotonía en la inclusión es equivalente a ser creciente, esto es, si $A, B \in IVFS(X)$, $A \subseteq_o B$ si y solo si $A(x) \leq_o B(x), \forall x \in X$.

Por ello, en este trabajo consideraremos la definición propuesta por Deschrijver [9]. En ella, elimina la cuarta condición de la Definición 3.2.

Definición 3.3: [9] Una t-norma en $L([0, 1])$ es una función $T : L([0, 1]) \times L([0, 1]) \rightarrow L([0, 1])$ conmutativa, asociativa, creciente en ambos argumentos con respecto a $\leq_{Lo}$ y cumpliendo que $T([1, 1], a) = a$ para todo $a \in L([0, 1])$.

Partiendo de esta definición, podríamos definir t-norma para intervalos como:

Definición 3.4: Sea $\leq_o$ un orden en $L([0, 1])$. Una t-norma en $L([0, 1])$ es una función $T : L([0, 1]) \times L([0, 1]) \rightarrow L([0, 1])$ conmutativa, asociativa, creciente en ambos argumentos con respecto a $\leq_o$ y tal que $T([1, 1], a) = a$ para todo $a \in L([0, 1])$.

Una de las t-normas más conocidas es la función mínimo. En nuestro caso, el mínimo también es t-norma para intervalos.

Proposición 3.2: Sea $\leq_o$ un orden admisible en $L([0, 1])$. La aplicación $T_M : L([0, 1]) \times L([0, 1]) \rightarrow L([0, 1])$ definida como $T_M([\underline{a}, \bar{a}], [\underline{b}, \bar{b}]) = \min\{[\underline{a}, \bar{a}], [\underline{b}, \bar{b}]\}$ para cualquier $[\underline{a}, \bar{a}], [\underline{b}, \bar{b}] \in L([0, 1])$ es una t-norma para intervalos.

Como hemos visto previamente, son interesantes aquellos casos donde podemos relacionar t-normas para intervalos con t-normas para números (ver [2], [3], [9]).

Definición 3.5: Sea T una t-norma en $L([0, 1])$. T se dice representable si existen dos t-normas, t_1, t_2 tal que $T(a, b) = [t_1(\underline{a}, \underline{b}), t_2(\bar{a}, \bar{b})]$, para todo $a, b \in L([0, 1])$.

En general es más sencillo usar t-normas representables para intervalos, ya que simplifican la dificultad del problema.

Proposición 3.3: Sean t_1, t_2 dos t-normas con $t_1 \leq t_2$. La función $T : L([0, 1]) \times L([0, 1]) \rightarrow L([0, 1])$ definida como $T(a, b) = [t_1(\underline{a}, \underline{b}), t_2(\bar{a}, \bar{b})]$ es una t-norma para intervalos con respecto al orden reticular.

Nota 3.1: Este resultado no es cierto para cualquier orden. Si consideramos un ejemplo de orden admisible como es el orden lexicográfico tipo 1, podemos obtener que $[0, 2, 0, 7] \leq_{Lex1} [0, 3, 0, 4]$. Sin embargo, considerando que t_1 y t_2 sean la t-norma mínimo, llegamos a que $T([0, 1, 0, 8], [0, 2, 0, 7]) = [0, 1, 0, 7]$ y $T([0, 1, 0, 8], [0, 3, 0, 4]) = [0, 1, 0, 4]$. Por tanto, $T([0, 1, 0, 8], [0, 2, 0, 7]) \not\leq_{Lex1} T([0, 1, 0, 8], [0, 3, 0, 4])$, y podemos concluir que T no es creciente en la segunda componente.

Para otros órdenes admisibles pueden conseguirse contraejemplos similares.

En la Proposición 3.3 hemos obtenido un método para obtener t-normas aplicadas a intervalos a partir de t-normas, sin embargo, hemos visto en la Nota 3.1 que esto no siempre es posible. De hecho, hemos encontrado una caracterización en el siguiente resultado:

Proposición 3.4: Sea $\leq_o$ un orden total en $L([0, 1])$ y sean t_1 y t_2 dos t-normas en $[0, 1]$. Si consideramos la aplicación T definida como $T(a, b) = [t_1(\underline{a}, \underline{b}), t_2(\bar{a}, \bar{b})]$, entonces tenemos que T es una t-norma para intervalos si y solo si T es creciente en el segundo argumento.

IV. CONVEXIDAD DE CONJUNTOS DIFUSOS INTERVALO-VALUADOS UTILIZANDO T-NORMAS

En este trabajo hemos considerado la definición de convexidad propuesta por Huidobro et al. [15]:

Definición 4.1: Sea $(X, \leq)$ un espacio ordenado y sea $\leq_o$ una relación de orden en $L([0, 1])$. Un conjunto difuso intervalo-valuado A en X se dice o -convexo, si para cada $x \leq y \leq z$ en X se cumple $A(x) \leq_o A(y)$ o $A(z) \leq_o A(y)$.

Claramente esta definición generaliza la idea de convexidad en conjuntos difusos propuesta por Zadeh [29], ya que si consideramos un conjunto difuso μ_A como conjunto difuso intervalo-valuado, esto es, $A(x) = [\mu_A(x), \mu_A(x)], \forall x \in X$, diremos que A es convexo si se cumple $A(x) \leq_o A(y)$ o $A(z) \leq_o A(y)$. Esto es equivalente a decir que $[\mu_A(x), \mu_A(x)] \leq_o [\mu_A(y), \mu_A(y)]$ o $[\mu_A(z), \mu_A(z)] \leq_o [\mu_A(y), \mu_A(y)]$, que con el uso de cualquier orden que refine al orden reticular, por ejemplo cualquiera de los órdenes admisibles, equivale a $\mu_A(x) \leq \mu_A(y)$ o $\mu_A(z) \leq \mu_A(y)$. Con esto recuperamos la definición de convexidad para conjuntos difusos propuesta por Zadeh, que venía dada por: μ_A es convexo si y solo si $\mu_A(y) \geq \min\{\mu_A(x), \mu_A(z)\}$ para $x < y < z$.

Cuando trabajamos con órdenes totales, esta definición es equivalente a $\min\{A(x), A(z)\} \leq_o A(y)$, y nosotros hemos podido comprobar satisfactoriamente que el mínimo es una t-norma para intervalos. Partiendo de esa idea, reemplazaremos el mínimo por otra t-norma para intervalos. Por lo que proponemos la siguiente definición:

Definición 4.2: Sea $(X, \leq)$ un espacio ordenado, sea $\leq_o$ una relación de orden total en $L([0, 1])$ y sea T una t-norma para intervalos. Un conjunto difuso intervalo-valuado A en X

se dice T - o -convexo, si para cada $x \leq y \leq z$ en X se cumple $T(A(x), A(z)) \preceq_o A(y)$.

Comprobemos la validez de la definición comprobando su compatibilidad con la intersección.

Proposición 4.1: Sea $(X, \leq)$ un espacio ordenado, sea $\preceq_o$ una relación de orden total en $L([0, 1])$ y sea T una t -norma para intervalos. La intersección de dos conjuntos difusos intervalo-valuados T - o -convexos es T - o -convexa.

Este resultado apoya nuestra definición de t -normas para intervalos mientras que las otras consideradas no conservaban esta propiedad. Veamos un ejemplo.

Ejemplo 4.1: Consideremos el orden lexicográfico tipo 1 y la Definición 3.1. Según esta definición, el operador $T^*(a, b) = [\min\{a, b\}, \min\{\bar{a}, \bar{b}\}]$ es una t -norma en $L([0, 1])$. Consideremos $X = \{x, y, z\}$ con $x < y < z$ y los conjuntos difusos intervalo-valuados A y B definidos por $A(x) = [0, 3, 1]$, $A(y) = [0, 3, 0, 8]$ y $A(z) = [0, 7, 0, 8]$, y $B(x) = [0, 4, 0, 5]$, $B(y) = [0, 5, 0, 7]$ y $B(z) = [0, 6, 0, 9]$.

La intersección resulta ser $(A \cap_{Lex1} B)(x) = [0, 3, 1]$, $(A \cap_{Lex1} B)(y) = [0, 3, 0, 8]$ y $(A \cap_{Lex1} B)(z) = [0, 6, 0, 9]$. Es fácil comprobar que A y B son T^* - $Lex1$ -convexos, pero $A \cap_{Lex1} B$ no, ya que $T^*(A \cap_{Lex1} B(x), A \cap_{Lex1} B(z)) = T^*([0, 3, 1], [0, 6, 0, 9]) = [0, 3, 0, 9] \not\preceq_{Lex1} [0, 3, 0, 8] = A \cap_{Lex1} B(y)$. Este contraejemplo es válido también para la Definición 3.2.

V. CONCLUSIONES

En este trabajo hemos presentado una definición de t -norma para intervalos. Utilizando este concepto, hemos generalizado la noción de convexidad para conjuntos difusos intervalo-valuados y hemos podido comprobar que la definición propuesta es compatible con la intersección de conjuntos difusos intervalo-valuados.

REFERENCIAS

- [1] E. Ammar and J. Metz. On fuzzy convexity and parametric fuzzy optimization. *Fuzzy Sets and Systems*, 49(2):135 – 141, 1992.
- [2] B. Bedregal and A. Takahashi. Interval t -norms as interval representations of t -norms. In *The 14th IEEE International Conference on Fuzzy Systems, 2005. FUZZ '05.*, pages 909–914, 2005.
- [3] B. Bedregal and A. Takahashi. The best interval representations of t -norms and automorphisms. *FuzzySetsandSystems*, 157:3220–3230, 2006.
- [4] G. Beliakov, H. Bustince, and T. Calvo. *A Practical Guide to Averaging Functions*. Springer, 2016.
- [5] H. Bustince and P. Burillo. Mathematical analysis of interval-valued fuzzy relations: Application to approximate reasoning. *Fuzzy Sets and Systems*, 113(2):205 – 219, 2000.
- [6] H. Bustince, J. Fernandez, A. Kolesárová, and R. Mesiar. Generation of linear orders for intervals by means of aggregation functions. *Fuzzy Sets and Systems*, 220:69 – 77, 2013.
- [7] C. Cornelis, G. Deschrijver, and E.E. Kerre. Implication in intuitionistic fuzzy and interval-valued fuzzy set theory. *International Journal of Approximate Reasoning*, 35:55 – 95, 2004.
- [8] S. Díaz, E. Induráin, V. Janiš, J. V. Llinares, and S. Montes. Generalized convexities related to aggregation operators of fuzzy sets. *Kybernetika*, 53:383 – 393, 2017.
- [9] G. Deschrijver. Arithmetic operators in interval-valued fuzzy set theory. *Information Sciences*, 177(14):2906–2924, 2007.
- [10] M. Gehrke, C. Walker, and E. Walker. Some comments on interval valued fuzzy sets. *International Journal of Intelligent Systems*, 11(10):751–759, 1996.
- [11] J.A. Goguen. L -fuzzy sets. *Journal of Mathematical Analysis and Applications*, 18(1):145 – 174, 1967.
- [12] M.B. Gorzalczy. A method of inference in approximate reasoning based on interval-valued fuzzy sets. *Fuzzy Sets and Systems*, 21:1 – 17, 1987.
- [13] I. Grattan-Guinness. Fuzzy membership mapped onto intervals and many-valued quantities. *Mathematical Logic Quarterly*, 22(1):149 – 160, 1976.
- [14] S.K. Gupta and D. Dangar. Duality for a class of fuzzy nonlinear optimization problem under generalized convexity. *Fuzzy Optimization and Decision Making*, 13:131 – 150, 2014.
- [15] P. Huidobro, P. Alonso, V. Janiš, and S. Montes. Orders preserving convexity under intersections for interval-valued fuzzy sets. In Marie-Jeanne Lesot, Susana Vieira, Marek Z. Reformat, João Paulo Carvalho, Anna Wilbik, Bernadette Bouchon-Meunier, and Ronald R. Yager, editors, *Information Processing and Management of Uncertainty in Knowledge-Based Systems*, pages 493 – 505, Cham, 2020. Springer International Publishing.
- [16] K.U. Jahn. Intervall-wertige mengen. *Mathematische Nachrichten*, 68:115 – 132, 1975.
- [17] Y. Kim, Z. Xi, and J. Lien. Disjoint convex shell and its applications in mesh unfolding. *Computer-Aided Design*, 90:180 – 190, 2017.
- [18] M. Li, W. Yang, and X. Zhang. Projection on convex set and its application in testing force closure properties of robotic grasping. In *Intelligent Robotics and Applications*, volume 6425, pages 240 – 251. Springer, 11 2010.
- [19] L. Liberti. Reformulation and convex relaxation techniques for global optimization. *Quarterly Journal of the Belgian, French and Italian Operations Research Societies*, 2(3):255 – 258, 2004.
- [20] R. Mesiar and M. Komorníková. Aggregation functions on bounded posets. In Chris Cornelis, Glad Deschrijver, Mike Nachtegaal, Steven Schockaert, and Yun Shi, editors, *35 Years of Fuzzy Set Theory: Celebratory Volume Dedicated to the Retirement of Etienne E. Kerre*, pages 3 – 17, Berlin, Heidelberg, 2011. Springer.
- [21] J. Ramik and M. Vlach. *Generalized Concavity in Fuzzy Optimization and Decision Analysis*. Kluwer Academic Publishers, 2002.
- [22] R. Sambuc. Function phi-floous, application a l'aide au diagnostic en pathologie thyroïdienne. *These de Doctorat en Medicine*, 1975.
- [23] D. Sarkar. Concavoconvex fuzzy set. *Fuzzy Sets and Systems*, 79:267 – 269, 1996.
- [24] Y. Syau and E. Lee. Fuzzy convexity and multiobjective convex optimization problems. *Computers & Mathematics with Applications*, 52(3):351 – 362, 2006.
- [25] M. Tofighi, O. Yorulmaz, K. Kose, D.C. Kahraman, R. Cetin-Atalay, and A.E. Cetin. Phase and tv based convex sets for blind deconvolution of microscopic images. *IEEE Journal of Selected Topics in Signal Processing*, 10(1):81 – 91, 2015.
- [26] I.B. Turksen and Z. Zhong. An approximate analogical reasoning schema based on similarity measures and interval-valued fuzzy sets. *Fuzzy Sets and Systems*, 34:323 – 346, 1990.
- [27] Z.S. Xu and R.R. Yager. Some geometric aggregation operators based on intuitionistic fuzzy sets. *International Journal of General Systems*, 35(4):417 – 433, 2006.
- [28] X. Yang. A property on convex fuzzy sets. *Fuzzy Sets and Systems*, 126:269 – 271, 2002.
- [29] L.A. Zadeh. Fuzzy sets. *Information and Control*, 8(3):338 – 353, 1965.
- [30] L.A. Zadeh. The concept of a linguistic variable and its application to approximate reasoning—I. *Information Sciences*, 8(3):199 – 249, 1975.
- [31] J. Zhang, Q. Zheng, X. Ma, and L. Li. Relationships between interval-valued vector optimization problems and vector variational inequalities. *Fuzzy Optimization and Decision Making*, 15:33 – 55, 2016.
- [32] Q. Zuo. Description of strictly monotonic interval and/or operations. *Reliable Computing*, pages 23 – 25, 1995.

XX Congreso Español sobre Tecnologías y Lógica Fuzzy

SESIÓN ESPECIAL:

TOMA DE DECISIONES CON INFORMACIÓN

DIFUSA



Preferencias no lineales en Toma de Decisión en Grupo. Amplificación de Valores Extremos

1 st Diego García-Zamora	2 nd Álvaro Labella	3 rd Rosa M. Rodríguez	4 st Luis Martínez
Departamento de Informática	Departamento de Informática	Departamento de Informática	Departamento de Informática
Universidad de Jaén	Universidad de Jaén	Universidad de Jaén	Universidad de Jaén
Jaén, España	Jaén, España	Jaén, España	Jaén, España
dgzamora@ujaen.es	alabella@ujaen.es	rmrodrig@ujaen.es	martin@ujaen.es

Resumen—Varios estudios han demostrado que el uso de escalas no lineales mejoran las decisiones obtenidas en problemas de Toma de Decisión en Grupo (TDG). Este trabajo está orientado a incorporar estas escalas no lineales en Procesos de Alcance de Consenso (PAC), los cuales son la herramienta fundamental para suavizar los conflictos que aparecen en los problemas de TDG. Para ello, utilizaremos automorfismos no lineales definidos en el intervalo unidad para deformar las preferencias de los expertos, expresadas mediante relaciones de preferencia difusas, con el objetivo de obtener escalas más realistas. Este trabajo introduce estas deformaciones no lineales como Amplificaciones de Valores Extremos (AVEs), analiza sus principales propiedades y presenta dos familias concretas de AVEs: una basada en la función seno y otra basada en polinomios. Por último estudiamos mediante un caso práctico cómo influye este enfoque no lineal basado en AVEs en dos modelos de consenso clásicos para TDG.

Palabras clave—Toma de Decisión en Grupo, Procesos de Alcance de Consenso, Amplificación de Valores Extremos, Preferencias no lineales.

I. INTRODUCCIÓN

Tradicionalmente se han utilizado escalas lineales para modelar las preferencias en los problemas de Toma de Decisión en Grupo (TDG). Sin embargo, estudios recientes han probado que se obtienen mejores decisiones usando escalas no lineales para representar las preferencias de los expertos [1], [2]. Sin embargo, ninguna de esas propuestas tiene en cuenta los conflictos entre expertos que normalmente aparecen en los problemas de TDG.

En este trabajo estudiaremos el efecto de deformar las preferencias de los expertos mediante escalas no lineales en Procesos de Alcance de Consenso (PACs) para problemas de TDG. Asumiremos que las preferencias de los expertos vienen dadas por medio de Relaciones de Preferencia Difusas (RPDs) y aplicaremos una deformación no lineal a cada una de estas preferencias para *ajustar* los valores iniciales a una escala más realista. Por último, estudiaremos el impacto de este enfoque no lineal en los modelos de consenso propuestos en [3], [4] analizando el grado de consenso obtenido y número de rondas empleado.

Esta contribución está dividida como sigue: En la Sección II se revisan brevemente los problemas de TDG y los PAC. La Sección III está dedicada a la noción principal de este trabajo, la Amplificación de Valores Extremos (AVE), definidas como aquellos automorfismos del intervalo $[0, 1]$ que incrementan

la distancia entre valores extremos de las preferencias dadas mediante RPDs. En la Sección IV se muestra cómo influyen estas escalas no lineales en los PAC. Finalmente, la Sección V concluye el trabajo.

II. TOMA DE DECISIÓN EN GRUPO Y PROCESOS DE ALCANCE DE CONSENSO

En esta sección se revisan brevemente los principales conceptos relativos a TDG y PACs.

Un problema de TDG es una situación en la que un grupo de expertos $E = \{e_1, e_2, \dots, e_m\}$, $2 \leq m \in \mathbb{N}$, tiene que elegir la mejor solución dentro de un conjunto de posibles alternativas $X = \{X_1, X_2, \dots, X_n\}$, $2 \leq n \in \mathbb{N}$. Sin pérdida de generalidad, podemos suponer que las opiniones de cada experto viene dada por medio de una RPD, la cual consiste en una matriz $P_k \in \mathcal{M}_{n \times n}([0, 1])$ donde cada entrada $p_{ij}^k \in [0, 1]$ representa el grado en el que el experto e_k prefiere la alternativa X_i sobre la X_j . Las RPDs verifican la condición de simetría $p_{ij}^k + p_{ji}^k = 1 \forall i, j \in \{1, 2, \dots, n\}$, $k \in \{1, 2, \dots, m\}$, conocida como reciprocidad aditiva.

Es posible que en el proceso de resolución de un problema de TDG surjan conflictos entre las opiniones de los expertos y que algunos consideren que sus opiniones no se han tenido suficientemente en cuenta durante el proceso [5], [6]. En estos casos se aplican PACs para lograr un acuerdo en la solución elegida [7], [8]. Un PAC es un proceso iterativo que usa una medida de consenso para calcular la cercanía entre las preferencias de los expertos [6] y finaliza cuando se alcanza un grado de consenso predefinido o se alcanza un determinado número de rondas. En este trabajo usaremos los modelos de consenso propuestos en [3], [4], ya que han demostrado un buen funcionamiento [8].

III. AMPLIFICACIÓN DE VALORES EXTREMOS

En esta sección se introduce la definición de AVE, así como algunos ejemplos concretos de familias de funciones que cumplen con los requisitos que caracterizan a estas AVEs.

Definición 1 (Amplificación de Valores Extremos). *Llamaremos Amplificación de Valores Extremos (AVE) a una función $D : [0, 1] \rightarrow [0, 1]$ verificando:*

1. *D es un automorfismo en el intervalo $[0, 1]$,*

2. D es una función de clase C^1 , esto es, D es derivable en $[0, 1]$ y su derivada D' es continua en $[0, 1]$,
3. D verifica $D(x) = 1 - D(1 - x) \forall x \in [0, 1]$,
4. $D'(0) > 1$ y $D'(1) > 1$,
5. D es cóncava en un entorno de 0 y convexa en un entorno de 1.

El propósito de las AVEs es transformar las preferencias de un problema de TDG de forma que los nuevos valores obedezcan una escala no lineal en la que las distancias entre los valores más extremos se ven incrementadas respecto a las diferencias originales.

Nota 1. Aunque el propósito y las características de las AVEs están orientadas al modelado de las escalas no lineales para las preferencias en PAC, el hecho de que permitan definir medidas de similaridad junto al enfoque particular que proporcionan a la hora de tratar la información más extrema recuerda al Hypermatching introducido por Yager y Petry [14]. En el futuro nos gustaría estudiar detalladamente cómo relacionar ambas propuestas.

La primera propiedad, junto con las condiciones de regularidad de la segunda, garantizan que la biyectividad de las AVEs. Esto es fundamental al comparar preferencias, puesto que no deseamos que dos valores diferentes de las preferencias originales tengan por imagen el mismo valor. La tercera propiedad está orientada a asegurar la reciprocidad aditiva de la matriz obtenida al aplicar la AVE a cada uno de los ítems de una RPD. La cuarta propiedad está relacionada con la amplificación de las distancias entre los valores extremos de las preferencias originales (los cercanos a 0 y a 1), mientras que la quinta permite deducir que a medida que nos acercamos a los extremos, la amplificación de las distancias será mayor. Formalmente:

Teorema 1. Sea $D : [0, 1] \rightarrow [0, 1]$ un AVE. Entonces

1. La función $d_D : [0, 1] \times [0, 1] \rightarrow [0, 1]$ definida por

$$d_D(x, y) = |D(x) - D(y)| \forall x, y \in [0, 1],$$

es una Disimilitud Restringida [9] y la función $S_D : [0, 1] \times [0, 1] \rightarrow [0, 1]$ definida por

$$S_D(x, y) = 1 - |D(x) - D(y)| \forall x, y \in [0, 1].$$

es una Función de Equivalencia Restringida [9],

2. Podemos encontrar tres intervalos $I_1, I_2, I_3 \subset [0, 1]$ tales que $0 \in I_1, 1 \in I_3$, y $I_1 < I_2 < I_3$ verificando

$$|D(y) - D(x)| > |y - x| \forall x, y \in I_1 : x \neq y,$$

$$|D(y) - D(x)| < |y - x| \forall x, y \in I_2 : x \neq y,$$

$$|D(y) - D(x)| > |y - x| \forall x, y \in I_3 : x \neq y,$$

3. El gráfico de D está por encima de la diagonal del cuadrado $[0, 1] \times [0, 1]$ para valores cercanos a 0 y está bajo la misma diagonal para los valores cercanos a 1,
4. Existen un entorno U_0 de 0 y un entorno U_1 de 1 tales que para cada par $x, y \in U_0^o, x < y$ existe $h_0 > 0$ tal que $|D(x) - D(x-t)| \geq |D(y) - D(y-t)| \forall t \in [0, h_0]$

y para cada par $x, y \in U_1^o, x < y$, existe $h_1 > 0$ tal que $|D(x-t) - D(x)| \leq |D(y-t) - D(y)| \forall t \in [0, h_1]$.

Demostración. La prueba de este resultado es consecuencia del Teorema del Valor Medio. Omitimos los detalles concretos por limitación de espacio. $\square$

Las cuatro tesis de este resultado permiten interpretar las características de las AVEs de la siguiente manera:

1. Transforman RPDs en RPDs.
2. Amplifican las distancias entre los valores extremos y reducen la distancia entre los valores intermedios.
3. Su gráfica tiene un patrón común (ver Figura 1).
4. La amplificación de las distancias se incrementa a medida que nos acercamos a los extremos.

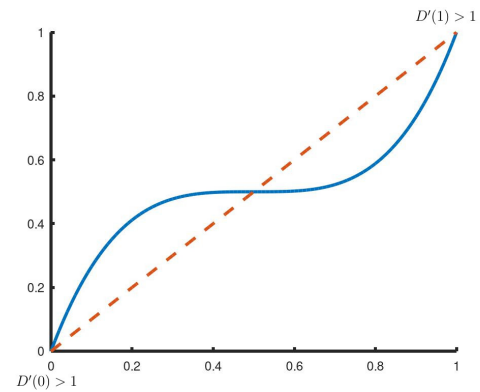


Figura 1: Esbozo del gráfico de una AVE

A continuación presentamos un par de familias de AVEs.

AVEs basadas en la función seno.

Sea $\alpha \in]0, \frac{1}{2\pi}]$. La función de clase C^∞ $s_\alpha : [0, 1] \rightarrow [0, 1]$ dada por

$$s_\alpha(x) = x - \alpha \cdot \sin(2\pi x - \pi) \forall x \in [0, 1]$$

es una AVE (ver la Figura 2).

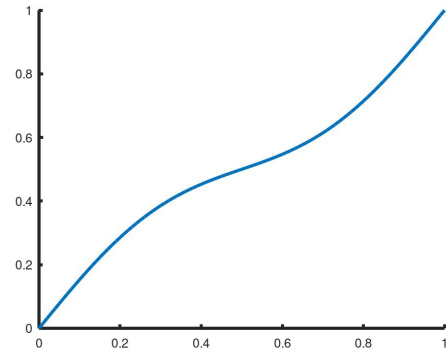


Figura 2: AVE $s_{0,09}$

AVEs basadas en polinomios.

Fijado $\alpha > 1$, la función $m_\alpha : [0, 1] \rightarrow [0, 1]$ dada por

$$m_\alpha(x) = \begin{cases} \frac{1}{2} - \frac{1}{2}(1-2x)^\alpha & 0 \leq x < \frac{1}{2} \\ \frac{1}{2} + \frac{1}{2}(2x-1)^\alpha & \frac{1}{2} \leq x \leq 1 \end{cases},$$

es una AVE (ver las Figuras 3 y 4).

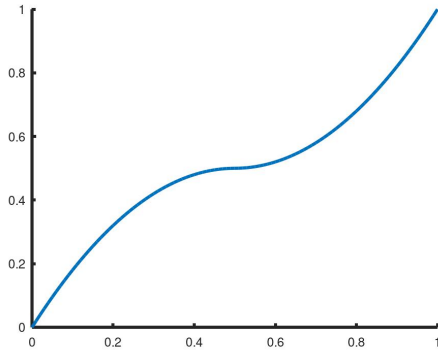


Figura 3: AVE m_2 .

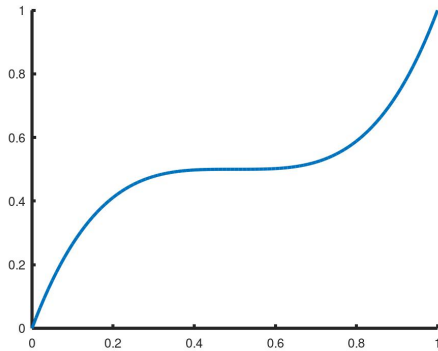


Figura 4: AVE $m_{3,39}$.

IV. AMPLIFICACIÓN DE VALORES EXTREMOS EN PROCESOS DE ALCANCE DE CONSENSO

En esta sección mostramos el funcionamiento de las AVEs cuando son utilizadas en PAC para TDG. Para ello, hemos implementado el enfoque no lineal de las AVEs en los PAC ampliamente utilizados en la literatura y presentados en [3], [4], que están implementados en el software AFRYCA [6].

El problema de TDG considerado consiste en grupo de 100 expertos que enuncian sus preferencias sobre las alternativas $X = \{X_1, X_2, X_3, X_4\}$ por medio de RPDs (ver [10]). Para hacer las simulaciones hemos usado los valores por defecto que AFRYCA define para los parámetros de cada modelo. El umbral de consenso deseado se ha establecido en 0.85 y el máximo número de rondas permitido es 15. Los resultados obtenidos se muestran en las Tablas I y II.

AFRYCA [6] proporciona una herramienta para visualizar los resultados de las distintas simulaciones mediante la técnica de escalado multidimensional [11] (ver Figuras 5 y 6). Esta representación muestra la opinión colectiva de los expertos en el centro del gráfico y a su alrededor las preferencias de los expertos.

Tabla I: Resultados con Herrera-Viedma et al. [3]

AVE	Orden de alternativas	Rondas	Consenso
Clásico	$x_1 \succ x_2 \succ x_4 \succ x_3$	6	0.87
$s_{0,08}$	$x_1 \succ x_2 \succ x_4 \succ x_3$	6	0.89
$s_{0,09}$	$x_1 \succ x_2 \succ x_4 \succ x_3$	6	0.92
m_2	$x_1 \succ x_2 \succ x_4 \succ x_3$	5	0.86
$m_{3,39}$	$x_1 \succ x_2 \succ x_4 \succ x_3$	5	0.91

Tabla II: Resultados con Quesada et al. [4]

AVE	Orden de alternativas	Rondas	Consenso
Clásico	$x_4 \succ x_1 \sim x_2 \sim x_3$	10	0.85
$s_{0,08}$	$x_4 \succ x_1 \succ x_2 \sim x_3$	7	0.86
$s_{0,09}$	$x_4 \succ x_1 \succ x_2 \sim x_3$	7	0.87
m_2	$x_4 \succ x_1 \succ x_2 \sim x_3$	7	0.86
$m_{3,39}$	$x_4 \succ x_1 \succ x_2 \sim x_3$	5	0.87

El modelo clásico [3] alcanzó un grado de consenso de 0,87 en 6 rondas. Para este modelo, las AVEs $s_{0,08}$ y $s_{0,09}$ han mejorado el grado de consenso alcanzado, aunque no han reducido el número de rondas. Las AVEs polinómicas m_2 y $m_{3,39}$ han reducido el número de rondas, obteniendo grados de consenso de 0,86 y 0,91 respectivamente en 5 rondas.

Por otro lado, el modelo clásico [4] necesitó 10 rondas para lograr un grado de consenso de 0,85. En este caso, ambas familias de AVEs han mejorado significativamente el modelo original. Las AVEs $s_{0,08}$, $s_{0,09}$ y m_2 han mejorado ligeramente el grado de consenso en 7 rondas, mientras que la AVE $m_{3,39}$ destaca por haber incrementado el grado de consenso en sólo 5 rondas. Las simulaciones indican que las AVEs mejoran ambos modelos ya que o bien disminuyen el número de rondas o bien incrementan el grado de consenso alcanzado.

Para analizar el comportamiento de las AVEs en un PAC desde un punto de vista teórico es necesario tener en cuenta dos aspectos clave. Por un lado, es un hecho probado que en los modelos de consenso los valores menos extremos de las preferencias aportan cohesión al grupo y favorecen el alcance del consenso colectivo [12], [13]. Por otro lado, si bien es cierto que las AVEs amplifican la distancia entre los valores extremos, cuando usamos una AVE los valores intermedios se acercan entre sí. Considerando estos dos hechos conjuntamente podemos deducir que al aplicar una AVE a las preferencias que van a ser usadas en un modelo de consenso, el modelo de consenso modificado por la AVE alcanzará generalmente el grado de consenso deseado más rápido que el modelo original puesto que los valores intermedios de las preferencias, esto es, la información *más importante* para el consenso, están *más cerca* desde el principio.

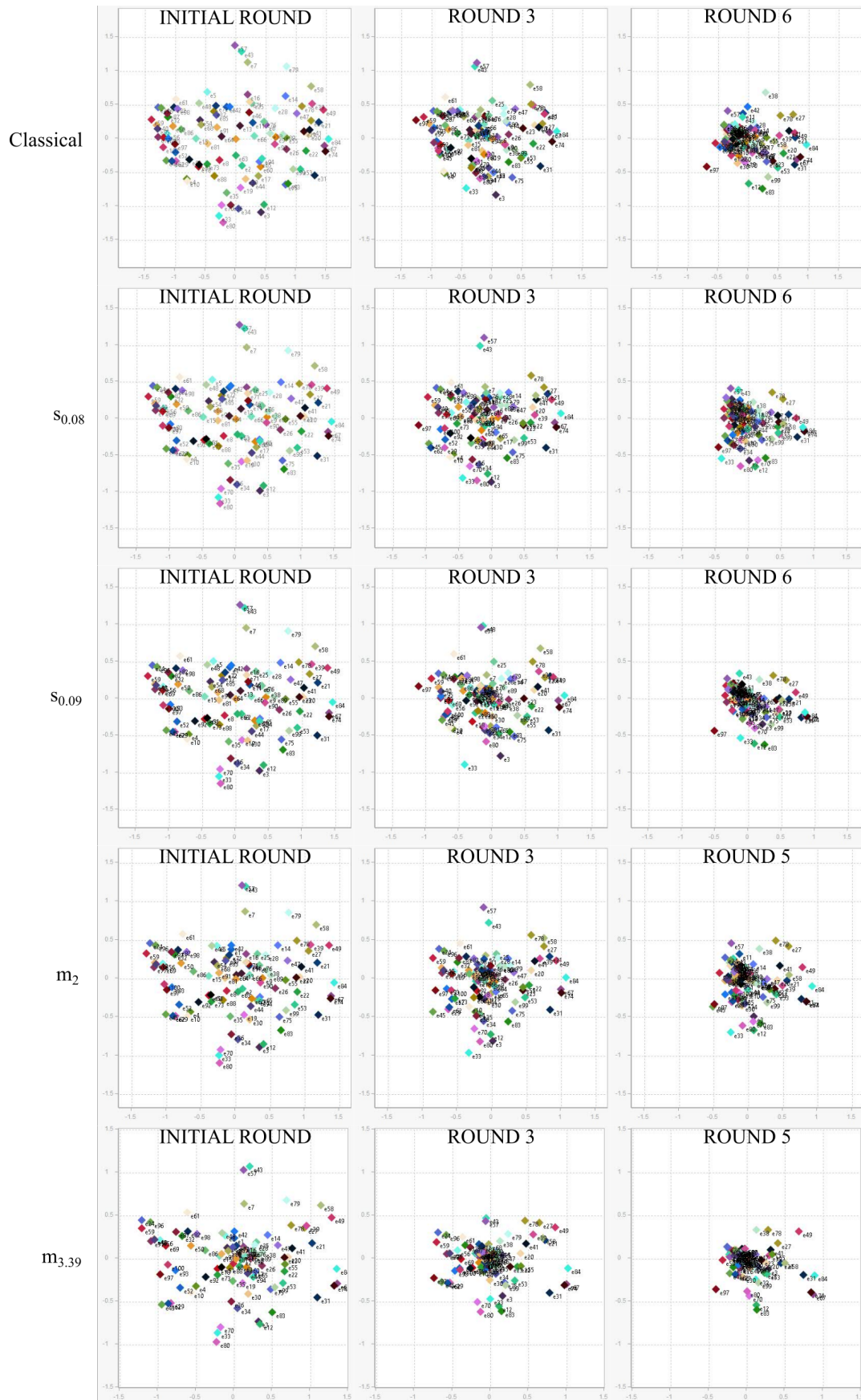


Figura 5: Herrera Viedma et al. [3] simulations.

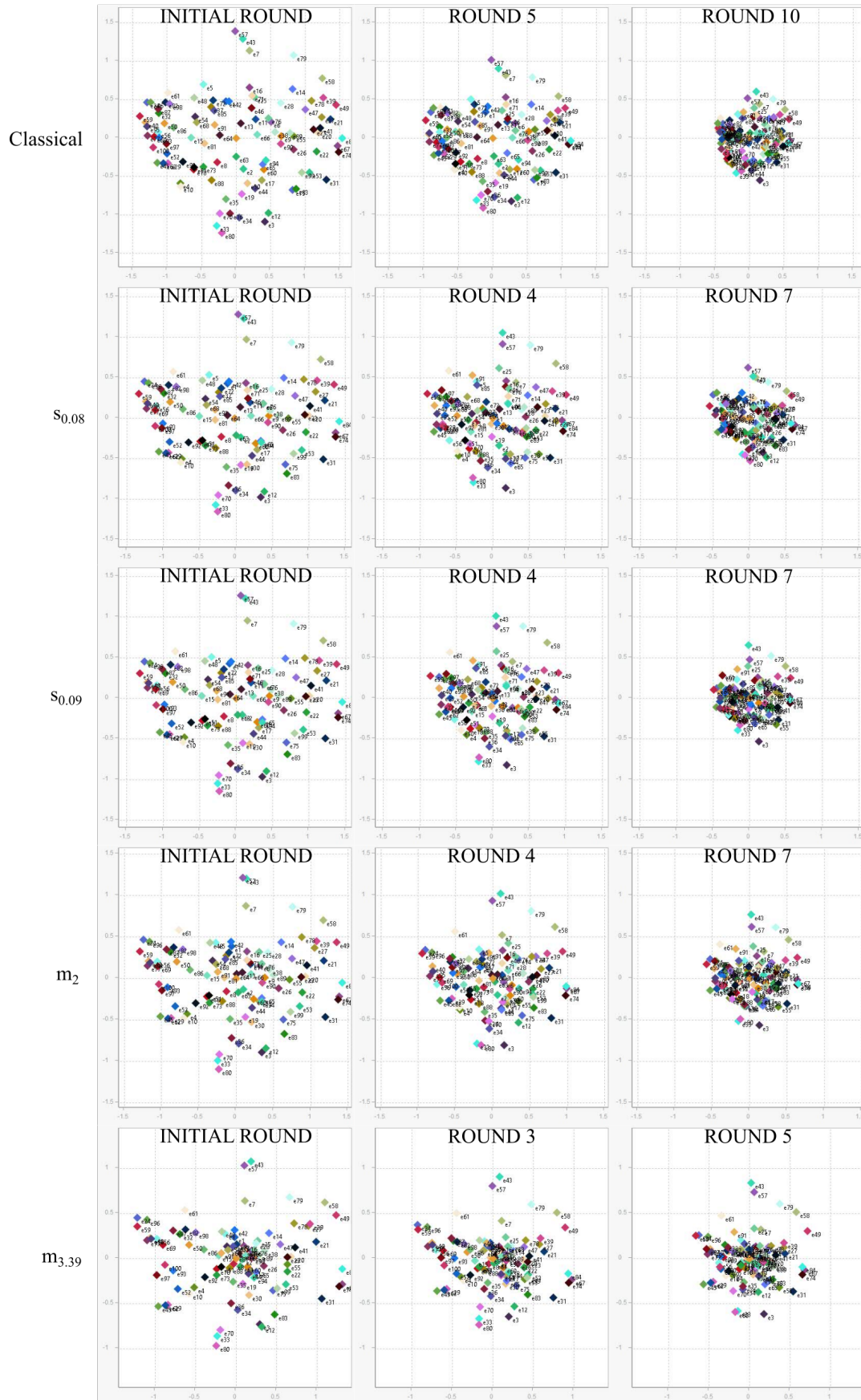


Figura 6: Quesada et al. [4] simulations.

Para que los resultados obtenidos sean fiables, es imprescindible tener en cuenta que aunque la AVE usada en el PAC modifique los valores iniciales de las preferencias, el orden de preferencia de las alternativas en cada RPD no sufre alteraciones significativas. Omitimos aquí los detalles concretos por motivos de espacio, pero es posible demostrar que las familias de AVEs propuestas en este trabajo, s_α , y m_α no introducen cambios significativos en el orden de las alternativas preferidas por cada experto.

V. CONCLUSIONES

Los PAC clásicos asumen escalas lineales para los valores de preferencias de los expertos. Este trabajo propone el uso de escalas no lineales en estos modelos con el objetivo de modelar las preferencias de los expertos de una forma más realista.

Este trabajo presenta las propiedades analíticas de estas escalas no lineales y estudia sus principales propiedades. Estas deformaciones no lineales han recibido el nombre de AVEs y se caracterizan por transformar RPDs en RPDs de forma que las distancias entre los valores extremos se incrementan mientras que las distancias entre los valores intermedios se reducen. Para mostrar el impacto de las AVEs en los PAC, hemos utilizado el software AFRYCA para simular el funcionamiento de dos modelos de consenso cuando se combinan con las AVEs definidas en este trabajo. Los resultados muestran que o bien se reducen el número de rondas necesarias para alcanzar el consenso o bien se incrementa el grado de consenso.

Como trabajos futuros definiremos nuevas AVEs e intentaremos optimizar los parámetros de las AVEs existentes para modelos de consenso concretos de acuerdo con diferentes métricas. También pretendemos analizar el comportamiento de las AVEs en otros modelos de consenso, así como aplicar este marco teórico para resolver problemas de decisión del mundo real.

AGRADECIMIENTOS

Este trabajo ha sido parcialmente subvencionado por el Ministerio de Economía y Competitividad mediante el proyecto PGC2018-099402-B-I00, y la Ayuda Postdoctoral Ramón y Cajal (RYC-2017-21978), y por el Ministerio de Ciencia, Innovación y Universidades mediante la Ayuda Predoctoral para la Formación de Profesorado Universitario (FPU2019/01203).

REFERENCIAS

- [1] J. Masthoff, "Group modeling: Selecting a sequence of television items to suit a group of viewers," *User Model. User-Adapt. Interact.*, vol. 14, pp. 37–85, 02 2004.
- [2] A. Delic, F. Ricci, and J. Neidhardt, "Preference networks and non-linear preferences in group recommendations," in *IEEE/WIC/ACM International Conference on Web Intelligence*, 10 2019, pp. 403–407.
- [3] E. Herrera-Viedma, F. Herrera, and F. Chiclana, "A consensus model for multiperson decision making with different preference structures," *Systems, Man and Cybernetics, Part A: Systems and Humans, IEEE Transactions on*, vol. 32, pp. 394 – 402, 06 2002.
- [4] F. J. Quesada, I. Palomares, and L. Martínez, "Managing experts behavior in large-scale consensus reaching processes with uninorm aggregation operators," *Applied Soft Computing*, vol. 35, pp. 873–887, 2015.
- [5] Á. Labella, Y. Liu, R. M. Rodríguez, and L. Martínez, "Analyzing the performance of classical consensus models in large scale group decision making: A comparative study," *Applied Soft Computing*, vol. 67, pp. 677–690, 2018.
- [6] I. Palomares, F. Estrella, L. Martínez, and F. Herrera, "Consensus under a fuzzy context: Taxonomy, analysis framework afryca and experimental case of study," *Information Fusion*, vol. 20, 11 2014.
- [7] R. M. Rodríguez, Á. Labella, G. D. Tré, and L. Martínez, "A large scale consensus reaching process managing group hesitation," *Knowledge-Based Systems*, vol. 159, pp. 86–97, 2018.
- [8] Á. Labella, H. Liu, R. M. Rodríguez, and L. Martínez, "A cost consensus metric for consensus reaching processes based on a comprehensive minimum cost model," *European Journal of Operational Research*, vol. 281, p. 316–331, 2020.
- [9] H. Bustince, E. Barrenechea, and M. Pagola, "Relationship between restricted dissimilarity functions, restricted equivalence functions and normal en-functions: Image thresholding invariant," *Pattern Recognition Letters*, vol. 29, no. 4, pp. 525 – 536, 2008.
- [10] sinbad2. Data set with 100 experts. [Online]. Available: <https://sinbad2.ujaen.es/afryca/sites/default/files/dataset/EVA100Preferences.pdf>
- [11] L. Blouvshtein and D. Cohen-Or, "Outlier detection for robust multi-dimensional scaling," *IEEE transactions on pattern analysis and machine intelligence*, vol. 41, no. 9, pp. 2273–2279, 2018.
- [12] A. Sirbu, V. Loreto, V. Servedio, and F. TRIA, "Cohesion, consensus and extreme information in opinion dynamics," *Advances in Complex Systems*, vol. 16, no. 06, p. 1350035, 2013.
- [13] A. Sirbu, V. Loreto, and V. Servedio, "Opinion dynamics with disagreement and modulated information," *Journal of Statistical Physics*, vol. 151, pp. 218–237, 2013.
- [14] R. R. Yager and F. E. Petry, "Hypermatching: Similarity matching with extreme values," *IEEE Transactions on Fuzzy Systems*, vol. 22, no. 4, pp. 949–957, 2014.

Estructuras de preferencia para relaciones recíprocas borrosas: caracterización y compatibilidad

Fabian Castiblanco

Faculty of Economics
Gran Colombia University
Bogotá, Colombia

fabianalberto.castiblanco@ugc.edu.co

Camilo Franco

Faculty of Engineering
Los Andes University
Bogotá, Colombia

c.franco31@uniandes.edu.co

J. Tinguaro Rodríguez

Faculty of Mathematics
Complutense University
Madrid, Spain

jtrodrig@mat.ucm.es

Javier Montero

Faculty of Mathematics
Complutense University
Madrid, Spain

monty@mat.ucm.es

Abstract—El empleo de relaciones borrosas $R : \mathbb{A}^2 \rightarrow [0, 1]$ está ampliamente extendido de cara a representar grados de preferencia entre alternativas en $\mathbb{A}$. En el trabajo recientemente aceptado [1] se reflexiona en primer lugar sobre la aparente incompatibilidad entre las relaciones recíprocas y el marco estándar de preferencias borrosas establecido en [2], proponiendo a continuación un marco general para estudiar el significado preferencial de las relaciones recíprocas a partir de la semántica, bien establecida, de las relaciones y estructuras de preferencia del modelo estándar. Una consecuencia de este análisis es que es posible dotar a las relaciones recíprocas de una estructura de preferencia propia con dos relaciones, preferencia estricta y ausencia de preferencia. Además, es posible dotar de un significado unívoco a esta estructura sobre la base de las estructuras de preferencia estándar cuando se cumplen ciertas condiciones de compatibilidad entre ambos modelos.

Index Terms—Relaciones de preferencia borrosa, relaciones recíprocas, estructuras de preferencia, semántica preferencial

I. INTRODUCCIÓN

Considérese un problema de decisión con conjunto de alternativas $\mathbb{A}$, y sea $R : \mathbb{A}^2 \rightarrow [0, 1]$ una relación de preferencia borrosa sobre $\mathbb{A}$. R es una relación de preferencia débil si es reflexiva, esto es, si $R(a, a) = 1 \forall a \in \mathbb{A}$, mientras que R es estricta si es antireflexiva, $R(a, a) = 0 \forall a \in \mathbb{A}$, y R se dice recíproca cuando $R(a, b) + R(b, a) = 1 \forall a, b \in \mathbb{A}$. Sean respectivamente $\mathcal{R}$ y $\mathcal{G}$ los conjuntos de todas las relaciones de preferencia borrosas débiles y recíprocas. Según se discute en [1], las relaciones de preferencia borrosas recíprocas presentan una incompatibilidad esencial con las relaciones de preferencia borrosa estándar, tanto débiles como estrictas, debido a que la condición de reciprocidad es incompatible con las propiedades de reflexividad y antireflexividad. Esto plantea dudas sobre el significado del predicado preferencial de una relación recíproca, así como sobre las situaciones preferenciales que un modelo de tipo recíproco puede expresar. Para estudiar esta cuestión, en [1] se plantea dotar a las relaciones recíprocas de una estructura de preferencia propia, y establecer el significado de esta en base al de las estructuras de preferencia estándar, siguiendo la formulación dada por

Fodor y Roubens [2]. Esto se lleva a cabo en base a la posibilidad de obtener relaciones recíprocas a partir de relaciones débiles mediante transformaciones $F : \mathcal{R} \rightarrow \mathcal{G}$ representables en términos de una función recíproca $f : [0, 1]^2 \rightarrow [0, 1]$, tal que para $R \in \mathcal{R}$ y $a, b \in \mathbb{A}$ se cumple que $F(R)(a, b) = f(R(a, b), R(b, a)) = f(x, y)$. Esta herramienta permite probar que cualquier relación recíproca $G \in \mathcal{G}$ puede descomponerse como $G = P + 0.5(I + J)$, donde $\langle P, I, J \rangle$ denota cualquier estructura de preferencia de una relación débil $R \in \mathcal{R}$ tal que $R \in F_g^{-1}(G)$, donde la transformación $F_g : \mathcal{R} \rightarrow \mathcal{G}$ se representa mediante la función $g(x, y) = \frac{1+x-y}{2}$. Esta descomposición sugiere que cualquier relación recíproca puede expresarse en términos de una componente de preferencia estricta P y otra componente de ausencia de preferencia $L = I + J$, y que por tanto la estructura de preferencia de una relación recíproca $G \in \mathcal{G}$ debe estar compuesta por dos relaciones, P_G y L_G , respectivamente representando esas componentes de preferencia estricta y ausencia de preferencia. Sin embargo, esta descomposición no es única a menos que se cumpla una compatibilidad perfecta entre las estructuras de preferencia de las relaciones débiles y recíprocas implicadas. Las condiciones bajo las que se cumple esta compatibilidad son estudiadas en [1] mediante un enfoque axiomático, que conduce al planteamiento de un sistema de ecuaciones funcionales que solo puede verificarse mediante estructuras compatibles. Esto permite probar que esta compatibilidad solo es posible cuando la semántica de las relaciones débiles se modela mediante una preferencia estricta asimétrica. Además, se muestra también que el sistema de ecuaciones mencionado admite al menos dos soluciones, basadas en diferentes transformaciones entre relaciones débiles y recíprocas. Así, es posible concluir que las relaciones recíprocas pueden admitir diferentes semánticas compatibles con el marco estándar de preferencias borrosas, de modo parecido a cómo una relación de preferencia borrosa débil puede ser dotada de diferentes semánticas a partir de diferentes soluciones a las ecuaciones del modelo estándar.

II. DEFINICIONES Y RESULTADOS

Recordemos en primer lugar que el modelo de preferencias borrosas estándar [2] se establece mediante un enfoque

Este trabajo ha sido parcialmente financiado por la Fundación Carolina (ayudas para estancias cortas postdoctorales), el Gobierno de España (proyecto PGC2018-096509-B-I00) y la Universidad Complutense de Madrid (grupo de investigación 910149).

axiomático que considera la existencia de funciones $p, i, j : [0, 1]^2 \rightarrow [0, 1]$, de manera que la estructura de preferencia $\langle P, I, J \rangle$ de una relación débil $R \in \mathcal{R}$ se obtiene como $P(a, b) = p(x, y)$, $I(a, b) = i(x, y)$ y $J(a, b) = j(a, b)$, donde $x = R(a, b)$ e $y = R(b, a)$, con $a, b \in \mathbb{A}$. Para representar la semántica mínima de preferencias, estas funciones p, i, j deben cumplir el sistema de ecuaciones funcionales

$$x = S_L(p(x, y), i(x, y)), \quad (1)$$

$$1 - y = S_L(p(x, y), j(x, y)). \quad (2)$$

donde S_L denota la t-conorma de Lukasiewicz bajo el automorfismo identidad. Existen múltiples soluciones al sistema anterior, siendo relevante para este trabajo la dada por

$$p(x, y) = \max\{x - y, 0\} \quad (3)$$

$$i(x, y) = \min\{x, y\} \quad (4)$$

$$j(x, y) = \min\{1 - x, 1 - y\} \quad (5)$$

que define una preferencia estricta asimétrica.

Asumiremos además que la representación de transformaciones $F : \mathcal{R} \rightarrow \mathcal{G}$ se lleva a cabo mediante funciones *recíprocas*, que se definen como funciones $f : [0, 1]^2 \rightarrow [0, 1]$ que verifican $f(x, y) + f(y, x) = 1$ para todo $x, y \in [0, 1]$, y son continuas, monótonas no decrecientes en x (y por tanto no crecientes en y) y tales que $f(1, 0) = 1$ y $f(0, 1) = 0$. En [1] se prueban diversos resultados sobre estas funciones recíprocas.

Por otro lado, los principales resultados en [1] que sustentan la mencionada descomposición semántica general de cualquier relación recíproca $G \in \mathcal{G}$ son los siguientes:

Teorema 1. *Sea $\langle p, i, j \rangle$ una solución del sistema de ecuaciones del modelo estándar dado por Ecs.(1) y (2). Entonces, la función $f : [0, 1]^2 \rightarrow [0, 1]$ dada por $f(x, y) = p(x, y) + 0.5(i(x, y) + j(x, y))$ es una función recíproca, y se cumple que $f(x, y) = g(x, y) = \frac{1+x-y}{2}$.*

Corolario 1. *Sea $G \in \mathcal{G}$ una relación recíproca, y sea $R \in \mathcal{R}$ una relación débil tal que $G(a, b) = g(x, y)$ para todo $a, b \in \mathbb{A}$, donde $x = R(a, b)$ e $y = R(b, a)$. Si $\langle P, I, J \rangle$ es una estructura de preferencia asociada a R mediante alguna solución del sistema dado por las Ecs.(1) y (2), entonces $G(a, b) = P(a, b) + 0.5(I(a, b) + J(a, b))$ para todo $a, b \in \mathbb{A}$.*

Así pues, toda relación recíproca puede ser descompuesta en términos de dos componentes, preferencia estricta P y ausencia de preferencia (estricta) $L = I + J$, en tanto I y J recogen las situaciones preferenciales que no pueden ser asociadas a preferencia entre alternativas. Esto lleva a asociar una nueva relación borrosa L a esa componente de ausencia de preferencia, dada por $L(a, b) = S_L(I(a, b), J(a, b))$, para todo $a, b \in \mathbb{A}$. Es fácil comprobar que esta relación L es reflexiva y simétrica. Además, es posible reestablecer la descomposición anterior en términos de disyunción entre preferencia estricta P y ausencia de preferencia L :

Corolario 2. *En las condiciones del Corolario 1, se cumple que $G(a, b) = S_L(P(a, b), 0.5L(a, b))$.*

Estos resultados sugieren que es natural asociar a una relación recíproca $G \in \mathcal{G}$ una estructura de preferencia $\langle P_G, L_G \rangle$. No obstante, en principio la descomposición de G en las relaciones P_G y L_G no es única, en tanto depende de la descomposición de una relación débil $R \in \mathcal{R}$ tal que $G = F_g(R)$ en la estructura $\langle P, I, J \rangle$, que es posible realizar mediante cualquier solución de las Eqs.(1) y (2). Esto además conlleva que la estructura $\langle P_G, L_G \rangle$ no sea necesariamente obtenible a partir de G . Para estudiar la posibilidad de que la estructura $\langle P_G, L_G \rangle$ pueda ser obtenida directamente a partir de G , y que al mismo tiempo sea armónica con la semántica de la estructura de preferencia débil $\langle P, I, J \rangle$ que la dota de significado, en [1] se recurre a un enfoque axiomático que asume la existencia de funciones $p_G, l_G : [0, 1] \rightarrow [0, 1]$ tales que $P_G(a, b) = p_g(u)$ y $L_G(a, b) = l_g(u)$ para todo $a, b \in \mathbb{A}$ y $u = G(a, b)$, con p_G no decreciente y l_G simétrica respecto a 0.5, esto es, $l_G(u) = l_G(1 - u)$, y que conduce a plantear el siguiente sistema de ecuaciones funcionales:

$$S_L(p_R(x, y), i_R(x, y)) = x, \quad (6)$$

$$S_L(p_R(x, y), j_R(x, y)) = 1 - y, \quad (7)$$

$$u = f(x, y), \quad (8)$$

$$p_G(u) = p_R(x, y), \quad (9)$$

$$l_G(u) = S_L(i_R(x, y), j_R(x, y)), \quad (10)$$

$$u = A(p_G(u), l_G(u)), \quad (11)$$

para todo $x, y \in [0, 1]$, y donde $A : \{(p_G(u), l_G(u)) | u \in [0, 1]\} \rightarrow [0, 1]$ es una función continua que permite recuperar la relación recíproca G a partir de su estructura de preferencia. El principal resultado en [1] sobre este sistema es el siguiente:

Teorema 2. *El sistema dado por Ecs.(6)-(11) admite solución si y solo si p_R, i_R, j_R vienen dadas por Ecs.(3)-(5).*

Una solución particular del sistema anterior, cuando la función recíproca empleada es $f = g$, viene dada como:

Teorema 3. *Si una solución del sistema Ecs.(6)-(11) es tal que p_R, i_R, j_R vienen dadas por Ecs.(3)-(5) y $f = g$, entonces $p_G(u) = \max\{2u - 1, 0\}$, $l_G(u) = 1 - |2u - 1|$ y $A(p_G(u), l_G(u)) = S_L(p_G(u), 0.5l_G(u))$ para todo $u \in [0, 1]$.*

En [1] se establecen otros resultados adicionales, como que el sistema anterior admite solución solo si f es estrictamente creciente en x y es invariante en la dirección (1,1) de $[0, 1]^2$, o que la función A es la dada en el teorema anterior si y solo si $f = g$, y se prueba la existencia de al menos otra solución al sistema referido.

REFERENCES

- [1] F. Castiblanco, C. Franco, J. T. Rodríguez and J. Montero (2021) "A characterization of reciprocal fuzzy preference structures and its compatibility with standard fuzzy preference structures," *Fuzzy Sets and Systems*, IN PRESS <https://doi.org/10.1016/j.fss.2021.02.001>.
- [2] J. Fodor, M. Roubens (1994) *Fuzzy Preference Modelling and Multicriteria Decision Support*. Dordrecht: Kluwer Academic Publishers.

Generalización de intervalos en matrices de riesgo imprecisas

María Martínez

Dept. of Statistics and O.R.
University of Oviedo
Oviedo, Spain
UO257879@uniovi.es

Juan Baz

Dept. of Statistics and O.R.
University of Oviedo
Gijón, Spain
bazjuan@uniovi.es

Raul Pérez-Fernández

Dept. of Statistics and O.R.
University of Oviedo
Gijón, Spain
perezfernandez@uniovi.es

Susana Montes

Dept. of Statistics and O.R.
University of Oviedo
Oviedo, Spain
montes@uniovi.es

Resumen—A pesar de que la fiabilidad humana puede tener mucha imprecisión asociada, debido a la incertidumbre natural relacionada con las decisiones tomadas por las personas, una de las principales herramientas en esta área, la matriz de riesgo, se define de manera muy precisa. Esto hace que los expertos no puedan tener ninguna duda a la hora de clasificar un posible error humano. Para evitar este problema, aquí se propone una relajación de la matriz de riesgo, en la que la valoración de los riesgos se dará no solo mediante números, sino también utilizando intervalos o incluso una generalización de estos, que llamaremos cajas. Dicho concepto es analizado en detalle, así como los órdenes obtenidos para el conjunto de cajas.

Index Terms—matriz de riesgo, órdenes entre intervalos, caja, órdenes entre cajas, métodos de elección

I. INTRODUCCIÓN

Las técnicas de análisis de fiabilidad se enfocaron inicialmente en aspectos técnicos del diseño y calidad de la maquinaria. Sin embargo, algunas investigaciones demostraron que el error humano era la causa más común de fallo en muchas situaciones (ver, por ejemplo, [1], [4], [7]). Desde que IEEE publicó un informe sobre fiabilidad humana en 1972, han surgido muchos trabajos en esta temática (ver, por ejemplo, [2], [8], [12], [14], [16], [26]).

Una herramienta fundamental es la matriz de evaluación de riesgos (ver [9]). Esta herramienta facilita la clasificación de diferentes tipos de errores en función de su riesgo asociado. Esta clasificación puede ayudar a priorizar los diferentes tipos de error.

Existen distintas versiones de matrices de evaluación de riesgos, pero en todas ellas los decisores se ven obligados a elegir entre los diferentes niveles de consecuencia/frecuencia, a pesar de que están tomando una decisión subjetiva y, en muchos casos, la elección del nivel apropiado se convierte en una tarea difícil. Para considerar un procedimiento más flexible, proponemos una metodología en la que los decisores pueden expresar valores intermedios. Bajo esta consideración, el valor asociado con un tipo de error ya no es un número sino un intervalo de la línea real o incluso una generalización de este que llamaremos caja. Así, para medir el riesgo asociado a los diferentes tipos de error, es necesario establecer un orden

entre cajas, que generalizará los órdenes considerados entre números e intervalos.

Así, la estructura de este trabajo es la siguiente: la sección 2 ofrece una visión general de los conceptos básicos necesarios para la comprensión del mismo. En la sección 3 introduciremos el concepto de caja, así como distintas formas de ordenarlas, lo que nos permitirá acabar proponiendo una matriz de riesgo imprecisa, que era el objetivo fundamental del trabajo. Para ver la forma de utilizar la metodología propuesta, planteamos un ejemplo de aplicación en la sección 4, donde la información proporcionada por los distintos órdenes es fusionada mediante métodos de ranking. Finalizamos con algunas conclusiones y puntos abiertos.

II. CONCEPTOS BÁSICOS

En esta sección, se recuerdan algunos conceptos básicos sobre matrices de riesgo y órdenes entre intervalos. Sirve además para fijar las notaciones que se van a utilizar a lo largo del trabajo.

A. Matriz de evaluación de riesgos

La matriz de riesgo es un elemento que permite cuantificar los riesgos disminuyendo el nivel de subjetividad en el momento de su evaluación, siempre que la parametrización y asignación de valores a los indicadores esté debidamente fundamentada. Son un instrumento que nace con la necesidad de poder cuantificar los errores y saber así sobre cuáles deberemos actuar con mayor rapidez.

Para obtener los valores de dicha matriz se deberá estudiar cuál es la frecuencia y las consecuencias de cada uno de los errores a examinar. Veamos cuál es su definición formal.

Definición 1: [9], [19] Sea M una matriz de dimensiones $n \times m$, se dice **matriz de riesgo** a la combinación de ciertas consecuencias (asociadas a las columnas de la matriz), ocurriendo en un cierto escenario y con una cierta probabilidad (asociadas a las filas de la matriz), lo cual significa que se necesitan solamente dos entradas para construir la matriz de riesgo.

- Los valores de dicha matriz reciben el nombre de **valores de riesgo** y quedarán determinados por los ejes de la matriz de riesgo, es decir por las consecuencias y la probabilidad de las mismas.

- Las consecuencias, la probabilidad y los valores de riesgo pueden estar divididos en diferentes niveles, según el caso particular con el que estudiemos dicha matriz.
- El proceso del cálculo de los valores de riesgo viene presentado por la implicación lógica: si la probabilidad es p y consecuencia es c , entonces el riesgo es r .

A pesar de esta definición común, la estructura de la matriz de riesgo es muy variable, pudiendo variar tanto sus dimensiones, como la posible nomenclatura. En el caso particular de matriz de riesgo aquí considerada (ver [18]), el eje de la frecuencia lo podremos categorizar como improbable (I), poco probable (Pp), posible (Pos), probable (P) y muy probable (Mp) y el eje de las consecuencias lo podremos clasificar como poco (P), normal (N), grande (G), elevada (E) y extrema (Ex). Las filas *Frec.* indican la frecuencia o probabilidad de ocurrencia y las columnas *Consec.* indican las consecuencias. Así pues, la matriz de riesgo estándar que vamos a considerar es la mostrada en la tabla I.

TABLE I
MATRIZ DE RIESGO ESTÁNDAR.

Consec.	Poco	Normal	Grande	Elevada	Extrema
Frec.					
Muy probable	5	10	15	20	25
Probable	4	8	12	16	20
Posible	3	6	9	12	15
Poco probable	2	4	6	8	10
Improbable	1	2	3	4	5

Típicamente se asigna un código de colores a los valores de la matriz de riesgos, según cual sea su valor. En concreto, se considera el color verde para valores de riesgo menores o igual que 3, el amarillo para valores mayores que 3 y menores o iguales que 7, el color naranja para valores mayores que 7 pero no mayores que 12 y el rojo para el resto. Es evidente que la zona con mayor riesgo se representa de color rojo. En esta zona están los problemas graves, aquellos sobre los que debemos encontrar una rápida solución. La zona naranja muestra los problemas que tienen una importancia elevada. La zona media, de color amarilla, es aquella donde se encuentran los problemas que tienen menos relevancia pero aún así son importantes. La zona de bajo riesgo, en color verde, contiene los riesgos controlados. Es habitual que no se actúe sobre los problemas que están en esta área.

B. Ordenación de intervalos en $\mathbb{R}$

En la subsección anterior hemos visto que los errores se cuantifican mediante números reales entre 1 y 25. No obstante, a la hora de flexibilizar esta matriz para recoger mejor la incertidumbre asociada a la toma de decisiones, nos van a aparecer en la misma intervalos. Recordemos que dados $x_1, x_2 \in \mathbb{R}$ con $x_1 \leq x_2$, llamaremos intervalo (cerrado) al conjunto definido como $X = [x_1, x_2] = \{y \in \mathbb{R} / x_1 \leq y \leq x_2\}$, denotando por $I(\mathbb{R})$ el conjunto de todos los intervalos cerrados en el conjunto de los números reales, es decir, $I(\mathbb{R}) = \{[x_1, x_2] / x_1, x_2 \in \mathbb{R} \text{ con } x_1 \leq x_2\}$.

Es sabido que si bien los números reales forman un conjunto totalmente ordenado con el orden usual ($a < b$ si $b - a$ es un número positivo), encontrar un orden total en el conjunto de los intervalos no es un problema tan inmediato. Así, en la literatura se han introducido diversas formas de comparar intervalos. En concreto, algunos ejemplos son los siguientes: orden débil [5], Maximin [22], [24], Maximax [21], Dominancia de intervalos [13], Hurwicz [17] y el orden producto (también llamado *lattice order* en inglés) [15]. De todos ellos, solo el orden producto es un verdadero orden. Recordemos que viene definido por: $X \leq_{LO} Y \Leftrightarrow x_1 \leq y_1 \text{ y } x_2 \leq y_2$. No obstante, y a pesar de ser un orden muy natural y habitual entre intervalos, es evidente que considerando dicho orden hay intervalos incomparables. Para evitar este problema, buscamos órdenes que sean completos y refinan el orden producto. Esto es posible gracias al concepto de orden admisible (orden total que refina el orden producto) introducido por Bustince et al. (ver [6]).

Dados $X, Y \in I(\mathbb{R})$ con $X = [x_1, x_2]$ e $Y = [y_1, y_2]$, son órdenes admisibles sobre $I(\mathbb{R})$: 1) Orden lexicográfico 1: $X \leq_{Lex1} Y$ si y solo si $x_1 < y_1$ o $(x_1 = y_1 \text{ y } x_2 \leq y_2)$; 2) Orden lexicográfico 2: $X \leq_{Lex2} Y$ si y solo si $x_2 < y_2$ o $(x_2 = y_2 \text{ y } x_1 \leq y_1)$ y 3) Orden de Xu y Yager ([25]): $X \leq_{XY} Y$ si y solo si $x_1 + x_2 < y_1 + y_2$ o $(x_1 + x_2 = y_1 + y_2 \text{ y } x_2 - x_1 \leq y_2 - y_1)$.

III. MATRIZ DE RIESGO IMPRECISA

En una matriz de riesgo clásica, el decisor tiene que decidir por una frecuencia poco probable o posible y no se permite una opción intermedia. Sin embargo, podría ser muy natural una clasificación del riesgo entre poco probable y posible. En ese caso, podría ser lógico considerar que el valor de riesgo es un número entre 8 y 12 si la consecuencia es elevada. Por lo tanto, el valor de riesgo podría venir dado por el intervalo cerrado $[8, 12]$.

Puede ocurrir una situación similar con respecto a la consecuencia. Por tanto, podríamos considerar situaciones intermedias y llegamos a una propuesta de matriz de riesgo 9×9 , que permite opciones intermedias para la frecuencia y la consecuencia, tal como la representada en la tabla II.

TABLE II
PRIMERA APROXIMACIÓN A LA MATRIZ DE RIESGO IMPRECISA.

Consec.	P	PN	N	NG	G	GE	E	EEx	Ex
Frec.									
MP	5	[5,10]	10	[10,15]	15	[15,20]	20	[20,25]	25
MPP	[4,5]	[8,10]	[8,12]	[12,15]	[12,16]	[16,20]	[16,20]	[20,25]	[20,25]
P	4	[4,8]	8	[8,12]	12	[12,16]	16	[16,20]	20
PPos	[3,4]	[6,8]	[6,9]	[9,12]	[9,12]	[12,16]	[12,16]	[16,20]	[15,20]
Pos	3	[3,6]	6	[6,9]	9	[9,12]	12	[12,15]	15
PosPP	[2,3]	[4,6]	[4,6]	[6,9]	[6,9]	[8,12]	[8,12]	[10,15]	[10,15]
PP	2	[2,4]	4	[4,6]	6	[6,8]	8	[8,10]	10
PPI	[1,2]	[3,4]	[3,4]	[3,6]	[3,6]	[4,8]	[4,8]	[5,10]	[5,10]
I	1	[1,2]	2	[2,3]	3	[3,4]	4	[4,5]	5

En la tabla II, se ha denotado con las letras correspondientes a las dos clases que la delimitan a cualquier categoría intermedia. Por ejemplo, si la frecuencia es entre poco probable y posible, se representa por PosPp, si la consecuencia es entre elevada y extrema, se denota por EEx, etc.

En el caso de que exista imprecisión en la consecuencia, pero la frecuencia sea fija, ya tendríamos resuelto el problema

de clasificar un riesgo, tal como hemos visto en la tabla II. En este caso el valor de riesgo sería un intervalo. Lo mismo ocurre si la consecuencia es fija y la imprecisión está en la frecuencia. Puesto que todo número x se puede identificar con el intervalo $[x, x]$, podríamos considerar los órdenes entre intervalos vistos en la sección anterior para ordenar los distintos errores según sus valores de riesgo. El problema viene cuando tanto la valoración de las consecuencias, como la de las frecuencias, es imprecisa. En tal caso no hemos podido valorar el riesgo, tal como se ha visto en la tabla II. Vamos a centrarnos en una zona determinada de la misma, para ir analizando la solución dada para este problema. En concreto nos centramos en la zona correspondiente al caso de consecuencias entre elevadas y extremas y de frecuencias entre poco probable y posible, la cual aparece ampliada en la tabla III.

TABLE III
ZOOM DE LA MATRIZ DE RIESGO.

Frecuencia	Pos PosPp Pp	Consecuencia		
		E	EEx	Ex
		12	[12, 15]	15
		[8, 12]	?	[10, 15]
		8	[8, 10]	10

Para la celda restante, podríamos pensar en considerar nuevamente un intervalo. Sin embargo, lo que tiene que ocurrir es que el valor debe estar entre $[8, 10]$ y $[12, 15]$ y además entre $[8, 12]$ y $[10, 15]$. Para recoger esta idea, aparece un nuevo concepto como intersección de dos “intervalos” de intervalos.

Definición 2: Sean a, b, c, d cuatro números reales con $a \leq b \leq c \leq d$. El conjunto $\{[x, y] \in I(\mathbb{R}) / [a, c] \leq_{LO} [x, y] \leq_{LO} [b, d], [a, b] \leq_{LO} [x, y] \leq_{LO} [c, d]\}$ se llama caja con extremos a, b, c y d y se denota por $[a, b, c, d]$.

Por lo tanto, está claro que cualquier elemento en la caja $[a, b, c, d]$ es un intervalo tal que está entre el intervalo $[a, b]$ y el intervalo $[c, d]$ y, al mismo tiempo, está entre el intervalo $[a, c]$ y el intervalo $[b, d]$ con respecto al orden producto. Denotaremos por $\mathcal{C}(\mathbb{R})$ al conjunto de todas las cajas en $\mathbb{R}$, es decir, $\mathcal{C}(\mathbb{R}) = \{[a, b, c, d] / a, b, c, d \in \mathbb{R}, a \leq b \leq c \leq d\}$. A continuación, presentamos una definición equivalente del concepto de caja.

Proposición 1: Sea $[a, b, c, d]$ en $\mathcal{C}(\mathbb{R})$. Se tiene que:

$$[a, b, c, d] = \{[x, y] \in I(\mathbb{R}) / a \leq x \leq b, c \leq y \leq d\}$$

De esta definición equivalente se puede deducir que una caja es un tipo particular de hiperrectángulo bidimensional o 2-ortótropo (ver [10]). Por otro lado, considerando ahora la identificación entre la caja $[a, b, c, d]$ y la 4-tupla (a, b, c, d) , el concepto de caja coincide con el de intervalo 4-dimensional, tal como ha sido definido por Bedregal et al. en [3]. En dicho trabajo se pone además de manifiesto que el conjunto de los intervalos 2-dimensionales coincide con el conjunto de los intervalos cerrados y el de los intervalos 1-dimensionales con el de los números reales. Así pues, todos los valores que tenemos hasta el momento en la matriz de riesgo (ver tabla II) puede considerarse que son cajas. Además, podríamos considerar cajas adecuadas para llenar los espacios vacíos

en esta matriz. Por ejemplo, el valor restante en el ejemplo considerado en la tabla III sería $[8, 10, 12, 15]$. Podríamos repetir este procedimiento a lo largo de todas las celdas vacías de la matriz de la tabla II. Sin embargo, para poder colorear esta matriz, debemos definir un orden apropiado en $\mathcal{C}(\mathbb{R})$. Una primera propuesta podría ser la siguiente.

Definición 3: Sea $[a, b, c, d]$ y $[\bar{a}, \bar{b}, \bar{c}, \bar{d}]$ en $\mathcal{C}(\mathbb{R})$. Diremos que $[a, b, c, d]$ es menor o igual que $[\bar{a}, \bar{b}, \bar{c}, \bar{d}]$ con respecto al orden producto, y se denota por $[a, b, c, d] \leq_{LO} [\bar{a}, \bar{b}, \bar{c}, \bar{d}]$, si y solo si $a \leq \bar{a}$, $b \leq \bar{b}$, $c \leq \bar{c}$ y $d \leq \bar{d}$.

Es fácil probar que $\leq_{LO}$ es un orden en $\mathcal{C}(\mathbb{R})$, pero no es una relación completa. Por ejemplo, $[2, 3, 4, 5] \not\leq_{LO} [2, 2, 5, 5]$ y $[2, 2, 5, 5] \not\leq_{LO} [2, 3, 4, 5]$. Como necesitamos ordenar todos los elementos en la matriz de riesgo, necesitamos definir un orden total en $\mathcal{C}(\mathbb{R})$. A partir de las ideas de los órdenes lexicográficos en $L(\mathbb{R})$ podemos obtener 24 órdenes distintos en $\mathcal{C}(\mathbb{R})$. Comenzamos con el de tipo 1, que nos servirá para generar el resto.

Definición 4: Sea $[a, b, c, d]$ y $[\bar{a}, \bar{b}, \bar{c}, \bar{d}]$ en $\mathcal{C}(\mathbb{R})$. Diremos que $[a, b, c, d]$ es menor o igual que $[\bar{a}, \bar{b}, \bar{c}, \bar{d}]$ con respecto al orden lexicográfico tipo 1, que se denota por $[a, b, c, d] \leq_{1234} [\bar{a}, \bar{b}, \bar{c}, \bar{d}]$, si y solo si se cumple una de las condiciones siguientes: $a < \bar{a}$ o ($a = \bar{a}$, $b < \bar{b}$) o ($a = \bar{a}$, $b = \bar{b}$, $c < \bar{c}$) o ($a = \bar{a}$, $b = \bar{b}$, $c = \bar{c}$, $d < \bar{d}$).

Se puede probar que la relación en $\mathcal{C}(\mathbb{R})$ es un orden total que refina el orden producto.

Puesto que en dicha definición no hemos considerado las relaciones entre los cuatro elementos que describen una caja y cualquier caja $[a, b, c, d]$ se puede identificar con una 4-tupla (a, b, c, d) en $\mathbb{R}^4$, el orden lexicográfico 1 en $\mathcal{C}(\mathbb{R})$ se podría generalizar de forma inmediata a $\mathbb{R}^4$ y también en ese espacio sería un orden lineal. De hecho, vamos a utilizarlo como generador de órdenes lineales en ese espacio, tal como puede verse en el siguiente resultado.

Proposición 2: Sea $\mathcal{A}$ una matriz cualquiera de rango completo con $\mathcal{A} \in \mathcal{M}_{4 \times 4}$. Se tiene que la relación binaria $\mathcal{R}$ definida sobre $\mathbb{R}^4 \times \mathbb{R}^4$ como $(a, b, c, d) \mathcal{R} (\bar{a}, \bar{b}, \bar{c}, \bar{d})$ si y solo si $\mathcal{A}(a, b, c, d)' \leq_{1234} \mathcal{A}(\bar{a}, \bar{b}, \bar{c}, \bar{d})'$ para cualesquiera $(a, b, c, d), (\bar{a}, \bar{b}, \bar{c}, \bar{d}) \in \mathbb{R}^4$, es un orden total en $\mathbb{R}^4$.

Este resultado puede considerarse una generalización inmediata del demostrado para el intervalo $[0, 1]$ en [11]. Además es inmediato que cualquier orden generado de esta forma, puede ser restringido al conjunto de cajas $\mathcal{C}(\mathbb{R})$. En concreto, es evidente que según vayamos considerando las 24 matrices de rango completo con un 1 en cada fila y el resto de los elementos nulos, podríamos generar 24 órdenes lineales distintos sobre $\mathcal{C}(\mathbb{R})$.

Corolario 1: Sea $\{i, j, k, l\}$ una permutación cualquiera del conjunto $\{1, 2, 3, 4\}$ y sea $\mathcal{A}_{ijkl}$ una matriz en $\mathcal{M}_{4 \times 4}$ tal que todos sus elementos son nulos salvo $a_{1i} = a_{2j} = a_{3k} = a_{4l} = 1$. La relación $\leq_{ijkl}$ definida como $[a, b, c, d] \leq_{ijkl} [\bar{a}, \bar{b}, \bar{c}, \bar{d}]$ si y solo si $\mathcal{A}_{ijkl}(a, b, c, d)' \leq_{1234} \mathcal{A}_{ijkl}(\bar{a}, \bar{b}, \bar{c}, \bar{d})'$, para cualesquiera $[a, b, c, d], [\bar{a}, \bar{b}, \bar{c}, \bar{d}] \in \mathcal{C}(\mathbb{R})$, es un orden total en $\mathcal{C}(\mathbb{R})$.

Si aplicamos la proposición 2 para la matriz

$$\mathcal{A} = \begin{pmatrix} -1 & 1 & 1 & 1 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \end{pmatrix}$$

obtenemos la generalización del orden de Xu-Yager para cajas.

Proposición 3: Dadas $[a, b, c, d]$ y $[\bar{a}, \bar{b}, \bar{c}, \bar{d}]$ dos cajas en $\mathcal{C}(\mathbb{R})$. La relación $\leq_{XY}$ sobre $\mathcal{C}(\mathbb{R})$ definida como: $[a, b, c, d] \leq_{XY} [\bar{a}, \bar{b}, \bar{c}, \bar{d}]$ si y sólo si verifica alguna de las siguientes condiciones:

- $a + b + c + d < \bar{a} + \bar{b} + \bar{c} + \bar{d}$
- $a + b + c + d = \bar{a} + \bar{b} + \bar{c} + \bar{d}$ y $(b - a) + (d - c) < (\bar{b} - \bar{a}) + (\bar{d} - \bar{c})$
- $a + b + c + d = \bar{a} + \bar{b} + \bar{c} + \bar{d}$, $(b - a) + (d - c) = (\bar{b} - \bar{a}) + (\bar{d} - \bar{c})$ y $a < \bar{a}$
- $a + b + c + d = \bar{a} + \bar{b} + \bar{c} + \bar{d}$, $(b - a) + (d - c) = (\bar{b} - \bar{a}) + (\bar{d} - \bar{c})$, $a = \bar{a}$ y $c \leq \bar{c}$

es un orden total que refina el orden producto.

También se puede probar que si nos restringimos a intervalos estos órdenes corresponden con sus homónimos en $L(\mathbb{R})$ y si nos restringimos a números, con el orden habitual sobre la recta real. Evidentemente una caja que es menor que otra con respecto a un orden, puede no serlo con respecto al otro. Si consideramos el orden lexicográfico tipo 1 y los colores verde ($\leq_{1234} 3$), amarillo ($>_{1234} 3$ y $\leq_{1234} 7$), naranja ($>_{1234} 7$ y $\leq_{1234} 12$) y rojo ($>_{1234} 12$), obtenemos la matriz de riesgo imprecisa que puede verse en la tabla IV.

TABLE IV
MATRIZ DE RIESGO IMPRECISA.

Cons.	P	PN	N	NG	G	GE	E	EEx	Ex
Cons.	5	[5,10]	10	[10,15]	15	[15,20]	20	[20,25]	25
MpP	[4,5]	[4,5,8,10]	[8,10]	[8,10,12,15]	[12,15]	[12,15,16,20]	[16,20]	[16,20,20,25]	[20,25]
P	4	[4,8]	8	[8,12]	12	[12,16]	16	[16,20]	20
PPos	[3,4]	[3,4,6,8]	[6,8]	[6,8,9,12]	[9,12]	[9,12,12,16]	[12,16]	[12,16,16,20]	[16,20]
Pos	[3,6]	[3,6]	6	[6,9]	9	[9,12]	12	[12,15]	15
PosPp	[2,3]	[2,3,4,6]	[4,6]	[4,6,6,9]	[6,9]	[6,9,9,12]	[8,12]	[8,10,12,15]	[10,15]
Pp	2	[2,4]	4	[4,6]	6	[6,8]	8	[8,10]	10
PpI	[1,2]	[1,2,4]	[2,4]	[2,4,6,9]	[2,6]	[2,4,6,8]	[4,8]	[4,8,8,10]	[8,10]
I	1	[1,2]	2	[2,3]	3	[3,4]	4	[4,5]	5

Si prefiriésemos trabajar con todos los órdenes a la vez, combinándolos, podríamos utilizar métodos de elección como el de Borda o Concorcet. Vamos a ver el procedimiento completo mediante un ejemplo.

IV. CASO DE ESTUDIO

Veremos de qué manera podemos aplicar estas formas de ordenar los distintos tipos de errores en un ejemplo ilustrativo (ver [23]). El objetivo es mostrar de qué forma quedan clasificados los errores según el orden que utilicemos. Así, podremos aplicar toda la teoría que hemos desarrollado para poder ver cuál es su utilidad real. En particular, estudiaremos y clasificaremos los errores más comunes que puede cometer un conductor de autobús. Podemos considerar los siguientes eventos:

- E1. Cruzar un semáforo en ámbar.
- E2. Cruzar un semáforo en rojo.
- E3. Conducir bajo los efectos de alcohol o de drogas.
- E4. No poner el intermitente al girar.
- E5. Saludar a otros conductores.
- E6. Confundirse en el cambio más de 50 céntimos.

- E7. Confundirse en el cambio menos de 50 céntimos.

Ahora, deberemos ver cuál es la frecuencia con la que se produce cada uno de los eventos, así como ver cuáles son las consecuencias de los mismos. Las frecuencias y las consecuencias de cada error se pueden determinar, por ejemplo, a partir de estadísticas que relacionen la asiduidad y el impacto que han tenido a lo largo de cierto periodo de tiempo. En base a esos datos, se podrían clasificar en las distintas categorías de las frecuencias y consecuencias (con la posibilidad de considerar intervalos).

En la tabla V aparece resumida toda la información: la etiqueta que relaciona cada error, con unas frecuencias y consecuencias que se han considerado razonables.

TABLE V
ESTIMACIÓN DE LAS FRECUENCIAS Y DE LAS CONSECUENCIAS DE CADA EVENTO.

Evento	Frecuencias	Consecuencias
E1	MpP	GE
E2	Pos	Ex
E3	PpI	Ex
E4	PosPp	EEx
E5	Mp	P
E6	PpI	PN
E7	PosPp	P

La notación utilizada es la misma que se ha utilizado en la matriz de riesgo imprecisa de la tabla IV.

Una vez obtenida la información de la tabla V, debemos considerar dicha matriz de riesgo para asignar a cada uno de los eventos (errores a cuantificar) el valor de riesgo asociado a los mismos, que vendrá dado mediante una caja, recordando que los números y los intervalos pueden ser vistos como tipos particulares de cajas. Así, dichas valoraciones son:

- E1. Cruzar un semáforo en ámbar: $\mathcal{C}_{E1} = [12, 15, 16, 20]$.
- E2. Cruzar un semáforo en rojo: $\mathcal{C}_{E2} = 15$.
- E3. Conducir bajo los efectos de alcohol o de drogas: $\mathcal{C}_{E3} = [5, 10]$.
- E4. No poner el intermitente al girar: $\mathcal{C}_{E4} = [8, 10, 12, 15]$.
- E5. Saludar a otros conductores: $\mathcal{C}_{E5} = 5$.
- E6. Confundirse en el cambio más de 50 céntimos: $\mathcal{C}_{E6} = [1, 2, 2, 4]$.
- E7. Confundirse en el cambio menos de 50 céntimos: $\mathcal{C}_{E7} = [2, 3]$.

Si aplicamos cada uno de los 24 órdenes lexicográficos que hemos introducido, la ordenación entre los distintos errores se repite en muchos de ellos, tal como puede verse en la tabla VI.

Nótese que tenemos solamente tres ordenaciones posibles. Aún así, se observa como la elección del orden entre cajas hace que la clasificación de los errores varíe de un caso a otro. Si queremos fusionar la información de varios órdenes, para no dar más relevancia a unas componentes de la caja que a otras, podríamos fusionar las distintas ordenaciones mediante métodos de elección. En particular, estudiaremos los métodos de Borda y de Condorcet. El primero se basa en la posición y el

TABLE VI
ORDENACIONES OBTENIDAS SEGÚN LOS DISTINTOS ÓRDENES
LEXICOGRÁFICOS.

Frec.	Ranking
8	$\mathcal{C}_{E1} \geq \mathcal{C}_{E2} \geq \mathcal{C}_{E4} \geq \mathcal{C}_{E3} \geq \mathcal{C}_{E5} \geq \mathcal{C}_{E7} \geq \mathcal{C}_{E6}$
8	$\mathcal{C}_{E1} \geq \mathcal{C}_{E2} \geq \mathcal{C}_{E4} \geq \mathcal{C}_{E3} \geq \mathcal{C}_{E5} \geq \mathcal{C}_{E6} \geq \mathcal{C}_{E7}$
8	$\mathcal{C}_{E2} \geq \mathcal{C}_{E1} \geq \mathcal{C}_{E4} \geq \mathcal{C}_{E3} \geq \mathcal{C}_{E5} \geq \mathcal{C}_{E7} \geq \mathcal{C}_{E6}$

segundo en la comparación por pares. En la mayoría de casos ambos criterios actúan de igual forma, aunque puede haber casos en los que la ordenación sea completamente contraria. Veremos a continuación qué es lo que sucede en este ejemplo.

A. Ranking Borda Count

Existen distintos rankings para clasificar órdenes, en este caso nos centraremos en el Ranking de Borda Count.

Para hacer una clasificación de todos estos órdenes que hemos señalado previamente, usaremos dicho ranking para fusionar la información dada por todos los órdenes considerados.

En este método (ver [20] para más información al respecto), basta mirar cuál es la frecuencia de cada una de las ordenaciones posibles. La representación usual es la que aparece en la matriz $\mathcal{O}$, llamada **matriz de votación**, en la que el elemento o_{ij} representa el número de veces que $a_i \geq a_j$. A partir de la matriz $\mathcal{O}$, bastará sumar por filas obteniendo un valor α_i . Para obtener el ranking final basta hacer una clasificación natural con los valores α_i obtenidos.

La matriz de votación asociada a este ejemplo puede verse en la tabla VII.

TABLE VII
MATRIZ DE VOTACIÓN.

$\mathcal{O}$	$\mathcal{C}_{E1}$	$\mathcal{C}_{E2}$	$\mathcal{C}_{E3}$	$\mathcal{C}_{E4}$	$\mathcal{C}_{E5}$	$\mathcal{C}_{E6}$	$\mathcal{C}_{E7}$
$\mathcal{C}_{E1}$	0	16	24	24	24	24	24
$\mathcal{C}_{E2}$	8	0	24	24	24	24	24
$\mathcal{C}_{E3}$	0	0	0	0	24	24	24
$\mathcal{C}_{E4}$	0	0	24	0	24	24	24
$\mathcal{C}_{E5}$	0	0	0	0	0	24	24
$\mathcal{C}_{E6}$	0	0	0	0	0	0	8
$\mathcal{C}_{E7}$	0	0	0	0	0	16	0

La puntuación que nos ofrece el ranking de Borda se denota por

$$\alpha_{E_i} = \sum_{j=1}^n o_{ij}$$

donde el elemento o_{ij} representa la i -ésima fila y la j -ésima columna de la matriz $\mathcal{O}$.

Así, las puntuaciones de Borda serían las siguientes: $\alpha_{E1} = 136$ y $\alpha_{E2} = 128$, $\alpha_{E3} = 72$, $\alpha_{E4} = 96$, $\alpha_{E5} = 48$, $\alpha_{E6} = 8$, $\alpha_{E7} = 16$. Se obtendría así el ranking según el Ranking de Borda Count por medio de la ordenación de estos números de menor a mayor:

$$\mathcal{C}_{E6} \leq \mathcal{C}_{E7} \leq \mathcal{C}_{E5} \leq \mathcal{C}_{E3} \leq \mathcal{C}_{E4} \leq \mathcal{C}_{E2} \leq \mathcal{C}_{E1}$$

Vemos en la figura 1 cómo queda la representación gráfica si usamos un diagrama de Hasse.

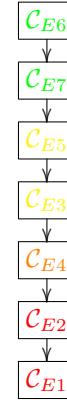


Fig. 1. Órdenes lexicográficos (ordenación por Borda).

Observamos como el diagrama tiene un aspecto lineal. El evento que tiene una mayor importancia y que domina a todos los demás es el evento que queda en la parte inferior, en este caso es el evento $\mathcal{C}_{E1}$, como ya hemos visto.

La interpretación real de todo lo que se ha comentado es que, si fuéramos propietarios de esta empresa de autobuses donde se producen estos errores, el primer error en el que nos deberíamos centrar en solucionar sería el $\mathcal{C}_{E1}$. Este error se corresponde con cruzar un semáforo en ambar. Habrá distintas formas de enfocar este problema y de intentar encontrar una solución, como el hecho de realizar cursos de concienciación, intentar disminuir las horas que un conductor está conduciendo el autobús de continuo, etc. Se pueden probar distintos métodos, y una vez que la frecuencia con la que se produce este error disminuya de manera considerable, podemos intentar solucionar el siguiente evento, que en este caso es $\mathcal{C}_{E2}$. Se procedería así sucesivamente hasta conseguir que todos los eventos tengan un riesgo bajo o, si esto no es posible, conseguir disminuirlo de forma considerable.

B. Ranking de Condorcet

En muchas situaciones reales habituales es difícil que cuando obtenemos distintos rankings y el objetivo es obtener un ranking consenso se tenga un ranking que aparezca más de la mitad de las ocasiones. Para solucionar este problema, Condorcet se apoya en el concepto de ganador por mayoría. Condorcet propone un ganador basado en la dominancia del mismo sobre el resto de candidatos. Si tal candidato existe, será elegido el ganador.

Para ver cuál es el ganador de Condorcet (para más información, ver de nuevo [20]), debemos buscar cuál es el candidato que comparado con el resto es el preferido por el mayor número de votantes. En este ejemplo, deberemos mirar cuál es el evento que tiene mayor importancia en la mayoría de los órdenes lexicográficos.

Se puede observar que el evento $\mathcal{C}_{E1}$ es elegido como primero un total de 16 veces sobre 24 posibles clasificaciones. Como 16 es más de la mitad del número de ordenaciones, este candidato será también preferido por pares a todos los demás.

Además en este caso, el único evento que le hace una competencia real a C_{E1} es C_{E2} , ya que este es el único que le supera en alguna ocasión (lo hace en las 8 restantes). Así el primer evento sobre el que deberíamos prestar atención es el C_{E1} . Para ver qué sucede con el resto de la clasificación, debemos seguir mirando qué sucede con el resto de pares. El segundo evento sobre el que nos deberíamos fijar es sobre C_{E2} ya que domina al resto de errores (salvo C_{E1} , como hemos mencionado anteriormente) de manera absoluta. De igual manera C_{E4} domina al resto de los eventos (salvo a C_{E1} y a C_{E2}), estamos en la misma situación con los eventos C_{E3} y C_{E5} . Todos ellos tienen dominancia sobre los eventos siguientes, C_{E3} tiene dominancia sobre C_{E5} , C_{E6} y C_{E7} y C_{E5} tiene dominancia sobre los eventos C_{E6} y C_{E7} . Finalmente, C_{E7} tiene dominancia sobre el evento C_{E6} .

Así, la clasificación según Condorcet quedaría de la siguiente manera: $C_{E6} \leq C_{E7} \leq C_{E5} \leq C_{E3} \leq C_{E4} \leq C_{E2} \leq C_{E1}$. En este caso, la clasificación de Borda y de Condorcet coinciden, lo cual ya habíamos comentado que suele ser habitual. Como ambos criterios coinciden, tanto el diagrama de Hasse, como la interpretación que podemos hacer en este caso es exactamente la misma que ocurría cuando utilizábamos el método de clasificación de Borda. Nótese que el método de Condorcet no siempre da lugar a una ordenación de los eventos (puede haber ciclos).

C. Orden de Xu-Yager

Por otro lado, si consideramos la generalización del orden de Xu-Yager, obtenemos la siguiente clasificación: $C_{E1} \geq_{XY} C_{E2} \geq_{XY} C_{E4} \geq_{XY} C_{E3} \geq_{XY} C_{E5} \geq_{XY} C_{E7} \geq_{XY} C_{E6}$.

En este caso, sucede lo contrario que cuando estudiamos los órdenes lexicográficos, tenemos una única clasificación posible para todos ellos, ya que los dos primeros puntos de la definición de $\leq_{XY}$ son suficientes para determinar la clasificación.

Las conclusiones según este orden coincide que se corresponden con las dadas por la fusión de los órdenes lexicográficos, tanto aplicando Borda, como Concordet.

V. CONCLUSIONES

Hemos introducido un concepto que generaliza al de intervalo y nos permite definir una matriz de riesgo bajo imprecisión, además de analizar distintas estructuras de orden sobre el conjunto de cajas. La nueva matriz generaliza a la usual, manteniendo coherencia con el significado de las etiquetas y los colores de la misma, pero permitiendo además hacer consideraciones menos estrictas. Los elementos de dicha matriz pueden ser ordenados con distintos órdenes totales y en un ejemplo aplicado vemos como dicha información puede ser fusionada aplicando los métodos de Borda y Concordet, que en este caso llegan a la misma conclusión.

Es inmediato observar que las cuestiones estudiadas plantean la necesidad de resolver algunos problemas bastante directos. Por ejemplo, considerar otro orden en la definición de caja, estudiar otros órdenes en el conjunto de cajas o utilizar otros métodos de elecciones para fusionar la ordenación dada por distintos criterios.

REFERENCIAS

- [1] O.N. Aneziris, I.A. Papazoglou, M.L. Mud, M. Damen, J. Kuiper, H. Baksteen, B.J. Ale, L.J. Bellamy, A.R. Hale, A. Bloemhoff, J.G. Post, and J. Oh. Towards risk assessment for crane activities. *Safety Science*, 46:872–884, 2008.
- [2] G.E. Apostolakis, V.M. Bier, and A. Mosleh. A critique of recent models for human error rate assessment. *Reliab. Eng. Syst. Saf.*, 22:201–217, 1988.
- [3] B. Bedregal, G. Beliakov, H. Bustince, T. Calvo, R. Mesiar, and D. Paternain. A class of fuzzy multisets with a fixed number of memberships. *Information Sciences*, 189:1–17, 2012.
- [4] L.J. Bellamy, T.A.W. Geyer, and J.A. Astley. Evaluation of the human contribution to pipework and in-line equipment failure frequencies. Technical Report 15, HSE Contract Research Report, UK Health and Safety Executive, Bootle, Merseyside, 1989.
- [5] K. Bogart, J. Bonin, and J. Mitás. Interval orders based on weak orders. *Discrete Applied Mathematics*, 60(1):93–98, 1995.
- [6] H. Bustince, J. Fernandez, A. Kolesárová, and R. Mesiar. Generation of linear orders for intervals by means of aggregation functions. *Fuzzy Sets and Systems*, 220:69 – 77, 2013.
- [7] E. Calixto, G.B.A. Lima, and P.R.A. Firmino. Comparing slim, spar-h and bayesian network methodologies. *Open Journal of Safety Science and Technology*, 3(2):31–41, 2013.
- [8] F. Castiglia and M. Giardina. Analysis of operator human errors in hydrogen refuelling stations: comparison between human rate assessment techniques. *Int. J. Hydrog. Energy*, 38:1166–1176, 2013.
- [9] L.A. Cox. What's wrong with risk matrices? *Risk Analysis*, 28:497–512, 2008.
- [10] H.S.M. Coxeter. *Regular Polytopes*. Dover, New York, 1973.
- [11] L. De Miguel, M. Sesma-Sara, M. Elkano, M. Asiain, and H. Bustince. An algorithm for group decision making using n-dimensional fuzzy sets, admissible orders and owa operators. *Information Fusion*, 37:126–131, 2017.
- [12] G. Desmorat, F. Guarnieri, D. Besnard, P. Desideri, and F. Loth. Pouring cream into natural gas: the introduction of common performance conditions into the safety management of gas networks. *Saf. Sci.*, 54:1–7, 2013.
- [13] P. Fishburn. *Interval ordenings*. Wiley, New York, 1987.
- [14] J. Forester, D. Bley, S. Cooper, E. Lois, N. Siu, A. Kolaczowski, and J. Wreathall. Expert elicitation approach for performing atheana quantification. *Reliability Engineering & System Safety*, 83(2):207–220, 2004.
- [15] J.A. Goguen. L-fuzzy sets. *Journal of Mathematical Analysis and Applications*, 18(1):145 – 174, 1967.
- [16] K.M. Groth and L.P. Swiler. Bridging the gap between hra research and hra practice: a bayesian network version of spar-h. *Reliab. Eng. Syst. Saf.*, 115:33–42, 2013.
- [17] L. Hurwicz. A class of criteria for decision-making under ignorance. *Cowles Comision Discussion Paper: Statistics 356*, 1951.
- [18] G.K. Kaya. *Good risk assessment practice in hospitals*. PhD thesis, University of Cambridge, 2018.
- [19] A.S. Markowski and M.S. Mannan. Fuzzy risk matrix. *Journal of hazardous materials*, 159(1):152–157, November 2008.
- [20] R. Pérez-Fernández and B. De Baets. The superdominance relation, the positional winner, and more missing links between borda and condorcet. *Journal of Theoretical Politics*, 31:46–65, 2021.
- [21] J.K. Satia and R.E. Lave. Markovian decision processes with uncertain transition probabilities. *Operations Research*, 21(3):728–740, 1973.
- [22] M. Sniedovich. Wald's maximin model: a treasure in disguise! *Journal of Risk Finance*, 9(3):287–291, 2008.
- [23] Virtual Scientific Meeting EUREKA 2016. *Transforming Risk Assessment Matrices via Receiver Operating Characteristic Curves*, Coahuila, México, 12 2016.
- [24] A. Wald. Statistical decision functions which minimize the maximum risk. *Annals of Mathematics*, 46(2):265–280, 1945.
- [25] Z.S. Xu and R.R. Yager. Some geometric aggregation operators based on intuitionistic fuzzy sets. *International Journal of General Systems*, 35(4):417–433, 2006.
- [26] B. Zimolong. Empirical evaluation of therp, slim and ranking to estimate heps. *Reliab. Eng. Syst. Saf.*, 35:1–11, 1992.

Consenso en Entornos Académicos Virtuales: Toma de Decisiones Grupales en el Trabajo de Clase Online

1st Aitor Manuel Orellana Romero
Dept. de Ingeniería Informática
Universidad de Cádiz
Cádiz, España
aitormanuel.oreromero@alum.uca.es

2nd Antonio Velez-Estevez
Dept. de Ingeniería Informática
Universidad de Cádiz
Cádiz, España
antonio.velez@uca.es

3rd Manuel Jesús Cobo Martín
Dept. de Ingeniería Informática
Universidad de Cádiz
Cádiz, España
manueljesus.cobo@uca.es

4th Ignacio Javier Pérez Gálvez
Dept. de Ingeniería Informática
Universidad de Cádiz
Cádiz, España
ignaciojavier.perez@uca.es

Abstract—En esta contribución se describe la implementación de una plataforma web que puede ser integrada en el campus virtual de cualquier centro educativo con el objetivo de centralizar la gestión de las tareas de toma de decisiones en grupo por parte de los estudiantes. Debido a la situación actual propiciada por el virus COVID-19, muchas clases que antes eran presenciales se están impartiendo de forma “online”. Por tanto, situaciones de toma de decisiones relativas al trabajo en grupo que hasta hace poco se resolvían directamente con un rápido debate en clase, ahora necesitan llevarse a cabo de forma no presencial. Dicha plataforma actúa como moderador virtual, permitiendo a los usuarios comunicarse a través de la web con el objetivo de tomar decisiones consensuadas.

Index Terms—Toma de decisiones en grupo, Consenso, Lógica difusa, Clase online

I. INTRODUCCIÓN

La toma de decisiones es una tarea cotidiana que todos realizamos a diario. Hay decisiones con poca trascendencia como la de escoger si tomar café o zumo en el desayuno, y otras más relevantes como decidir si cambio de trabajo o mantengo mi puesto actual. La actividad docente no está excluida de este tipo de tareas, tanto desde el punto de vista del profesor como desde las tareas intrínsecas del estudiante [14]. Decidir qué día de la semana hacer una tarea, la fecha de un examen o escoger el tema de un trabajo de entre un conjunto de posibilidades, son decisiones que se toman frecuentemente en la universidad.

Asimismo, cuando se trabaja en grupo y la decisión afecta a varias personas, la situación empieza a ser más compleja. Este caso es conocido como el problema de la toma de decisiones en grupo. Resolver este problema de forma presencial suele ser relativamente sencillo; se discuten conjuntamente las diferentes alternativas para intentar llegar a una solución conjunta que sea “lo mejor para todos”. Sin embargo, analizándolo bien, se pueden dar situaciones en las que cada persona

perteneciente al grupo tenga unos intereses o motivaciones diferentes, incluso que dé más importancia a unos criterios de elección diferentes al resto de compañeros [9].

Cuando los alumnos se encuentran en una de estas situaciones, sumando a la ecuación que el entorno de estudio no sea presencial, necesitan herramientas o entornos virtuales que les ayuden con el proceso de toma de decisiones en grupo, de forma que no se paralice el trabajo sencillamente porque no hay consenso en alguna decisión que afecta al grupo en su conjunto. Por el momento, los estudiantes disponen de herramientas que les facilitan la comunicación, como foros, chats, o incluso wikis que fomentan el trabajo colaborativo. Sin embargo, no tienen ningún soporte para ayudarles a llegar a acuerdos grupales.

Cuando hablamos de toma de decisiones en grupo, hay varias formas de llegar a un acuerdo [5]. Desde un punto de vista más dominante, en el que habría un líder que toma la decisión y todos los demás le hacen caso (casuística poco recomendable en el entorno docente), hasta modelos más participativos en los que se hace lo que dice la mayoría (usando un sistema de votación en el que la alternativa que tenga más votos será la elegida como decisión del grupo). Sin embargo, tanto la primera como la segunda opción pueden dejar a algunos miembros del equipo sin sentirse cómodos con la decisión tomada, creando diferentes bandos en el grupo de trabajo. Lo ideal sería alcanzar algún tipo de consenso [2].

Dicho esto, el objetivo general de este proyecto es el de implementar un modelo de consenso usando la lógica difusa [1], [11], [12], [22], en el que todos los miembros del grupo se sientan igualmente representados por la decisión tomada. Para ello, no trataremos el concepto de consenso como un concepto “crisp” binario (hay consenso total o no lo hay), si no que se establecerá una medida del nivel de consenso alcanzado por el grupo en base a sus preferencias individuales (valores entre

cero y uno), y se establecerá un umbral para considerar si hay consenso suficiente o la negociación debe continuar [2], [16], [17]. De esta forma, usamos el término “Soft consensus” para referirnos a que el concepto de consenso puede tener diferentes grados y para poder medir el nivel de consenso que existe en cada momento en los procesos de toma de decisiones.

Para guiar la negociación de forma que el proceso sea convergente (que cada vez las posiciones de los estudiantes estén más cerca), se establecerá un mecanismo de consenso, que actuará como moderador virtual enviando un mensaje personalizado a cada uno de los miembros del grupo para que relaje un poco sus preferencias individuales en la dirección que marque el grupo en su conjunto.

De esta forma, tendremos dos procesos diferenciados que actúan conjuntamente de forma secuencial. El proceso de consenso actúa para lograr alcanzar el máximo grado de consenso posible entre las opiniones de los usuarios. Cuando todos han expresado sus opiniones, el sistema calcula el grado de consenso existente. Si es satisfactorio, entonces se aplica el proceso de selección de cara a obtener la solución final. Por el contrario, si el grado de consenso medio no es satisfactorio, entonces el sistema insta a los usuarios a modificar sus opiniones de cara a aumentar la proximidad en sus preferencias. De esta manera, un proceso de toma de decisiones puede verse como un proceso dinámico e iterativo en el que los usuarios van acercando sus posiciones hasta maximizar el consenso.

En la plataforma quedarán registradas todas las interacciones de los alumnos con el sistema, de forma que el profesor dispondrá de nuevos indicadores para la evaluación de la competencia de trabajo en grupo en caso de ser evaluable según el plan docente de cada asignatura [6], [15], [18].

El resto de este documento se estructura de la siguiente forma. La Sección 2 muestra el estado del arte. En la Sección 3 se describe la plataforma implementada. Finalmente, la Sección 4 presenta las conclusiones obtenidas y los trabajos futuros.

II. ESTADO DEL ARTE

En esta sección vamos a mostrar los conceptos técnicos en los que se basa el modelo implementado.

A. Conceptos Básicos

Formalmente, el problema subyacente a un proceso de toma de decisiones en grupo se puede definir de la siguiente manera:

Sea $X = \{x_1, x_2, \dots, x_n\}$ ($n \geq 1$) un conjunto de alternativas posibles y, teniendo en cuenta los valores de preferencia, $P = \{p_1, \dots, p_m\}$, proporcionados por un grupo de expertos $E = \{e_1, \dots, e_m\}$, ¿cómo deben ordenarse los valores del conjunto X de mejor a peor alternativa posible?

Por lo general, para resolver el problema, los procesos de toma de decisiones en grupo siguen los siguientes pasos [7]:

- 1) **Introducción de preferencias en el sistema:** Los expertos proporcionan sus preferencias al sistema. Las

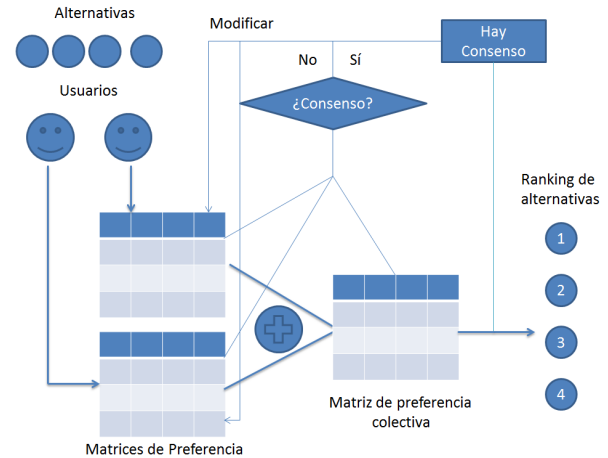


Fig. 1. Proceso de toma de decisiones con medidas de consenso.

preferencias definen directa o indirectamente un orden sobre el conjunto de alternativas.

- 2) **Cálculo de la matriz colectiva de preferencias:** La información de las preferencias proporcionadas por todos los usuarios es agregada en una sola pieza de información. La matriz colectiva representa la media de las preferencias proporcionadas.
- 3) **Proceso de selección de alternativas:** Usando la matriz colectiva y los operadores de selección deseados, se genera el ranking final de las alternativas.

El esquema comentado arriba tiene la desventaja de que no permite a los usuarios debatir ni llegar a ningún consenso antes de tomar la decisión final. Para solucionar este problema se utilizan las *medidas de consenso* [3]. Usando las matrices de preferencia de los expertos involucrados en el proceso de decisión, las medidas de consenso permiten determinar si los expertos opinan de forma parecida o si, por el contrario, tienen opiniones encontradas. De esta forma, si los expertos no llegan a un consenso, se les puede permitir que hablen y modifiquen sus preferencias con el objetivo de que se pongan de acuerdo. Si, por el contrario, todos están de acuerdo, se calcula el ranking de alternativas y el proceso de decisión termina. En la Fig. 1, podemos ver un esquema de como se definiría un proceso de toma de decisiones con medidas de consenso.

En un proceso de toma de decisiones, los expertos pueden proporcionar sus preferencias de diferentes formas. El procedimiento elegido es muy importante ya que establecerá la forma en que se deben realizar las operaciones necesarias para la toma de decisiones. Los métodos más comunes en la literatura son los siguientes [4], [19]:

- **Órdenes de preferencia:** El experto e_k proporciona sus preferencias utilizando una lista ordenada de preferencias $O^k = \{o^k(1), \dots, o^k(n)\}$ donde $o^k(\cdot)$ se define como una función de permutación sobre el conjunto de índices $\{1, \dots, n\}$ del conjunto de alternativas. De esta forma, las alternativas aparecen ordenadas de mejor a peor opción.

- **Funciones de utilidad:** El experto e_k comunica sus preferencias representadas como un conjunto de n valores de utilidad $U^k = \{u_i^k, i = 1, \dots, n\}$, $u_i^k \in [0, 1]$ donde u_i^k representa la evaluación que el experto e_k proporciona a la alternativa x_i .
- **Relaciones de preferencia difusa:** El experto e_k proporciona sus preferencias mediante una relación $P^k \subset X \times X$ cuya función de pertenencia es $\mu_{P^k} : X \times X \rightarrow [0, 1]$. $\mu_{P^k} = p_{ij}^k$ establece el grado de preferencia de la alternativa x_i sobre x_j .

En este trabajo, vamos a utilizar las relaciones de preferencia difusas como formato de representación de preferencias.

B. Medidas de consenso y proximidad

Para calcular el consenso de un proceso de toma de decisiones que utiliza relaciones de preferencia difusa, podemos seguir los pasos expuestos en el artículo de Mata et al. [13] y que detallamos a continuación:

- 1) Para cada par de expertos e_i y e_j , calculamos las matrices de similaridad sm_{ij} . Para ello, aplicamos la siguiente función de similaridad para cada uno de los valores de preferencia de cada dos expertos:

$$s(p_i^{lk}, p_j^{lk}) = 1 - |(p_i^{lk} - p_j^{lk})| \quad (1)$$

donde $s(p_i^{lk}, p_j^{lk})$ muestra la similaridad entre las preferencias de las alternativas x_i sobre x_k para los expertos e_i y e_j .

- 2) Una vez calculadas todas las matrices se agregan en una única matriz de consenso colectiva. Para ello podemos utilizar el operador de media:

$$sm_c = \phi(sm_{ij}), \forall i, \forall j, i \neq j, i < j \quad (2)$$

- 3) Utilizando la matriz de consenso colectiva sm_c , podemos calcular tres medidas distintas de consenso, cada una representativa de un nivel diferente:

- a) *Nivel 1, consenso entre pares de alternativas:* Cada valor de la matriz sm_c nos muestra el consenso alcanzado para cada par de alternativas:

$$cp^{lk} = cm^{lk} \forall l, k = 1, \dots, n, \wedge l \neq k \quad (3)$$

donde n es el número de alternativas del proceso de toma de decisiones.

- b) *Nivel 2, consenso en cada alternativa:* Para cada alternativa x_l , puede calcularse el nivel de consenso alcanzado, ca^l , usando la matriz cp tal y como muestra la siguiente expresión:

$$ca^l = \frac{\sum_{k=1, l \neq k}^n (cp^{lk} + cp^{kl})}{2(n-1)} \quad (4)$$

- c) *Nivel 3, consenso general del proceso:* Finalmente, podemos agregar los valores de consenso de cada una de las alternativas para obtener un valor de consenso global:

$$cr = \sum_{l=1}^n ca^l / n \quad (5)$$

También es interesante calcular la distancia que hay entre las preferencias de cada uno de los expertos a la matriz colectiva global. De esta forma, podemos ver si las opiniones del experto son similares o no a la de los demás y en que grado. Estas medidas de proximidad [10], al igual que las de consenso, se pueden calcular en tres niveles distintos:

- 1) *Nivel 1, proximidad en cada par de alternativas:* El nivel de proximidad para cada par de alternativas (x_l, x_k) , pp_i , del experto e_i , puede calcularse de la siguiente forma:

$$pp_i^{lk} = s(p_i^{lk}, p_c^{lk}) \quad (6)$$

donde p_c es la matriz colectiva.

- 2) *Nivel 2, proximidad para cada alternativa:* De manera análoga que en el consenso, podemos calcular el nivel de proximidad del experto a cada una de las alternativas mediante la siguiente expresión:

$$pa_i^l = \frac{\sum_{k=1, l \neq k}^n (pp_i^{lk} + pp_i^{kl})}{2 \cdot (n-1)} \quad (7)$$

- 3) *Nivel 3, Proximidad general:* El nivel de proximidad general de las preferencias del experto e_i puede calcularse usando la siguiente expresión:

$$pr_i = \frac{pa_i^l}{n} \quad (8)$$

Estas medidas serán útiles a la hora de identificar aquellos usuarios mas alejados de la opinión colectiva del grupo, y que por tanto, deberían acercar posturas con el mismo para maximizar el grado de consenso.

C. Métodos de agregación de información

Para calcular la matriz colectiva de preferencias es necesario agregar la información proporcionada por los expertos. Para ello debemos usar algún operador de agregación. A continuación expondremos algunos operadores que pueden usarse para completar esta tarea:

- el operador de media.
- el de media ponderada.
- el operador de media de pesos ordenados (OWA) [20], [21].

Para calcular la matriz de preferencias colectiva utilizando el operador de media podemos utilizar la siguiente expresión:

$$C_{ij} = \frac{p_{ij}^1 + \dots + p_{ij}^n}{m} \quad (9)$$

D. Operadores de selección

Para el proceso de selección, se utilizan los operadores de selección. Este tipo de operadores son capaces de obtener un ranking de alternativas a partir de una matriz colectiva de preferencias. Dos ejemplos de este tipo de operadores son los operadores de dominancia y no dominancia, GDD y GNDD respectivamente [8]. El operador GDD calcula el grado en que una alternativa domina a otra mientras que el de no dominancia se encarga de determinar qué alternativas no son dominadas por otras.

El operador GDD se calcula mediante la siguiente expresión:

$$GDD_i = \phi(c_{i1}, c_{i2}, \dots, c_{i(i-1)}, c_{i(i+1)}, \dots, c_{in}) \quad (10)$$

donde c es la matriz de preferencia colectiva y ϕ representa el operador de media.

El operador GNDD puede calcularse utilizando la siguiente expresión:

$$GNDD_i = \phi(c_{1i}^s, c_{2i}^s, \dots, c_{(i-1)i}^s, c_{(i+1)i}^s, \dots, c_{ni}^s) \quad (11)$$

donde

$$c_{ji}^s = \max\{c_{ji} - c_{ij}, 1\}$$

III. PLATAFORMA DE AYUDA A LA TOMA DE DECISIONES GRUPALES EN EL TRABAJO DE CLASE ONLINE

En esta sección vamos a mostrar el funcionamiento de la plataforma desarrollada.

A. Tecnologías empleadas para el desarrollo de la plataforma

A nivel tecnológico, la plataforma ha sido desarrollada teniendo como cimiento de programación el lenguaje PHP 8 y, a mas alto nivel, se ha utilizado el framework de aplicaciones web PHP Laravel.

Dicha plataforma cuenta entre sus dependencias con una extensión de seguridad diseñada por el propio ámbito de Laravel llamado Laravel JetStream. Dicha extensión nos ofrece una funcionalidad complementaria para la autenticación de usuarios y todo el conjunto referente a su gestión dentro de la plataforma. Incluyéndose el uso de los “teams” para los roles de los usuarios.

Para el envío de emails a los usuarios con la información sobre las decisiones a tomar o ya tomadas, se ha utilizado la dependencia PHPMailer por un servidor SMTP (Simple Mail Transfer Protocol) con un correo previamente definido y adecuado para esta envergadura.

El estilo de la página viene predefinido por una base de Bootstrap 5 con modificaciones e inclusión de clases de CSS propias.

La plataforma se basa en un modelo de tres capas, definida como Modelo-Vista-Controlador (MVC). La base de datos que utiliza el modelo es una base de datos relacional MySQL.

B. Ejemplo de uso y funcionalidad ofrecida por la Plataforma

En la pantalla inicial (ver Fig. 2), se puede configurar un nuevo proceso de decisión en grupo ajustando las alternativas y los expertos, que son divididos en diferentes roles según los tipos de decisiones que van a tomar. Una vez configurado el proceso, los usuarios pueden proceder a insertar sus preferencias individuales (ver Fig.3).

Cuando el usuario se identifica en la plataforma, tendrá una sección llamada “Mis Participaciones” donde puede ver todos los procesos de decisión pendientes en los que se ha visto involucrado y el estado en el que se encuentran.

Otra de las operaciones que ofrece la plataforma es el panel de administración (ver Fig.4). El administrador de cada



Fig. 2. Pantalla inicial



Fig. 3. Inserción de preferencias sobre las alternativas

proceso podrá ver qué usuarios han registrado sus preferencias, el estado de la decisión y en el caso de que el nivel de consenso sea suficiente, la alternativa escogida como mejor opción.

Mis proyectos administrados						
Título Proyecto	Descripción Proyecto	Fecha Fin	Votados Ya / Nº encuestados	Nº Alternativas	Consenso	Estado
Donde Viajar	Decision para saber donde viajar	2021-04-12	1 / 3	3	Barcelona	Actualizar Procesar resultados
regalo sara	que comprar para el cumpleaños de sara	2021-04-15	1 / 3	3		Actualizar Procesar resultados
Comida Domingos	Donde ir a comer el domingo	2021-04-21	0 / 3	3	Dominos	Actualizar Procesar resultados
a	a	2021-04-25	1 / 3	2		Actualizar Procesar resultados
b	b	2021-04-25	0 / 3	2	a	Actualizar Procesar resultados
b	b	2021-04-25	0 / 0	2		Actualizar Procesar resultados
b	b	2021-04-25	1 / 3	2		Actualizar Procesar resultados
Manual Drogas	Que ejercicio hacer del manual de drogas	2021-05-07	1 / 3	3	web de television	Actualizar Procesar resultados

Fig. 4. Panel de Administración

Una vez que todos los usuarios han insertado sus preferencias, al administrador se le enviará un mensaje con los resultados obtenidos (ver Fig.5) y, en caso de no alcanzar el nivel de consenso requerido, tendrá la opción de iniciar una nueva ronda de valoraciones de preferencias, mostrando previamente a cada experto, no solo el consenso alcanzado a nivel de alternativas, sino también la proximidad con respecto a la opinión colectiva para que vea si es necesario actualizar sus preferencias (y en qué dirección) como mecanismo de retroalimentación que trate de conseguir que el proceso sea convergente y en cada iteración el nivel de consenso sea mayor que en la anterior.

En cambio, si la acción escogida es “Terminar Valoración”,

debido a que el nivel de consenso ya es satisfactorio, el proyecto pasará a un estado de “Valoración Final” donde ya los expertos tendrán en el apartado “Decisiones Tomadas” un mensaje con el resultado sobre la decisión final consensuada.

Inicio / Procesado Votos

Datos del Proyecto con identificador : 4

Nombre	Descripción	Estado
Tema Trabajo Tema 1 PW	Vamos a elegir de que va a tratar el trabajo de programación del Tema 1 de PW	Iniciado

Resultados de las Alternativas

Alternativa	Resultado
Tienda Online de zapatillas	40,00% Compartir
Red Social académica	50,00% Compartir
Gestión de reservas de un centro deportivo	50,00% Compartir

Elección escogida

Tienda Online de zapatillas

La mayor decisión de esta tanda de votación ha sido la alternativa **Tienda Online de zapatillas** con un consenso de 60% como solución al problema planteado de **Tema Trabajo Tema 1 PW**

Si está satisfecho de la elección pulse **Terminar Valoración** de lo contrario determina la opción **Repetir Valoración**

Repetir Valoración Terminar Valoración

Fig. 5. Estado actual del proceso de toma de decisiones

Finalmente, en la Fig. 6, se puede observar la versión de la aplicación móvil de la plataforma funcionando bajo el sistema operativo Android.

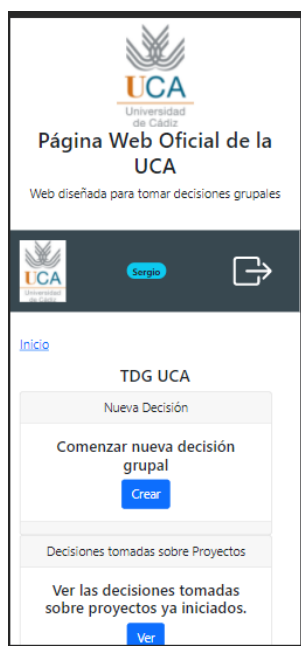


Fig. 6. Aplicación Android

IV. CONCLUSIONES Y TRABAJOS FUTUROS

En esta contribución se ha presentado una plataforma web de toma de decisiones en grupo, orientada al caso práctico del trabajo grupal en clase en un entorno docente no presencial propiciado por el Covid-19. De esta manera, la plataforma creada actúa como moderador virtual, ayudando a tomar decisiones sensatas para todos los participantes.

Como trabajos futuros, se propone la extensión de la plataforma para aceptar diferentes estructuras de representación de preferencias. También se plantea la posibilidad de implementar un sistema de control de consistencia y de

estimación de preferencias no rellenas. Por último, se desea implementar una plataforma dinámica, que permita que los elementos del problema puedan cambiar una vez comenzado el proceso de negociación.

AGRADECIMIENTOS

Este trabajo ha sido financiado por la Agencia Estatal de Investigación a través del proyecto “iScience” PID2019-105381GA-I00 / AEI / 10.13039/501100011033 , por el Programa de Proyectos de Innovación y Mejora Docente de la Universidad de Cádiz (UCA/R100REC/2020) y por el Programa de Fomento e Impulso de la actividad Investigadora de la Universidad de Cádiz.

REFERENCIAS

- [1] Bellman, R.E., Zadeh, L.A.: Decision-making in a fuzzy environment. *Management science* **17**(4), B-141 (1970)
- [2] Cabrerizo, F.J., Moreno, J.M., Pérez, I.J., Herrera-Viedma, E.: Analyzing consensus approaches in fuzzy group decision making: advantages and drawbacks. *Soft Computing* **14**(5), 451–463 (2010)
- [3] Cabrerizo, F.J., Moreno, J.M., Pérez, I.J., Herrera-Viedma, E.: Analyzing consensus approaches in fuzzy group decision making: advantages and drawbacks. *Soft Computing* **14**(5), 451–463 (2010)
- [4] Chiclana, F., Herrera, F., Herrera-Viedma, E.: Integrating three representation models in fuzzy multipurpose decision making based on fuzzy preference relations. *Fuzzy sets and Systems* **97**(1), 33–48 (1998)
- [5] Collins, B.E., Guetzkow, H.S.: A social psychology of group processes for decision-making. Wiley New York (1964)
- [6] García-Cabrero, B., Luna-Serrano, E., Ponce-Ceballos, S., Cisneros-Cohenour, E., Cordero-Arroyo, G., Espinoza-Díaz, Y., García-Vigil, M.H.: Las competencias docentes en entornos virtuales: Un modelo para su evaluación. *Revista Iberoamericana de Educación a Distancia* **21**(1), 343–365 (2018)
- [7] Herrera, F., Alonso, S., Chiclana, F., Herrera-Viedma, E.: Computing with words in decision making: foundations, trends and prospects. *Fuzzy Optimization and Decision Making* **8**(4), 337–364 (2009)
- [8] Herrera, F., Herrera-Viedma, E., Verdegay, J.L.: A sequential selection process in group decision making with a linguistic assessment approach. *Information Sciences* **85**(4), 223–239 (1995)
- [9] Herrera-Viedma, E., Cabrerizo, F.J., Kacprzyk, J., Pedrycz, W.: A review of soft consensus models in a fuzzy environment. *Information Fusion* **17**, 4–13 (2014)
- [10] Herrera-Viedma, E., Martínez, L., Mata, F., Chiclana, F.: A consensus support system model for group decision-making problems with multigranular linguistic preference relations. *Fuzzy Systems, IEEE Transactions on* **13**(5), 644–658 (2005)
- [11] Huang, C.L., Lo, C.C., Chao, K.M., Younas, M.: Reaching consensus: A moderated fuzzy web services discovery method. *Information and Software Technology* **48**(6), 410–423 (2006)
- [12] Kacprzyk, J.: Group decision making with a fuzzy linguistic majority. *Fuzzy sets and systems* **18**(2), 105–118 (1986)
- [13] Mata, F., Martínez, L., Herrera-Viedma, E.: An adaptive consensus support model for group decision-making problems in a multigranular fuzzy linguistic context. *IEEE Transactions on fuzzy Systems* **17**(2), 279–290 (2009)
- [14] Morente-Molinera, J., Castellón-Fuentes, J., Herrera-Viedma, E., López-Herrera, A.: A decision making web platform for educational centers. In: *INTED2013 Proceedings*. pp. 1368–1376. IATED (2013)
- [15] Pazmiño, D.X.S.: Evaluación de las competencias digitales en la formación inicial docente de los y las estudiantes de la Facultad de Filosofía, Letras y Ciencias de la Educación de la Universidad Central del Ecuador. Ph.D. thesis, Universitat d’Alacant-Universidad de Alicante (2021)
- [16] Pérez, I.J., Cabrerizo, F.J., Alonso, S., Dong, Y., Chiclana, F., Herrera-Viedma, E.: On dynamic consensus processes in group decision making problems. *Information Sciences* **459**, 20–35 (2018)
- [17] Pérez, I.J., Cabrerizo, F.J., Alonso, S., Herrera-Viedma, E.: A new consensus model for group decision making problems with non-homogeneous experts. *IEEE Transactions on Systems, Man, and Cybernetics: Systems* **44**(4), 494–498 (2013)

- [18] Pueyo, Á.P., Berrocal, O.M.C., Bernardino, C.H., Martín, J.J.B., Cobo, D.V., González, L.P., Álvarez, I.H., Fernández, L.C., Benito, R.M., Nava, M.M., et al.: Programar y evaluar competencias básicas en 15 pasos, vol. 301. Graó (2013)
- [19] Tanino, T.: On group decision making under fuzzy preferences. In: Multiperson Decision Making Models using Fuzzy Sets and Possibility Theory, pp. 172–185. Springer (1990)
- [20] Yager, R.R.: On ordered weighted averaging aggregation operators in multicriteria decisionmaking. *Systems, Man and Cybernetics, IEEE Transactions on* **18**(1), 183–190 (1988)
- [21] Yager, R.R.: Quantifier guided aggregation using OWA operators. *International Journal of Intelligent Systems* **11**(1), 49–73 (1996)
- [22] Zadeh, L.A.: The concept of a linguistic variable and its application to approximate reasoning-I. *Information sciences* **8**(3), 199–249 (1975)

XX Congreso Español sobre Tecnologías y Lógica Fuzzy

SESIÓN ESPECIAL:

FUNCIONES DE AGREGACIÓN Y PREAGREGACIÓN

Y SUS APLICACIONES



Mapas de características adicionales en redes neuronales profundas a partir de agregaciones basadas en orden

I. Dominguez-Catena, D. Paternain, M. Galar

Institute of Smart Cities

Universidad Pública de Navarra, Pamplona, España

{iris.dominguez, mikel.galar, daniel.paternain}@unavarra.es

Resumen—En los últimos años hemos observado avances enormes en el área del Aprendizaje Automático, especialmente a través del uso de Redes Neuronales Profundas. Uno de los ejemplos más relevantes es la clasificación de imágenes, donde las Redes Neuronales Convolucionales han demostrado ser una herramienta muy precisa. Aunque las funciones de agregación, como los operadores OWA, ya han sido empleadas en combinación con redes neuronales, en este trabajo proponemos y exploramos una nueva forma de integrar operadores OWA en redes neuronales. Para ello, introducimos los operadores OWA dentro de una nueva capa en una red neuronal convolucional. Realizamos varios experimentos introduciendo la capa en una red basada en VGG-13 y mostramos como la capa introduce nuevo conocimiento en la red.

Index Terms—Redes Neuronales, Redes Neuronales Convolucionales, aprendizaje profundo, operadores OWA

I. INTRODUCCIÓN

Uno de los problemas más estudiados en aprendizaje automático es el de clasificación supervisada de imágenes [1], [2]. En este tipo de problemas intentamos desarrollar un modelo matemático que aprende sobre un conjunto de imágenes etiquetadas, y que después sea capaz de etiquetar nuevos ejemplos apropiadamente. Actualmente, la técnica más común es el uso de redes neuronales convolucionales (CNNs), el foco de atención de este trabajo.

Las medidas ponderadas ordenadas (OWA) [3], [4] son un tipo de agregación paramétrica comúnmente utilizada en el campo del aprendizaje automático y la lógica difusa. En trabajos previos, los operadores OWA han sido integrados en aprendizaje profundo principalmente como un método para combinar las salidas de diferentes clasificadores (ensembles) [5], [6], [7], [8]. También han sido empleados con resultados interesantes en las capas de pooling de las CNNs [9].

A diferencia de los trabajos anteriores, intentamos emplear operadores OWA en las capas internas de una CNN, con el objetivo de incrementar la información disponible para la siguiente capa, añadiendo muy pocos parámetros a la red. Nuestro objetivo es generar información sin coste en la red, obteniendo los mapas de características en un punto de la red y añadiendo información derivada que sería difícil de

conseguir a través de operadores convolucionales normales. Para esto, proponemos implementar una capa de agregaciones OWA a nivel de canal, que aprenda los pesos de varios operadores junto con los parámetros de la red, y los aplique para añadir mapas de características virtuales a la información ya existente.

Para comprobar el funcionamiento de esta propuesta hemos considerado una arquitectura base de tipo VGG13 [10], y hemos insertado capas OWA en ella. Probamos diferentes configuraciones de la capa sobre los datasets de clasificación de imágenes CIFAR10 y CIFAR100 [11].

El resto del trabajo está organizado de la siguiente manera. La Sección II describe la literatura relevante a nuestra propuesta. La Sección III revisa algunos conceptos preliminares sobre OWA y CNN. Después, la Sección IV especifica nuestra metodología para la inserción de la capa OWA y explica el funcionamiento interno de esta. La Sección V presenta los experimentos que hemos diseñado para probar la capa y los detalles de implementación. A continuación, la Sección VI recopilamos los resultados experimentales y los analizamos. Finalmente, la Sección VII concluye este trabajo y propone algunas líneas de trabajo futuras.

II. LITERATURA RELACIONADA

En la literatura se han explorado previamente varias maneras de combinar operadores OWA y redes neuronales [5], [6], [7], [8], [9]. La manera más habitual es emplear OWAs sobre la salida de las redes, agregando sus resultados [5], [6], [7], [8]. La otra técnica habitual es la sustitución de las agregaciones en las capas de pooling por OWAs [9].

El empleo de operaciones de agregación basadas en medidas difusas [12], como las integrales Choquet y Sugeno (de las cuales los operadores OWA son un caso particular), para agregar ensembles de redes neuronales ha sido estudiado en múltiples ocasiones recientemente [5], [6], [7], [8]. En estos sistemas, se entrenan una serie de clasificadores independientemente, y se emplean operadores de agregación sobre los resultados. En este paso final es donde se pueden emplear operadores OWA y otros operadores de agregación basados en medidas difusas.

Los operadores basados en medidas difusas también se han utilizado como operadores de reducción en las capas de

pooling de CNNs [9]. En este caso, las agregaciones habituales empleadas en las capas de pooling se reemplazan directamente por otros operadores. La idea es obtener representaciones más fieles de la información original tras la reducción de dimensionalidad que se da en estas capas.

Finalmente, la inspiración principal para este trabajo se encuentra en [13], donde los autores proponen la creación de lo que ellos denominan una “Capa Difusa”. Esta capa, insertada en diferentes puntos de una CNN, realiza seis operaciones OWA predeterminadas (máximo, mínimo, máximo suavizado, mínimo suavizado, media y un operador aleatorio) sobre los canales de la red, ordenados en base a una medida de la entropía de cada canal. Mientras que los autores aplicaban su método a un problema de segmentación de imágenes [14], experimentalmente hemos comprobado que intentar trasladar la técnica a problemas generales de clasificación de imágenes no obtenía buenos resultados. En nuestra opinión, en el caso de las redes empleadas para clasificación, existe demasiada información codificada en el orden de los mapas de características, que se pierde al aplicar los operadores OWA. Por tanto, diseñamos nuestra propuesta con la idea de aumentar la información en la red, en vez de reemplazarla, concatenando los nuevos mapas de características a los ya presentes. De esta manera, empleamos la salida de nuestra capa OWA como un complemento a la salida de las convoluciones estándar, proporcionando a las siguientes capas información que sería de otra forma difícil de obtener (información global a partir de las métricas del canal). Además, incluimos los pesos de los operadores OWA como parámetros de la red, aprendiéndolos en vez de mantenerlos fijos como en [13].

III. PRELIMINARES

III-A. Operadores OWA

Los *operadores OWA* fueron propuestos inicialmente por Yager [3]. Estos operadores son mapeos $F: \mathcal{R}^n \rightarrow \mathcal{R}$ basados en una colección de pesos $W = [w_1, \dots, w_n]$, con la condición de $w_i \in [0, 1]$ para todo $i = 1, \dots, n$ y $\sum_{i=1}^n w_i = 1$, y definidos como:

$$F(a_1, \dots, a_n) = \sum_{j=1}^n w_j b_j \quad (1)$$

donde b_j representa el j -ésimo elemento más grande de a_i .

Algunos ejemplos notables de operadores OWA serían el máximo ($W = [1, 0, \dots, 0]$), mínimo ($W = [0, \dots, 0, 1]$), y la media aritmética ($W = [\frac{1}{n}, \dots, \frac{1}{n}]$).

III-B. Redes Neuronales Convolucionales Profundas

Las redes neuronales convolucionales (CNNs) modifican la arquitectura habitual de las redes neuronales para especializarse en información espacial [15]. El uso más común, el procesamiento de imágenes, supone reconocer las relaciones 2D espaciales de la información de entrada, y emplear operaciones convolucionales que sólo toman en consideración píxeles vecinos de la imagen. Esto también supone perder algo de información de gran escala de la imagen, al sólo trabajarse con píxeles cercanos.

Otra importante característica de las CNNs son las capas de *pooling* [1]. Estas capas, como las convolucionales, reconocen la estructura espacial de las imágenes y mapas de características derivados, pero en vez de agregar mapas (como las convolucionales) operan sobre un solo canal, resumiendo la información por bloques y reduciendo el tamaño de cada mapa de características independientemente.

Algunas arquitecturas de CNN conocidas que han sido desarrolladas para clasificación de imágenes son LeNet [16], la familia VGG [10] y ResNet [2]. En este trabajo nos centraremos en VGG, pero nuestra metodología podría ser extrapolada a casi cualquier arquitectura CNN.

IV. METODOLOGÍA

IV-A. Capa OWA

Nuestra capa OWA propuesta funcionará tomando una entrada de N imágenes, con una resolución de I filas por J columnas y C_{in} canales de profundidad (mapas de características de entrada), y agregando C_f nuevos canales a los originales, con $C_f \in [0, C_{in}]$. La salida será de N imágenes con la misma resolución $I \times J$, pero $C_{out} = C_{in} + C_f$ canales de profundidad, $C_{out} \geq C_{in}$. Para generar esos C_f nuevos mapas de características, aplicaremos C_f operadores OWA sobre los canales de entrada. Estos operadores OWA compartirán la misma función de ordenación, que utilizará métricas calculadas por canal para reordenarlos. Después, cada uno de estos operadores OWA generará un nuevo mapa de características como una combinación lineal de los canales ordenados, a partir de un vector de pesos propio. Estos vectores de pesos, uno por cada uno de los C_f OWAs aprendidos, se aprenderán y actualizarán como parámetros de la red. Profundizamos más sobre esto en la Sección IV-C. La arquitectura general de la capa se muestra en la Figura 1.

IV-B. Ordenación de canales

Al decidir trabajar por canal (y no por píxel), resulta vital definir precisamente la función de ordenación, esto es, en base a que métrica ordenaremos los canales. Dado un canal X de tamaño $I \times J$, consideramos las siguientes métricas:

- *Entropía del canal.* Empleamos la fórmula de entropía de Shannon [17] aplicada a los valores de todos los píxeles del canal,

$$H(X) = - \sum_{i=1}^I \sum_{j=1}^J x_{ij} \log x_{ij} \quad (2)$$

Como esta función está diseñada para trabajar sobre vectores de elementos en el rango $x_{ij} \in [0, 1]$ con $\sum_{i=1}^I \sum_{j=1}^J x_{ij} = 1$, primero aplicamos la función *softmax* a la entrada X para normalizar el mapa de características,

$$\text{Softmax}(x_{kl}) = \frac{e^{x_{kl}}}{\sum_{i=1}^I \sum_{j=1}^J e^{x_{ij}}} \quad (3)$$

Intuitivamente, podemos entender la entropía como una medida de desorden, de la cantidad de información

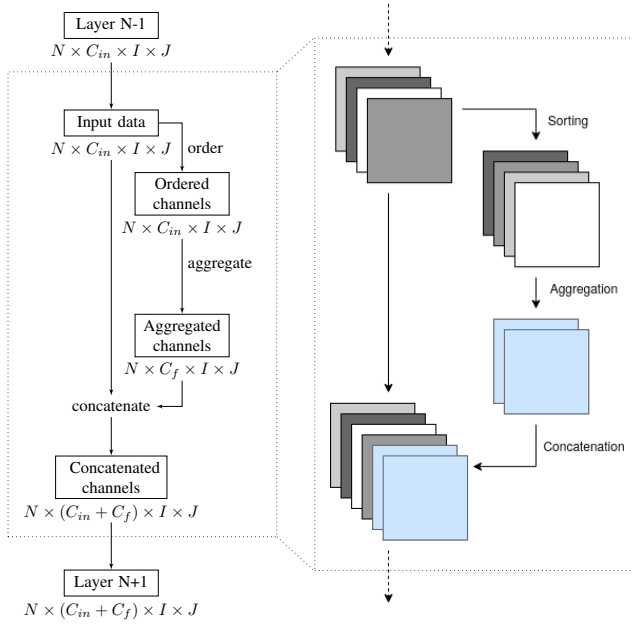


Figura 1. Propuesta de estructura para la capa OWA.

codificada en un canal. Un valor más elevado de entropía se corresponde con más uniformidad en los valores de entrada, mientras que un valor más pequeño se corresponde con un mayor contraste e información en la entrada.

- **Suma de valores.** Consideramos también la suma simple de activaciones en el canal,

$$S(X) = \sum_{i=1}^I \sum_{j=1}^J x_{ij} \quad (4)$$

- **Variación total** [18]. Considerando las características espaciales de la imagen, calculamos las diferencias entre cada píxel y su vecindad, tanto horizontal como verticalmente, y sumamos el valor absoluto de las diferencias en la imagen.

$$TV_v(X) = \sum_{i=2}^I \sum_{j=1}^J |x_{i,j} - x_{i-1,j}| \quad (5)$$

$$TV_h(X) = \sum_{i=1}^I \sum_{j=2}^J |x_{i,j} - x_{i,j-1}| \quad (6)$$

$$TV(X) = TV_v(X) + TV_h(X) \quad (7)$$

La Variación Total (TV), como se define en [18], es una medida que nos dice cuánta variación existe entre los píxeles de una imagen y sus vecinos. Esta medida será elevada para imágenes con muchos bordes nítidos y contraste elevado, y baja para imágenes muy uniformes.

- **Mediana de los valores.** Un operador OWA clásico,

$$M(X) = \text{median}(x_{11}, \dots, x_{IJ}) \quad (8)$$

donde el operador mediana devuelve el $\text{ceil}(I \cdot J/2)$ -ésimo elemento más grande de X si $I \cdot J$ es impar o la media aritmética de los $I \cdot J/2$ -ésimo y $I \cdot J/2 + 1$ -ésimos elementos más grandes de X , si $I \cdot J$ es par.

- **Máximo de los valores.** En este contexto, el valor del píxel más activado del canal,

$$\text{MAX}(X) = \max(x_{11}, \dots, x_{IJ}) \quad (9)$$

Hemos considerado, además, dos métodos de ordenación de referencia que no se basan en los valores del canal:

- **No ordenación.** Mantenemos los canales en el orden que vienen, convirtiendo el OWA en una agregación estándar.
- **Ordenación aleatoria.** Esta ordenación nos permitirá introducir ruido equivalente al resultado de aplicar la capa OWA, permitiéndonos descartar que se obtengan mejoras en la precisión debido a fenómenos de tipo regularización.

IV-C. Agregación ponderada

Con los canales de entrada ya ordenados, realizaremos una agregación ponderada, empleando C_f vectores de pesos (uno por OWA), cada uno de ellos con un peso por cada canal de entrada (C_{in} en total). En nuestra propuesta, los inicializaremos de forma aleatoria siguiendo una distribución $U(0,1)$. Trataremos estos pesos como parámetros de la red, y por tanto se aprenderán a través del método de retropropagación habitual.

Estos pesos no se encuentran directamente restringidos, y conforme son aprendidos pueden llegar a tomar cualquier valor. Para ajustarnos a la definición de OWA dada en la Sección III-A, donde para cada $w_i, i \in 1, \dots, C_f$, requerimos $w_i \in [0,1]$ y $\sum_{i=1}^{C_f} w_i = 1$, aplicamos algunas transformaciones antes de realizar la agregación. Primero, empleamos una función ReLU para convertir los pesos negativos en 0, y después normalizamos dividiendo el vector por la suma de sus valores, de forma que sumen en total 1.

$$\text{ReLU}(x) = \max(x, 0) \quad (10)$$

$$w_j = \frac{\text{ReLU}(x_j)}{\sum_{i=1}^{C_f} \text{ReLU}(x_i)} \quad (11)$$

El resultado es un vector de pesos OWA correcto que puede ser aplicado directamente a los canales.

V. MARCO EXPERIMENTAL

V-A. Datasets

Como datasets de pruebas hemos escogido CIFAR10 y CIFAR100 [11]. Ambos son datasets de clasificación de imágenes muy conocidos, compuestos cada uno por 60.000 imágenes en color en una resolución de 32x32 píxeles. En el caso de CIFAR10 se reparten en 10 clases balanceadas y fácilmente distinguibles (6.000 por clase), mientras que CIFAR100 etiqueta sus imágenes en 100 clases (600 ejemplos por clase). Ambos distribuyen las imágenes en un conjunto de entrenamiento de 50.000 imágenes y un conjunto de prueba de 10.000, balanceados por clase.

La elección de estos datasets está motivada por el número de pruebas de configuraciones diferentes que queremos realizar. Estos datasets de pequeño tamaño nos facilitan realizar múltiples repeticiones y entrenamientos por cada configuración, de manera que podamos obtener resultados promediados estables.

V-B. Arquitectura

Decidimos trabajar con una arquitectura de referencia de tipo VGG [10]. Se trata de una arquitectura relativamente simple pero con muy buenos resultados, y el entrenamiento es lo suficientemente rápido como para realizar múltiples repeticiones. Al tener una estructura lineal, nos ofrece múltiples puntos de inserción para nuestra capa OWA sobre los que evaluar sus características.

De las posibles configuraciones de esta familia hemos seleccionado la VGG13, que experimentalmente nos ofrece buenos resultados para CIFAR sin ser excesivamente costosa de entrenar.

La red en particular consta de 10 capas convolucionales, compuestas por una convolución, una capa de normalización por lotes y una activación no lineal ReLU, seguidas por un clasificador final. Estas 10 capas se reparten en 5 bloques, cada uno de ellos delimitado por capas MaxPool que reducen la resolución del mapa de características por la mitad. El clasificador final se compone de 3 capas densamente conectadas en la estructura original, pero en nuestro caso lo reducimos a una sola, considerando la pequeña resolución de las imágenes de CIFAR, siguiendo el ejemplo de [19].

Para nuestros experimentos, consideraremos como potenciales puntos de inserción de capas OWA los puntos justo antes de cada capa convolucional, a excepción de la primera, resultando en 9 posibles puntos de inserción. La red de referencia será siempre la original sin ninguna inserción, y introduciremos capas en los puntos de inserción para generar las configuraciones de estudio. Esta estructura se refleja en la Tabla I.

V-C. Detalles de Implementación

La implementación de estos experimentos se ha realizado en PyTorch 1.3.1 y Fastai 1.0.58.

Para todas las configuraciones empleamos los mismos hiperparámetros, en concreto, un ratio de aprendizaje máximo de $1e^{-2}$ con una política de entrenamiento 1cycle [20]. Este parámetro se ha determinado empleando la herramienta *lr_finder* de Fastai, optimizándolo para la red de referencia (sin capas OWA). Para todos los experimentos se ha empleado un tamaño de lote de 1024.

Adicionalmente hemos empleado aumentación de datos, siguiendo el ejemplo de [21]. En particular, hemos realizado volteos horizontales con una probabilidad de 0.5, y padding de 4 píxeles (empleado espejado para rellenar las regiones externas) seguido de un recorte aleatorio a la resolución original (32×32).

V-D. Evaluación

Para la evaluación de los resultados de los experimentos, optamos por repetir el entrenamiento de cada configuración

Tabla I
ARQUITECTURA DE LA RED*.

Nombre	Tamaño de núcleo	Paso	Tamaño de salida
input_data	-	-	$32 \times 32 \times 3$
conv1_1	3×3	1	$32 \times 32 \times 64$
OWA ₁	-	-	$32 \times 32 \times (64 + C_f)$
conv1_2	3×3	1	$32 \times 32 \times 64$
maxpool	2×2	2	$16 \times 16 \times 64$
OWA ₂	-	-	$16 \times 16 \times (64 + C_f)$
conv2_1	3×3	1	$16 \times 16 \times 128$
OWA ₃	-	-	$16 \times 16 \times (128 + C_f)$
conv2_2	3×3	1	$16 \times 16 \times 128$
maxpool	2×2	2	$8 \times 8 \times 128$
OWA ₄	-	-	$8 \times 8 \times (128 + C_f)$
conv3_1	3×3	1	$8 \times 8 \times 256$
OWA ₅	-	-	$8 \times 8 \times (256 + C_f)$
conv3_2	3×3	1	$8 \times 8 \times 256$
maxpool	2×2	2	$4 \times 4 \times 256$
OWA ₆	-	-	$4 \times 4 \times (256 + C_f)$
conv4_1	3×3	1	$4 \times 4 \times 512$
OWA ₇	-	-	$4 \times 4 \times (512 + C_f)$
conv4_2	3×3	1	$4 \times 4 \times 512$
maxpool	2×2	2	$2 \times 2 \times 512$
OWA ₈	-	-	$2 \times 2 \times (512 + C_f)$
conv5_1	3×3	1	$2 \times 2 \times 512$
OWA ₉	-	-	$2 \times 2 \times (512 + C_f)$
conv5_2	3×3	1	$2 \times 2 \times 512$
maxpool	2×2	2	$1 \times 1 \times 512$
flatten	-	-	512
linear	-	-	10

* Las capas marcadas como OWA_x son los posibles puntos de inserción para las nuevas capas OWA.

50 veces, cada una de ellas de cero (reiniciando la red) y entrenando por 30 épocas. De estos resultados recogemos la precisión en test final de cada repetición, y empleamos la media y desviación estándar de esas precisiones como medidas principales. Para comparar los resultados de las configuraciones modificadas respecto de la original empleamos el test no paramétrico de Mann-Whitney U [22]. Calculamos este test considerando como hipótesis nula que la referencia obtiene precisiones mayores o iguales que la versión modificada.

La referencia es siempre una versión sin modificar de la red. En el tercer experimento, además, entrenamos la red con dos ordenaciones de canal de referencia (no ordenarlos y ordenarlos aleatoriamente).

V-E. Configuración de los experimentos

Dada la gran cantidad de combinaciones de parámetros posibles hemos dividido el trabajo en 3 experimentos principales. Cada experimento pretende determinar uno de los principales parámetros de la capa OWA: posición, cantidad de operadores y métrica de orden. Cada experimento se repite con CIFAR10 y con CIFAR100.

1. Posición de la capa. En el primer experimento, mantenemos constante el número de operadores aprendidos ($C_f = 16$) y generamos configuraciones para todas las posiciones de capa posibles (OWA₁ a OWA₉). Como métrica de orden consideramos la suma de activaciones.
2. Número de operadores. En el segundo experimento probamos una variedad de números de operadores ($C_f = 4$,

$C_f = 8$, $C_f = 16$ y $C_f = 32$) para las dos mejores posiciones del experimento anterior, ambas con las mismas métricas.

3. Métrica de ordenación. En el tercer experimento nos centramos en las métricas de orden, fijando las mejores configuraciones de posición y número de operadores del experimento anterior. En este experimento, además, estudiamos las matrices de pesos de los operadores OWA aprendidos en la red.

VI. ESTUDIO EXPERIMENTAL

VI-A. Posición de la capa

Los resultados del primer experimento se recogen en la Tabla II. Podemos observar como existe una fuerte dependencia entre el punto de inserción y la precisión obtenida. En el caso de CIFAR10 observamos como todos los puntos de inserción entre OWA₂ y OWA₅ obtienen una mejora de precisión estadísticamente significativa respecto de la media, con el mejor en OWA₄. En CIFAR100 podemos observar resultados similares, con las mejores configuraciones en OWA₃ y OWA₅.

Sospechamos que esta tendencia a obtener mejores resultados en las capas inferiores está ligada al tamaño de imagen muy reducido de nuestro dataset, de 32×32 píxeles. Esto, en combinación con la arquitectura VGG, hace que las capas superiores tengan tamaños de imagen realmente pequeños (4×4 y menores a partir de OWA₆), haciendo que las métricas de capa no aporten información respecto de las convoluciones estándar.

Tabla II
RESULTADOS EN FUNCIÓN DEL PUNTO DE INSERCIÓN.

Capa	CIFAR10 prec.	CIFAR100 prec.
referencia	92.44 ± 0.17	69.74 ± 0.27
OWA ₁	92.40 ± 0.19	69.85 ± 0.29*
OWA ₂	92.53 ± 0.19*	69.87 ± 0.32*
OWA ₃	92.52 ± 0.18*	69.97 ± 0.27*
OWA ₄	92.55 ± 0.18*	69.95 ± 0.24*
OWA ₅	92.51 ± 0.17*	69.97 ± 0.28*
OWA ₆	92.45 ± 0.20	69.75 ± 0.33
OWA ₇	92.44 ± 0.17	69.82 ± 0.25
OWA ₈	92.44 ± 0.18	69.79 ± 0.30
OWA ₉	92.49 ± 0.20	69.78 ± 0.25

Los resultados señalados con * mejoran la referencia con p-valor < 0.05.

VI-B. Número de operadores

Los resultados del segundo experimento se resumen en la Tabla III. A partir de los resultados del primer experimento, decidimos seguir explorando los puntos de inserción en OWA₃ y OWA₄ y mantenemos la métrica de ordenación de suma de activaciones.

En estos resultados observamos una cierta tendencia a favorecer mayores números de operadores aprendidos, pero sin una variación particularmente significativa. Todos los resultados para $C_f = 8$ y $C_f = 16$, en ambos datasets, obtienen resultados estadísticamente significativos con p-valor < 0.05.

Tabla III
RESULTADOS SEGÚN NÚMERO DE OPERADORES.

Capa	C_f	CIFAR10 prec.	CIFAR100 prec.
referencia	-	92.44 ± 0.17	69.74 ± 0.27
OWA ₃	4	92.49 ± 0.19	69.82 ± 0.32
	8	92.51 ± 0.21*	69.88 ± 0.26*
	16	92.52 ± 0.18*	69.97 ± 0.27*
	32	92.57 ± 0.16*	69.82 ± 0.28
OWA ₄	4	92.47 ± 0.17	69.81 ± 0.33
	8	92.50 ± 0.19*	69.91 ± 0.31*
	16	92.55 ± 0.18*	69.95 ± 0.24*
	32	92.51 ± 0.17*	69.90 ± 0.29*

Los resultados señalados con * mejoran la referencia con p-valor < 0.05.

VI-C. Métricas de orden

Los resultados del tercer experimento se recogen en la Tabla IV. En este experimento se han empleado, en base a los resultados de los anteriores experimentos, el punto de inserción OWA₃ con $C_f = 32$ para CIFAR10 y el punto OWA₃ con $C_f = 16$ para CIFAR100. Podemos observar que la suma de activaciones obtiene mejores resultados de precisión que el resto de métricas, seguida de cerca por la variación total, tanto en CIFAR10 como en CIFAR100.

Las dos medidas de referencia, la ordenación aleatoria y la no ordenación de canales, obtienen resultados similares a la referencia sin capa OWA, de forma consistente con la hipótesis de que la capa OWA introduce nueva información a través de la función de ordenación.

Tabla IV
RESULTADOS SEGÚN LA MÉTRICA DE ORDEN.

Orden	CIFAR10 prec.	CIFAR100 prec.
referencia	92.44 ± 0.17	69.74 ± 0.27
activ_sum	92.57 ± 0.16*	69.97 ± 0.27*
total_var	92.55 ± 0.19*	69.91 ± 0.24*
max_activ	92.51 ± 0.21*	69.74 ± 0.28
median_activ	92.48 ± 0.19	69.84 ± 0.31
entropy	92.47 ± 0.16	69.80 ± 0.25
random	92.45 ± 0.17	69.76 ± 0.26
no_sorting	92.43 ± 0.19	69.79 ± 0.30

Los resultados señalados con * mejoran la referencia con p-valor < 0.05.

VI-D. Matrices de pesos

Resulta de especial interés estudiar las matrices de pesos obtenidas de los entrenamientos. En la Figura 2 mostramos 8 ejemplos de matrices de pesos de las capas OWA aprendidas, todas ellas sobre la misma configuración base (inserción en OWA₂, $C_f = 8$) y variando la métrica de orden. El tamaño de estas matrices es de 8×64 , siendo $C_f = 8$ el número de operadores aprendidos (cada uno representado en una línea de la imagen) y $C_{in} = 64$ el número de pesos de cada operador, correspondientes a los canales de entrada a capa.

Podemos apreciar como el sistema converge a patrones bastante claros para la mayor parte de las agregaciones. Estos patrones se corresponden con operaciones de tipo máximo y

mínimo suavizados, donde la mayor parte del peso se reparte en las capas de entrada con mayor o menor valor de métrica asociados.

En concreto, se observa que en general el sistema converge a operaciones de tipo mínimo suavizado para todos los operadores excepto la entropía, aunque en todos los casos con algún operador de tipo máximo suavizado intercalado. En el caso de la variación total y la suma de activaciones se observan operadores con valores más concentrados, mientras que para la mediana los operadores se acercan más a una media. En el caso de la no ordenación y la ordenación aleatoria, como es de esperar, no se aprecian patrones claros.

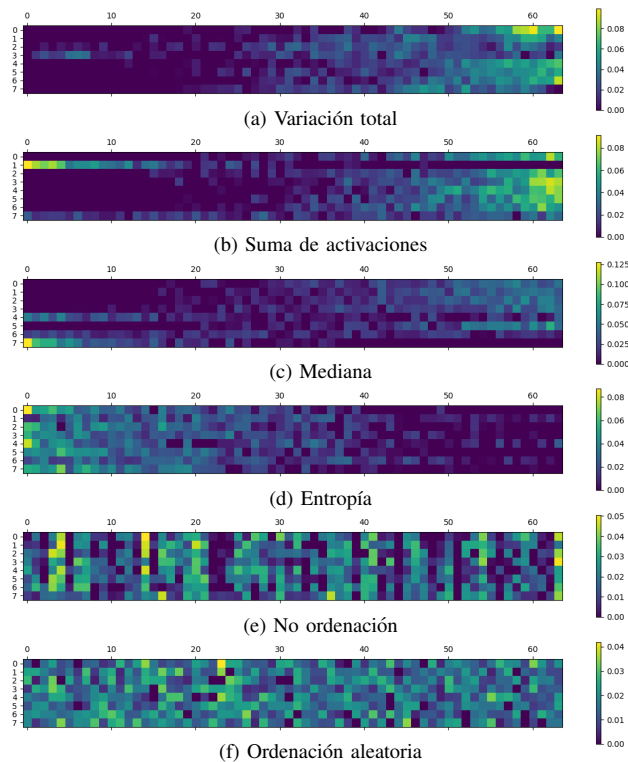


Figura 2. Operadores OWA aprendidos por la capa propuesta según diferentes métricas de orden. El eje vertical se corresponde con los diferentes operadores aprendidos, mientras que el horizontal con los pesos de los canales de entrada.

VII. CONCLUSIONES Y TRABAJO FUTURO

En este trabajo hemos propuesto la inserción de operadores OWA dentro de CNNs como un método para aumentar la información de los mapas de características. Si bien los resultados obtenidos no se posicionan en el estado del arte, consideramos que prueban sin dudas el potencial de esta técnica.

En el futuro, sería importante analizar si este enfoque puede aplicarse en redes más complejas, como ResNet [2] y arquitecturas similares, que se adaptan mejor a ciertos problemas. Es necesario, en general, una investigación más profunda sobre cómo la red neuronal está obteniendo la ventaja que hemos constatado en nuestros experimentos.

REFERENCIAS

- [1] A. Krizhevsky, I. Sutskever, and G. E. Hinton, "Imagenet classification with deep convolutional neural networks," in *Advances in Neural Information Processing Systems* 25, F. Pereira, C. J. C. Burges, L. Bottou, and K. Q. Weinberger, Eds. Curran Associates, Inc., 2012, pp. 1097–1105.
- [2] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016, pp. 770–778.
- [3] R. R. Yager, "On ordered weighted averaging aggregation operators in multicriteria decisionmaking," *IEEE Transactions on Systems, Man, and Cybernetics*, vol. 18, no. 1, pp. 183–190, 1988.
- [4] R. R. Yager, "Families of owa operators," *Fuzzy Sets and Systems*, vol. 59, no. 2, pp. 125 – 148, 1993.
- [5] G. J. Scott, R. A. Marcum, C. H. Davis, and T. W. Nivn, "Fusion of deep convolutional neural networks for land cover classification of high-resolution imagery," *IEEE Geoscience and Remote Sensing Letters*, vol. 14, no. 9, pp. 1638–1642, 2017.
- [6] G. J. Scott, K. C. Hagan, R. A. Marcum, J. A. Hurt, D. T. Anderson, and C. H. Davis, "Enhanced fusion of deep neural networks for classification of benchmark high-resolution image data sets," *IEEE Geoscience and Remote Sensing Letters*, vol. 15, no. 9, pp. 1451–1455, 2018.
- [7] D. T. Anderson, G. J. Scott, M. Islam, B. Murray, and R. Marcum, "Fuzzy choquet integration of deep convolutional neural networks for remote sensing," in *Computational Intelligence for Pattern Recognition*. Springer, 2018, pp. 1–28.
- [8] X. Du and A. Zare, "Multiple instance choquet integral classifier fusion and regression for remote sensing applications," *IEEE Transactions on Geoscience and Remote Sensing*, vol. 57, no. 5, pp. 2741–2753, 2019.
- [9] C. A. Dias, J. C. S. Bueno, E. N. Borges, G. Lucca, H. Santos, G. P. Dimuro, H. Bustince, P. L. J. D. Junior, S. S. C. Botelho, and E. Palmeira, "Simulating the behaviour of choquet-like (pre) aggregation functions for image resizing in the pooling layer of deep learning networks," in *International Fuzzy Systems Association World Congress*. Springer, 2019, pp. 224–236.
- [10] K. Simonyan and A. Zisserman, "Very deep convolutional networks for large-scale image recognition," *arXiv 1409.1556*, 2014.
- [11] A. Krizhevsky, "Learning multiple layers of features from tiny images." *Master's thesis, Department of Computer Science, University of Toronto*, 2009.
- [12] J. Keller, D. Liu, and D. Fogel, *Fundamentals of Computational Intelligence: Neural Networks, Fuzzy Systems, and Evolutionary Computation*. Wiley, 2016.
- [13] S. R. Price, S. R. Price, and D. T. Anderson, "Introducing fuzzy layers for deep learning," in *2019 IEEE International Conference on Fuzzy Systems (FUZZ-IEEE)*, 2019, pp. 1–6.
- [14] E. Shelhamer, J. Long, and T. Darrell, "Fully convolutional networks for semantic segmentation," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 39, no. 4, pp. 640–651, 2017.
- [15] Y. LeCun, B. Boser, J. S. Denker, D. Henderson, R. E. Howard, W. Hubbard, and L. D. Jackel, "Backpropagation applied to handwritten zip code recognition," *Neural Computation*, vol. 1, no. 4, pp. 541–551, 1989.
- [16] Y. Lecun, L. Bottou, Y. Bengio, and P. Haffner, "Gradient-based learning applied to document recognition," *Proceedings of the IEEE*, vol. 86, no. 11, pp. 2278–2324, 1998.
- [17] C. E. Shannon, "A mathematical theory of communication," *The Bell System Technical Journal*, vol. 27, no. 3, pp. 379–423, 1948.
- [18] L. I. Rudin, S. Osher, and E. Fatemi, "Nonlinear total variation based noise removal algorithms," *Physica D: Nonlinear Phenomena*, vol. 60, no. 1, pp. 259 – 268, 1992.
- [19] S. Liu and W. Deng, "Very deep convolutional neural network based image classification using small training sample size," in *2015 3rd IAPR Asian Conference on Pattern Recognition (ACPR)*, 2015, pp. 730–734.
- [20] L. N. Smith, "A disciplined approach to neural network hyperparameters: Part 1 – learning rate, batch size, momentum, and weight decay," *arXiv 1803.09820*, 2018.
- [21] K. He, X. Zhang, S. Ren, and J. Sun, "Identity mappings in deep residual networks," *arXiv 1603.05027*, 2016.
- [22] H. B. Mann and D. R. Whitney, "On a test of whether one of two random variables is stochastically larger than the other," *The annals of mathematical statistics*, pp. 50–60, 1947.

Learning OWA weights by combining fuzzy quantifiers with empirical data

Álvaro Cristóbal
University of Zaragoza
Zaragoza, Spain
alvcrimar@gmail.com

Ignacio Huitzil
Artificial Intelligence Research Institute (IIIA-CSIC)
Bellaterra, Spain
ihuitzil@iiia.csic.es

Fernando Bobillo
I3A, University of Zaragoza
Zaragoza, Spain
fbobillo@unizar.es

Abstract—Ordered Weighted Averaging (OWA) is a popular family of aggregation operators that has been used in many practical applications. This paper addresses the problem of learning the weighting vector of OWA operators and uses a hybrid approach combining a sample learning method and a function-based method. The idea is to search for the parameters of a fuzzy quantifier that minimizes the error on a given set of examples. We also perform an experimental study in the field of smart cities.

Index Terms—aggregation, fuzzy quantifiers, smart cities

I. INTRODUCTION

Aggregation of different criteria into a single value is a very common operation in many real-world applications. For example, online travel agencies typically provide a simple way for customers to compare hotels by combining the scores obtained in different criteria (such as location, price, or cleanness) into a single value. The interest in aggregation operators is not expected to decrease in the next years. For instance, in the field of smart cities, with high numbers of sensors providing pieces of information that need to be combined somehow, aggregation operators seem crucial.

Ordered Weighted Averaging (OWA) operators [1] are a very popular family of aggregation operators that has been successfully used in many applications [2]. OWA operators are parameterized with a vector of weights. While the choice of the weights is critical in the behaviour of the operators, determining the concrete values is a common problem in practice. Among the many existing solutions, we are interested in quantifier-guided aggregation [3], where the weights are computed from fuzzy quantifiers.

To illustrate the interest in quantifier-guided aggregation, we will mention some examples in the field of fuzzy ontologies, which are fuzzy extensions of the current de-facto standard for knowledge representation. For example, the fuzzy ontology language *Fuzzy OWL 2* [4] and the fuzzy ontology reasoner *fuzzyDL* [5] support quantifier-guided aggregation using right-shoulder and linear functions. Some recent applications using fuzzy ontologies also take advantage of quantifier-guided aggregation. In particular, the beer recommender system *GimmeHop* [6] and *Fuzzy BIM* [7] support flexible queries about

beers and a semantic Building Information Modeling (BIM), respectively. In both cases, user preferences can be combined using OWA operators, built in a transparent way for the user from right-shoulder, linear, and power functions.

Although there have been many approaches to determine the weights of an OWA operator, the comparisons between the existing alternatives are mostly theoretical. Instead, it would be interesting to evaluate the performance of different methods on real-world datasets. In this paper, we use empirical data to learn the parameters of different functions (fuzzy quantifiers) used in quantifier-based aggregation. Therefore, our method can be seen as a combination of a sample learning method and a function-based method [8]. Moreover, we evaluate the behaviour of our learning strategy using smart cities data.

The remaining of this paper is organized as follows. Section II provides some background on aggregation operators and quantifier-based aggregation. Then, Section III describes our approach, and Section IV discusses the result of our empirical evaluation. Finally, Section V sets out some conclusions and ideas for future work.

II. AGGREGATION OPERATORS

Aggregation Operators (AOs) are mathematical functions that are used to combine different pieces of information (typically, membership degrees to fuzzy sets) [9], [10]. There is no standard definition of AO. Following [9], an AO takes n numerical values $x_1, x_2, \dots, x_n$ (the values of n different criteria) and returns another numerical value, i.e., given a domain $\mathbb{D}$ (such as $[0, 1]$ or $\mathbb{R}$), an AO of dimension n is a mapping $@ : \mathbb{D}^n \rightarrow \mathbb{D}$. A classical example of AO is the weighted mean. We will write $@_{\mathbf{W}}$ to denote the usual case where an AO is parameterized with a vector of n weights $\mathbf{W} = [w_1, \dots, w_n]$ such that $w_i \in [0, 1]$ and $\sum_{i=1}^n w_i = 1$.

A very important family of AOs are the *Ordered Weighted Averaging* (OWA) operators [1]. OWA operators provide a parameterized class of mean type AOs. Formally, given a weighting vector $\mathbf{W}$, an OWA operator of dimension n is an AO such that:

$$@_{\mathbf{W}}^{\text{OWA}}(x_1, \dots, x_n) = \sum_{i=1}^n w_i x_{\sigma(i)} \quad (1)$$

where σ is a permutation such that $x_{\sigma(1)} \geq x_{\sigma(2)} \geq \dots \geq x_{\sigma(n)}$, i.e., $x_{\sigma(i)}$ is the i -th largest of the values $x_1, \dots, x_n$

We were partially supported by the projects TIN2016-78011-C4-3-R (AEI/FEDER, UE), PID2020-113903RB-I00 (AEI/FEDER, UE), and DGA/FEDER.

to be aggregated. Note that, because of this reordering step, a weight w_i is not associated with a specific argument but with an ordered position of the aggregate. By choosing different weights, OWA operators can implement different AOs, such as arithmetic mean, k-th maximum, k-th minimum, median or order statistic, among others.

Example 1: The average temperatures of Barcelona, Basel, Logroño, Madrid, and Milan in February 2018 are 6.7, 1.4, 5.5, 6, and 7.5 °C, respectively. Given the weighting vector $\mathbf{W} = [0.0048, 0.9952, 0, 0, 0]$, the aggregated value using OWA $@_{\mathbf{W}}^{\text{owa}}(6.7, 1.4, 5.5, 6, 7.5)$ is given by:

$$0.0048 \cdot 7.5 + 0.9952 \cdot 6.7 + 0 \cdot 6 + 0 \cdot 5.5 + 0 \cdot 1.4 = 6.70384 \blacksquare$$

A common practical problem is how to compute the weights of an OWA operator, and several solutions have been proposed in the literature [8], [11], [12]. According to X. Liu, existing approaches can be classified in 5 categories [8]: optimization-based methods, sample learning methods fitting to empirical data, function-based methods, argument dependent methods, and preference methods.

The family of function-based methods include methods to build the weights from an orness value. For example, the vector of weights $\mathbf{W}$ can be defined starting from a desired value for the orness in two recursive ways, a Left Recursive Form and a Right Recursive Form [13], or using Faulhaber's formulas [14]. However, the most popular example of function-based methods is *quantifier-based aggregation*.

In quantifier-based aggregation, the vector of weights $\mathbf{W}$ can be defined using a fuzzy quantifier $Q : [0, 1] \rightarrow [0, 1]$. We will focus on *Regular Increasing Monotone* (RIM) quantifiers [15], characterized by the idea that as the proportion increases, the degree of satisfaction does not decrease. More formally, RIMs satisfy the boundary conditions $Q(0) = 0$ and $Q(1) = 1$, and are monotone increasing, i.e., $x_1 \leq x_2$ implies $Q(x_1) \leq Q(x_2)$. A RIM Q can be used to define an OWA weighting vector W_Q of dimension n , where each weight is computed as follows:

$$w_i = Q\left(\frac{i}{n}\right) - Q\left(\frac{i-1}{n}\right) \quad (2)$$

Note that indeed $w_i \in [0, 1]$ and $\sum_i w_i = 1$.

In this paper, we will consider the following functions to build RIMs:

- *Right-shoulder* (or *window* [16]), illustrated in Figure 1 (a). Given $q_1, q_2 \in [0, 1]$ such that $q_1 < q_2$:

$$\mathbf{right}(q_1, q_2) = \begin{cases} 0 & x \leq q_1 \\ \frac{x - q_1}{q_2 - q_1} & x \in [q_1, q_2] \\ 1 & x \geq q_2 \end{cases} \quad (3)$$

If $q_1 = q_2 \neq 1$, we have a *step* function [16]:

$$\mathbf{right}(q_1, q_2) = \begin{cases} 0 & x \leq q_1 \\ 1 & x > q_1 \end{cases} \quad (4)$$

If $q_1 = q_2 = 1$, we also have a *step* function:

$$\mathbf{right}(q_1, q_2) = \begin{cases} 0 & x < q_1 \\ 1 & x \geq q_1 \end{cases} \quad (5)$$

- *Linear*, illustrated in Figure 1 (b). Given $q_1, q_2 \in [0, 1]$ with $q_1 \in (0, 1)$:

$$\mathbf{linear}(q_1, q_2) = \begin{cases} \frac{q_2}{q_1} \cdot x & x \leq q_1 \\ \frac{(1 - q_2)x + (q_2 - q_1)}{1 - q_1} & x > q_1 \end{cases} \quad (6)$$

If $q_1 = 0$:

$$\mathbf{linear}(q_1, q_2) = \begin{cases} 0 & x = 0 \\ (1 - q_2)x + q_2 & x > 0 \end{cases} \quad (7)$$

If $q_1 = 1$:

$$\mathbf{linear}(q_1, q_2) = \begin{cases} q_2 \cdot x & x < 1 \\ 1 & x = 1 \end{cases} \quad (8)$$

- *Power*, illustrated in Figure 1 (c). Given $q \in (0, \infty)$:

$$\mathbf{power}(q) = x^q \quad (9)$$

Example 2: The weighting vector in Example 1, with $n = 5$ weights, can be computed from the quantifier $Q = \mathbf{right}(0.1991, 0.3851)$:

- $w_1 = Q(1/5) - Q(0) = 0.0048 - 0 = 0.0048$
- $w_2 = Q(2/5) - Q(1/5) = 1 - 0.0048 = 0.9952$
- $w_3 = Q(3/5) - Q(2/5) = 1 - 1 = 0$
- $w_4 = Q(4/5) - Q(3/5) = 1 - 1 = 0$
- $w_5 = Q(5/5) - Q(4/5) = 1 - 1 = 0 \blacksquare$

III. LEARNING THE PARAMETERS OF THE FUZZY QUANTIFIERS

We assume that we have a set of examples $E = \langle e_1, \dots, e_m \rangle$. Each example e_j contains the values of n input variables (x_{ij}) and the value of an output variable (y_j) :

$$e_j = \langle x_{1j}, x_{2j}, \dots, x_{nj}, y_j \rangle \quad (10)$$

For each e_j , we compute aggregated values using different weighting vectors $\mathbf{W}$:

$$z_j^{\mathbf{W}} = @_{\mathbf{W}}^{\text{owa}}(x_{1j}, x_{2j}, \dots, x_{nj}) \quad (11)$$

For each $\mathbf{W}$, we compute the error between the expected value y_j and the aggregated value $z_j^{\mathbf{W}}$. In particular, we consider the Mean Absolute Percentage Error (MAPE), which is computed as a percentage in $[0, 100]$ where the smaller the percentage, the smaller the error:

$$\text{MAPE}(E, \mathbf{W}) = \frac{100}{m} \sum_{j=1}^m \left| \frac{y_j - z_j^{\mathbf{W}}}{y_j} \right| \quad (12)$$

Finally, we choose the quantifier type and parameters that lead to the weighting vector $\mathbf{W}$ minimizing the MAPE:

$$\underset{\mathbf{W}}{\text{argmin}} \text{MAPE}(E, \mathbf{W}) \quad (13)$$

Example 3: Revisiting Example 1, assume that we want to predict the temperature in Zaragoza from the temperatures of the other 5 cities. If the expected value of the

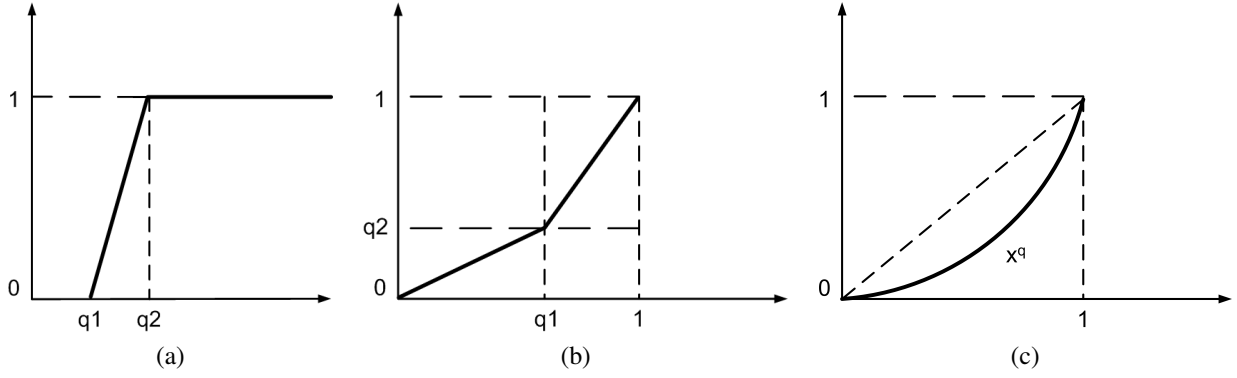


Fig. 1. (a) Right-shoulder function; (b) Linear function; (c) Power function

temperature in Zaragoza in February 2018 is 6.8 °C, then the term of the summation corresponding to this example is $|(6.8 - 6.70384)/6.70384| = 0.0143$. ■

The key of our approach is that W is built using different types of RIM functions and with different parameters. If the search space is not very large, and computing the weights and the aggregation using OWA are not computationally expensive operations. Therefore, it could be possible to compute the best parameters by brute force, at least for not very large datasets.

- Algorithm 1 shows how to compute the parameters of a linear RIM quantifier that minimizes the MAPE using brute force. The algorithm loops over the values of q_1 and q_2 by adding an increment Δ . The best MAPE is initialized to 100, the highest possible value. When a smaller MAPE is found, the parameters q_1 and q_2 that made it possible are stored. Finally, the pair of parameters $\langle q_1, q_2 \rangle$ is returned.
- A brute-force algorithm to compute the parameters of a right-shoulder RIM quantifier is similar, but the loop over all values of q_2 starts from q_1 so that $q_1 \leq q_2$.
- Finally, for the power function, a single loop is needed, as there is just one parameter q , ranging in $(0, \infty)$.

As a final remark, note that we do not split out set of examples E into training and test sets.

For large datasets, it could be possible to use heuristic methods, such as Monte Carlo algorithms, local search, or evolutionary algorithms. For example, Algorithm 2 shows how to compute the parameters of a linear RIM quantifier using a Monte Carlo algorithm. The idea is to generate pseudo-random numbers as the possible values of q_1 and q_2 , repeating the experiments several times, and storing the values that minimize the MAPE.

To conclude this section, let us note that it is trivial to consider the case where there are missing values, i.e., for some examples e_j some of the values x_{ij} are unknown. In this case, rather than having the same vector W for all the examples, we could compute for each example e_j a weighting vector of dimension n_j , where n_j is the number of non-missing values for the input variables, using the same function type and parameters for all the examples.

Algorithm 1 Brute-force algorithm to compute the parameters of a LINEAR quantifier minimizing the MAPE.

Input: A dataset E with examples as in Equation 10.

Output: Parameters $\langle q_1, q_2 \rangle$ of a linear function.

```

1:  $bestMapeL \leftarrow 100$ 
2: for  $q_1 \leftarrow 0$  to 1 by  $\Delta$  do
3:   for  $q_2 \leftarrow 0$  to 1 by  $\Delta$  do
4:      $Q \leftarrow \text{linear}(q_1, q_2)$ 
5:      $W \leftarrow \text{compute a vector from } Q \text{ using Eq. 2}$ 
6:      $mape \leftarrow \text{compute MAPE}(E, W) \text{ using Eq. 12}$ 
7:     if  $mape < bestMapeL$  then
8:        $bestMapeL \leftarrow mape$ 
9:        $bestQ1 \leftarrow q_1$ 
10:       $bestQ2 \leftarrow q_2$ 
11:    end if
12:  end for
13: end for
14: return  $\langle bestQ1, bestQ2 \rangle$ 

```

IV. EVALUATION

This section discusses an evaluation of our approach in the field of smart cities.

a) Datasets: We consider two datasets: Temperatures and Tourism.

- The *Temperatures* dataset includes the temperatures in 7 cities: Barcelona, Basel, Buenos Aires, Logroño, Madrid, Milan, and Zaragoza. Note that all of them are cities in Northern hemisphere except Buenos Aires, which thus has opposite seasons. Table 1 shows the URLs where the temperatures were retrieved. The dataset has 12 rows, with the monthly average temperatures on one year.
- The *Tourism* dataset includes information about London¹: the total number of visits, the total number of nights, and the total spend. The values cover the period 2002–2019 (so the effects of the COVID-19 pandemic are

¹<http://data.london.gov.uk/dataset/number-international-visitors-london>

TABLE I
URLS WITH THE TEMPERATURES OF EACH SMART CITY

City	URL
Barcelona	http://opendata-ajuntament.barcelona.cat/data/es/dataset/temperatures-hist-bcn
Basel	https://www.meteoblue.com/es/tiempo/archive/export/basilea_suiza_2661604
Buenos Aires	http://data.buenosaires.gob.ar/dataset/registro-temperatura-ciudad
Logroño	http://datos.gob.es/en/catalogo/a17002943-estaciones-meteorologicas-sos-rioja1
Madrid	http://es.climate-data.org/europe/espana/comunidad-de-madrid/madrid-92/enero-1
Milan	http://dati.comune.milano.it/dataset/ds305-ambientemeteo-temperature-mese-2008-2014
Zaragoza	http://datosclima.es/Aemet2013/Temperatura2013.php

Algorithm 2 Monte Carlo algorithm to compute the parameters of a linear quantifier minimizing the MAPE.

Input: A dataset E with examples as in Equation 10, and the number of repetitions MAX_REPETITIONS.

Output: Parameters $\langle q_1, q_2 \rangle$ of a linear function.

```

1:  $bestMapeL \leftarrow 100$ 
2:  $repetition \leftarrow 0$ 
3: repeat
4:    $q_1 \leftarrow$  random number in  $[0, 1]$ 
5:    $q_2 \leftarrow$  random number in  $[0, 1]$ 
6:    $Q \leftarrow \text{linear}(q_1, q_2)$ 
7:    $W \leftarrow$  compute a vector from  $Q$  using Eq. 2
8:    $mape \leftarrow$  compute  $MAPE(E, W)$  using Eq. 12
9:   if  $mape < bestMapeL$  then
10:     $BestMapeL \leftarrow mape$ 
11:     $BestQ1 \leftarrow q_1$ 
12:     $BestQ2 \leftarrow q_2$ 
13:   end if
14:    $repetition \leftarrow repetition + 1$ 
15: until  $repetition = \text{MAX\_REPETITIONS}$ 
16: return  $\langle bestQ1, bestQ2 \rangle$ 

```

not observed) and are aggregated by quarters. Therefore, there are 72 rows.

b) Experiments: We consider two experiments with the Temperatures dataset, and three with the Tourism dataset:

- E1. Prediction of the temperature in Zaragoza from the temperatures of 5 cities: Barcelona, Basel, Logroño, Madrid, and Milan.
- E2. Prediction of the temperature in Zaragoza from the temperatures of 6 cities: Barcelona, Basel, Buenos Aires, Logroño, Madrid, and Milan. This experiment is similar to E1 but taking Buenos Aires into account.
- E3. Prediction of the total number of visits from the total number of nights and the total spend.
- E4. Prediction of the total number of nights from the total number of visits and the total spend.
- E5. Prediction of the total spend from the total number of visits and the total number of nights.

In the Tourism dataset we apply a normalization step, since the three variables have a different range of values. For each

variable, we divide each value by the maximum value plus a 5 %, obtaining a value in $[0, 1]$.

c) Parameters:

- In our brute-force algorithms, we use increments $\Delta = 0.001$ and $\Delta = 0.0001$.
- For the power quantifier, we take 20 as an upper bound for the value of the parameter, i.e. $q \in (0, 20]$. This choice was made after checking experimentally that higher values produce very small changes in a vector of 5 weights.
- In our Monte Carlo algorithms, we repeat the experiments $5 \cdot 10^5$ times. This choice was made to have a similar running time as in the brute-force algorithm with $\Delta = 0.001$. We also noticed that in different runs of the algorithm, the MAPE did not change if we rounded to two decimals.

d) Environmental setup: Our code was implemented in Java 1.8. All experiments were performed on a laptop computer with Intel Core i7-8750H, 16 GB RAM, 1 TB HDD + 256 GB SSD under Windows 10, 64-bits.

e) Results: Table II includes the results. For each experiment, for each algorithm, and for each quantifier type, we show the best MAPE, the best parameters, and the corresponding weighting vector. For each experiment and algorithm, we show the total running time (in seconds) to optimize the parameters of all quantifier types.

Figures 2, 3, and 4 illustrate a summary of the results for the brute-force algorithm in experiment E1, for different types of quantifiers (right-shoulder, linear, and power, respectively) and an increment $\Delta = 0.1$. MAPE values are rounded to the next integer. In Figure 4 the rows and the columns indicate the integer part and the fractional part of q , respectively, and the integer part is not shown if higher than 10.

f) Discussion: To start with, it is worth to clarify that our main objective is not to solve some prediction problems but to find the parameters corresponding to the best OWA weights. Indeed, to solve these prediction problems, more complex machine learning strategies are very likely to perform with better results.

The first interesting observation is that the MAPE is very similar regardless of the algorithm for a given quantifier type. In all the experiments with the Tourism dataset (E3, E4, and E5), the MAPE is actually the same (and so is the weighting vector) regardless of the quantifier type. In the other experiments, the differences in the MAPE are smaller than 0.007.

TABLE II
RESULTS OF THE EXPERIMENTS

Experiment	Algorithm	Quantifier	MAPE	Parameters	Weights	Time
E1	Brute-force $\Delta = 0.0001$	Right-shoulder	5.3888	0.1991 0.3851	[0.0048 0.9952 0 0 0]	47
		Linear	9.4604	0.4434 0.9319	[0.4203 0.4203 0.1104 0.0245 0.0245]	
		Power	10.0735	0.3014	[0.6156 0.143 0.0986 0.0777 0.065]	
	Brute-force $\Delta = 0.001$	Right-shoulder	5.3956	0.199 0.4	[0.005 0.995 0 0 0]	0.6
		Linear	9.4614	0.443 0.932	[0.4208 0.4208 0.1096 0.0244 0.0244]	
		Power	10.0741	0.301	[0.616 0.1429 0.0985 0.0776 0.065]	
	Monte Carlo	Right-shoulder	5.3893	0.1995 0.3007	[0.0048 0.9952 0 0 0]	0.6
		Linear	9.461	0.4438 0.9327	[0.4203 0.4203 0.1110 0.0242 0.0242]	
		Power	10.0736	0.3013	[0.6157 0.143 0.0986 0.0776 0.065]	
E2	Brute-force $\Delta = 0.0001$	Right-shoulder	5.2896	0.166 0.513	[0.0019 0.4803 0.4803 0.0375 0 0]	60
		Linear	12.6139	0.0836 0.0082	[0.0981 0.1804 0.1804 0.1804 0.1804 0.1804]	
		Power	13.4133	102.431	[0.1078 0.1474 0.1673 0.1816 0.1931 0.2028]	
	Brute-force $\Delta = 0.001$	Right-shoulder	5.2896	0.166 0.513	[0.0019 0.4803 0.4803 0.0375 0 0]	0.7
		Linear	12.6139	0.125 0.053	[0.0981 0.1804 0.1804 0.1804 0.1804 0.1804]	
		Power	13.4136	10.243	[0.1078 0.1474 0.1673 0.1816 0.1931 0.2028]	
	Monte Carlo	Right-shoulder	5.29	0.166 0.5129	[0.0019 0.4805 0.4805 0.0372 0 0]	0.7
		Linear	12.6139	0.1447 0.0743	[0.0981 0.1804 0.1804 0.1804 0.1804 0.1804]	
		Power	13.4133	1.2431	[0.1078 0.1474 0.1673 0.1816 0.1931 0.2028]	
E3	Brute-force $\Delta = 0.0001$	Right-shoulder	9.007	0.1492 0.5278	[0.9266 0.0734]	112
		Linear	9.007	0.0535 0.861	[0.9266 0.0734]	
		Power	9.007	0.11	[0.9266 0.0734]	
	Brute-force $\Delta = 0.001$	Right-shoulder	9.007	0.235 0.521	[0.9266 0.0734]	1.5
		Linear	9.007	0.142 0.874	[0.9266 0.0734]	
		Power	9.007	0.11	[0.9266 0.0734]	
	Monte Carlo	Right-shoulder	9.007	0.3208 0.5142	[0.9266 0.0734]	1.3
		Linear	9.007	0.4624 0.921	[0.9266 0.0734]	
		Power	9.007	0.11	[0.9266 0.0734]	
E4	Brute-force $\Delta = 0.0001$	Right-shoulder	9.8532	0.1265 0.501	[0.9973 0.0027]	105
		Linear	9.8532	0.0263 0.9948	[0.9973 0.0027]	
		Power	9.8532	0.0038	[0.9974 0.0026]	
	Brute-force $\Delta = 0.001$	Right-shoulder	9.8532	0.126 0.501	[0.9973 0.0027]	1.4
		Linear	9.8532	0.251 0.996	[0.9973 0.0027]	
		Power	9.8533	0.004	[0.9972 0.0028]	
	Monte Carlo	Right-shoulder	9.8532	0.0612 0.5012	[0.9973 0.0027]	1.4
		Linear	9.8532	0.1962 0.9957	[0.9973 0.0027]	
		Power	9.8532	0.0039	[0.9973 0.0027]	
E5	Brute-force $\Delta = 0.0001$	Right-shoulder	16.582	0.5 0.5	[0 1]	108
		Linear	16.582	0.5 0	[0 1]	
		Power	16.582	20	[0 1]	
	Brute-force $\Delta = 0.001$	Right-shoulder	16.582	0.5 0.5	[0 1]	1.3
		Linear	16.582	0.5 0	[0 1]	
		Power	16.582	20	[0 1]	
	Monte Carlo	Right-shoulder	16.582	0.6822 0.842	[0 1]	1.2
		Linear	16.582	0.5584 0	[0 1]	
		Power	16.582	20	[0 1]	

Brute-force with $\Delta = 0.0001$ always has the smaller MAPE. Brute-force with a higher increment is slightly worse in 5 cases (in E1 and in E2 for two quantifier types) and Monte Carlo is slightly worse in 4 cases (in E1 and in E2 for one quantifier type). If we compare the two worse algorithms, Monte Carlo wins strictly speaking in 4 cases and loses in 1, but the MAPEs are actually equal if we round to 2 decimals.

However, reasoning times can be more different. The fastest algorithms are brute-force with $\Delta = 0.001$ and Monte Carlo algorithm, both of them with almost the same time. On the other hand, brute-force with $\Delta = 0.0001$ is pretty much slower, and the increase in the running time is not compensated with a significant decrease in the MAPE.

If we compare the quantifier types, right-shoulder is always the best function in the Temperatures dataset, whereas in the Tourism dataset the same MAPE is always obtained regardless

of the quantifier type.

Now let us give a closer look to each dataset. In the Temperatures dataset, the MAPE is in general clearly smaller in E1 than in E2. This result is expected, because E2 introduces data (from Buenos Aires) making the prediction harder. However, experiments with the right-shoulder function are an exception, and the MAPE is actually slightly smaller in E2. We can notice that in these cases the weight associated to the smallest value to be aggregated is 0, and the weight associated to the largest value is very small, about 0.002. Because Buenos Aires is the only city from the Southern hemisphere, its temperature will usually be either the highest or the lowest one, but it will have a small influence in the aggregated value because of the weights. Note that after considering Buenos Aires, the size of the weighting vector increases. Therefore, to build the weighting vectors, the quantifier functions are evaluated

q1q2	0	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9	1
0	13	13	13	11	10	10	11	12	14	16	19
0.1		13	13	10	9	10	12	13	15	18	20
0.2			13	6	6	9	11	13	15	18	21
0.3				6	6	11	13	15	17	21	24
0.4					6	17	17	19	20	25	28
0.5						17	17	20	22	27	31
0.6							17	25	25	31	36
0.7								25	25	36	42
0.8									25	66	66
0.9										66	66
1											66

Fig. 2. MAPE for right-shoulder quantifiers in E1

q1q2	0	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9	1
0	19	17	15	14	12	11	10	11	11	12	13
0.1	20	19	17	15	13	11	10	10	11	12	13
0.2	21	20	19	17	15	13	11	10	11	12	13
0.3	24	22	20	19	17	15	13	11	10	10	11
0.4	28	25	23	21	19	16	14	12	10	9	10
0.5	31	27	25	23	21	19	16	14	13	11	10
0.6	36	30	28	26	23	21	19	16	14	12	11
0.7	42	35	32	29	26	23	21	19	16	14	12
0.8	66	50	42	36	31	28	24	21	19	16	14
0.9	66	51	43	38	33	30	26	24	21	19	16
1	66	53	45	39	35	31	28	26	23	21	19

Fig. 3. MAPE for linear quantifiers in E1

in the points $\{0, 0.2, 0.4, 0.6, 0.8, 1\}$ rather than in the points $\{0, 0.25, 0.5, 0.75, 1\}$.

In the Tourism dataset, the best strategy seems using almost exclusively one of the two input variables. Indeed, in E5, the weighting vector is $[0, 1]$. In the other experiments E3 and E4, the highest value to be aggregated has a weight greater than 0.9. The MAPE is worse than in the Temperatures dataset, and is particularly high when predicting the total spend (E5).

V. CONCLUSIONS AND FUTURE WORK

In this paper we have followed a hybrid approach to learn the weights of OWA operators by choosing the parameters of some functions, commonly used in quantifier-guided aggregation, that minimize the error over a set of samples. We studied

q	.0	.1	.2	.3	.4	.5	.6	.7	.8	.9
0		12	11	10	11	13	14	15	16	18
1	19	19	20	21	22	22	23	23	24	24
2	25	25	26	26	27	28	28	29	30	31
3	31	32	33	33	34	34	35	35	36	36
4	37	37	38	38	38	39	39	40	40	40
5	41	41	42	42	42	43	43	43	44	44
6	44	45	45	45	46	46	46	47	47	47
7	47	48	48	48	49	49	49	49	50	50
8	50	50	51	51	51	51	52	52	52	52
9	53	53	53	53	54	54	54	54	54	55
10	55	55	55	55	56	56	56	56	56	57

Fig. 4. MAPE for power quantifiers in E1

right-shoulder, linear, and power functions, and proposed two alternatives for the search on the parameter space: brute force and a Monte Carlo algorithm.

Our approach has several advantages over learning the OWA weights directly. On the one hand, the number of parameters is smaller, reducing the search space. On the other hand, the results of more interpretable, as fuzzy quantifiers can be more easily understood by humans.

We have also discussed the results of an empirical evaluation on two datasets in the field of smart cities. We found significant differences in the running times of the algorithms but not on the error. Furthermore, we observed that in one dataset the error was smaller for right-shoulder functions, whereas in the other dataset the results were independent on the function type.

There are a lot of directions for our future work. Firstly, it would be interesting to consider more types of RIMS. Secondly, we could consider more complex algorithms to search for the best parameters. We could also take into account alternative approaches than quantifier-guided aggregation. Finally, experiments on more real-world datasets are desirable.

REFERENCES

- [1] R. R. Yager, "On ordered weighted averaging aggregation operators in multicriteria decision making," *IEEE Transactions on Systems, Man and Cybernetics*, vol. 18, no. 1, pp. 183–190, 1988.
- [2] R. R. Yager, J. Kacprzyk, and G. Beliakov, Eds., *Recent developments in the Ordered Weighted Averaging operators: Theory and practice*, ser. Studies in Fuzziness and Soft Computing. Springer, 2011, vol. 265.
- [3] R. R. Yager, "Quantifier guided aggregation using OWA operators," *International Journal of Intelligent Systems*, vol. 11, no. 1, pp. 49–73, 1996.
- [4] F. Bobillo and U. Straccia, "Fuzzy ontology representation using OWL 2," *International Journal of Approximate Reasoning*, vol. 52, no. 7, pp. 1073–1094, 2011.
- [5] —, "The fuzzy ontology reasoner fuzzyDL," *Knowledge-Based Systems*, vol. 95, pp. 12–34, 2016.
- [6] I. Huitzil, F. Alegre, and F. Bobillo, "GimmeHop: A recommender system for mobile devices using ontology reasoners and fuzzy logic," *Fuzzy Sets and Systems*, vol. 401, pp. 55–77, 2020.
- [7] I. Huitzil, M. Molina-Solana, J. Gómez-Romero, and F. Bobillo, "Minimalistic fuzzy ontology reasoning: An application to Building Information Modeling," *Applied Soft Computing*, vol. 103, p. 107158, 2021.
- [8] X. Liu, "A review of the OWA determination methods: Classification and some extensions," in *Recent Developments in the Ordered Weighted Averaging Operators: Theory and Practice*, ser. Studies in Fuzziness and Soft Computing. Springer, 2011, vol. 265, pp. 49–90.
- [9] V. Torra and Y. Narukawa, *Modeling decisions - Information fusion and aggregation operators*. Springer, 2007.
- [10] G. Beliakov, H. Bustince, and T. Calvo, *A practical guide to averaging functions*, ser. Studies in Fuzziness and Soft Computing. Springer, 2016, vol. 329.
- [11] Z. Xu, "An overview of methods for determining OWA weights," *International Journal of Intelligent Systems*, vol. 20, no. 8, pp. 843–865, 2005.
- [12] R. Fullér, "On obtaining OWA operator weights: A sort survey of recent developments," in *Proceedings of the 5th International Conference on Computational Cybernetics (ICCC 2007)*, 2007, pp. 241–244.
- [13] L. Troiano and R. R. Yager, "Recursive and iterative OWA operators," *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, vol. 13, no. 6, pp. 579–600, 2005.
- [14] O. Duarte and S. Téllez, "A family of OWA operators based on Faulhaber's formulas," arXiv 1801.10545, 2018.
- [15] R. R. Yager, "Connectives and quantifiers in fuzzy sets," *Fuzzy Sets and Systems*, vol. 40, no. 1, pp. 39–75, 1991.
- [16] —, "Families of OWA operators," *Fuzzy Sets and Systems*, vol. 59, no. 2, pp. 125–148, 1993.

Integrales de Choquet y capacidades 2-aditivas en la construcción de indicadores sintéticos

Bonifacio Llamazares
Departamento de Economía Aplicada
Universidad de Valladolid
Valladolid
boni@eco.uva.es

Patrizia Pérez-Asurmendi
Departamento de Análisis Económico y Economía Cuantitativa
Universidad Complutense de Madrid
Madrid
patrizip@ucm.es

I. RESUMEN

Los indicadores sintéticos (también denominados compuestos) se han convertido en una herramienta muy valiosa en diversos ámbitos científicos dado que permiten resumir en un único valor la información proporcionada por diversos indicadores. A la hora de elaborar un indicador sintético es conveniente seguir una serie de pasos (véase [1, Tabla 1]), entre los que se encuentra la selección de un procedimiento de agregación. En la elección de dicho procedimiento resulta fundamental tener en cuenta tanto la posible correlación existente entre los indicadores utilizados como la compensabilidad que se desea permitir entre ellos (véanse [2, p. 363], [1, p. 21]).

En este contexto, uno de los métodos más habituales para agregar los valores de los indicadores consiste en utilizar funciones de agregación. Entre la gran variedad de funciones existentes destacan, por su sencillez y propiedades, las medias aritméticas y geométricas.¹ Aunque estas funciones están presentes en la construcción de muchos indicadores sintéticos, su uso ha sido objeto de diversas críticas.

En el caso de la media aritmética, la principal tiene que ver con el hecho de que valores elevados en algunos indicadores pueden compensar valores bajos de otros. Ello ha originado que, en algunos casos, haya sido sustituida por la media geométrica. Por ejemplo, el Índice de Desarrollo Humano (IDH), que hasta el año 2009 se basaba en la media aritmética, se construye desde el año 2010 mediante la media geométrica debido a que esta función permite una menor compensabilidad entre los indicadores que la media aritmética. Por lo que respecta a la media geométrica, sus principales debilidades son la imposibilidad de utilizar el 0 como valor mínimo de referencia (algo bastante habitual en la práctica al transformar los valores de los indicadores mediante determinadas normalizaciones), lo cual dificulta la capacidad informativa de la escala utilizada, y que no es invariante a cambios de origen [3].

En el contexto descrito, la integral de Choquet [4] emerge como un instrumento eficaz para solventar los problemas anteriormente mencionados. Además, permite tener en cuenta

tanto la posible correlación existente entre los indicadores utilizados como la compensabilidad que se desea permitir entre ellos. En la definición de la integral de Choquet juega un papel fundamental el concepto de capacidad. Entre la amplia variedad de capacidades que se pueden construir, las capacidades 2-aditivas [5] son, probablemente, el mejor compromiso entre baja complejidad y riqueza del modelo [6].

En trabajos anteriores (véanse [7], [8]) se ha desarrollado un modelo basado en capacidades 2-aditivas que permite tener en cuenta las interacciones que habitualmente existen entre los indicadores empleados. En el presente trabajo se extiende el modelo anterior de manera que, además de tener en cuenta las interacciones existentes entre los indicadores, es posible regular la compensabilidad que se permite entre ellos.

Como aplicación práctica, el modelo propuesto se aplica al Índice de Sociedad Sostenible (SSI),² reemplazando, como función de agregación, la media geométrica por la integral de Choquet y comparando los resultados obtenidos según ambas agregaciones.

BIBLIOGRAFÍA

- [1] M. Nardo, M. Saisana, A. Saltelli, S. Tarantola, A. Hoffmann, and E. Giovannini, *Handbook on Constructing Composite Indicators: Methodology and User Guide*. OECD Publishing, 2008.
- [2] A. Saltelli, M. Nardo, M. Saisana, and S. Tarantola, "Composite indicators — The controversy and the way forward," in *Statistics, Knowledge and Policy: Key Indicators to Inform Decision Making (Proceedings of the World Forum on Key Indicators)*. Palermo (Italy): OECD, November 10-13, 2004 2005, pp. 359–372.
- [3] M. Grabisch, "The application of fuzzy integrals in multicriteria decision making," *European Journal of Operational Research*, vol. 89, no. 3, pp. 445–456, 1996.
- [4] G. Choquet, "Theory of capacities," *Annales de l'Institut Fourier*, vol. 5, pp. 131–295, 1953.
- [5] M. Grabisch, " k -order additive discrete fuzzy measures and their representation," *Fuzzy Sets and Systems*, vol. 92, no. 2, pp. 167–189, 1997.
- [6] M. Grabisch and C. Labreuche, "A decade of application of the Choquet and Sugeno integrals in multi-criteria decision aid," *Annals of Operations Research*, vol. 175, no. 1, pp. 247–286, 2010.
- [7] B. Llamazares and P. Pérez-Asurmendi, "Criteria interaction and Shapley values," in *XIII Encuentro de la Red Española de Elección Social (REES)*, Granada, 18-19 de noviembre de 2016.
- [8] B. Llamazares and P. Pérez-Asurmendi, "Elaboración de indicadores sintéticos mediante la integral de Choquet," in *XXV Jornadas de ASEPUMA*, La Coruña, 8-9 de junio de 2017.

Patrizia Pérez-Asurmendi agradece el apoyo recibido por la Fundación Areces (Proyecto CISP18A6216).

¹En ambas funciones, la importancia de los indicadores en el indicador sintético se establece con el vector de pesos utilizado en su cálculo.

²Este indicador sintético permite medir la sostenibilidad en un sentido amplio dado que tiene en cuenta tres dimensiones del bienestar: el bienestar humano, el bienestar medioambiental y el bienestar económico.

Extensión multidimensional de la integral de Choquet discreta y su aplicación en redes neuronales recurrentes

Mikel Ferrero-Jaurrieta, Iosu Rodríguez-Martínez, Graçaliz Pereira Dimuro, Javier Fernández, Humberto Bustince

Departamento de Estadística, Informática y Matemáticas

Universidad Pública de Navarra

Campus Arrosadía, s/n, 31006 Pamplona, España

{mikel.ferrero, iosu.rodriguez, gracializ.pereira, fcojavier.fernandez, bustince}@unavarra.es

Abstract—En este trabajo presentamos una definición de la integral de Choquet discreta n -dimensional, para fusionar datos vectoriales. Como aplicación, utilizamos estas nuevas integrales de Choquet discretas multidimensionales en la fusión de información secuencial en las redes neuronales recurrentes, mejorando los resultados obtenidos mediante el método de agregación tradicional.

Index Terms—Integral de Choquet, Función de Agregación, Información Multivariante, Redes Neuronales Recurrentes, LSTM.

I. INTRODUCCIÓN

El proceso de fusión de información resulta un procedimiento fundamental a la hora de combinar o agregar distintas estructuras de información en una sola [1]. Su utilización es necesaria en diversos campos, como por ejemplo: toma de decisión multi-criterio [2], economía y finanzas [3], estadística, procesamiento de imágenes [4], aprendizaje automático [5], etc. Recientemente también se ha aplicado en el aprendizaje profundo, por ejemplo en las capas de *pooling* de redes neuronales convolucionales [6].

En la literatura se analizan multitud de métodos de agregación de información. Algunos de los más habituales son las medias aritméticas ponderadas [7] o los órdenes estadísticos [8]. Sin embargo, frecuentemente los criterios y datos considerados interactúan entre ellos y es conveniente utilizar operadores de agregación que tengan en cuenta este hecho. En este sentido, en la literatura se han utilizado las integrales difusas [9], las cuales se basan en medidas difusas. Estas medidas [10] permiten tener en cuenta la relación existente entre los elementos a agregar, valorando la relevancia de posibles coaliciones entre los datos [5].

Una de las integrales difusas más utilizadas es la integral de Choquet [11]. Hasta ahora, en la literatura se han presentado distintas generalizaciones de la integral de Choquet [12], [13], [5], [14] para datos unidimensionales.

Frecuentemente los elementos a agregar son datos o instancias con varias variables, es decir, información multivariante, estructurada en forma de vectores. Por ejemplo, un modelo utilizado actualmente en la inteligencia artificial donde se manejan datos multivariantes son las Redes Neuronales

Recurrentes [15].

Las Redes Neuronales Recurrentes son un tipo de red neuronal artificial que se encargan de modelar información de tipo secuencial o temporal, como las series temporales o el procesamiento de lenguaje natural [16]. Dichas redes constan de una arquitectura en la que en cada instante, los valores de salida de la capa del instante anterior se conectan con la información del instante actual. Para la conexión de dichos datos multivariantes, usualmente se suelen sumar los vectores como forma de agregación de la información multivariante secuencial. Entre los datos recurrentes generados por la red y los datos provenientes del dataset puede haber interacción entre los mismos.

El objetivo de este trabajo es presentar una extensión multidimensional de la integral de Choquet, es decir, una función que agregue m datos n -dimensionales, teniendo en cuenta las posibles coaliciones entre los mismos.

Para mostrar la utilidad de nuestra extensión, presentamos su uso en la modelización de las posibles coaliciones entre datos en el proceso de agregación de una red neuronal recurrente. En concreto, la modificación de la arquitectura que presentamos en este trabajo consiste en una red neuronal recurrente del tipo memoria de corto y largo plazo (LSTM) [17]. En los pasos en los que se agrega la información recurrente con la información inicial estudiamos la utilización de la integral de Choquet discreta multidimensional. De forma análoga, completamos el estudio con el uso de combinaciones lineales de otras funciones de agregación.

Este trabajo se organiza de la siguiente manera: en la Sección II recordamos los conceptos preliminares necesarios para comprender el resto del trabajo. En la Sección III se introduce la nueva definición de la integral de Choquet extendida a datos multivariantes. En la Sección IV se introduce la modificación en la arquitectura. En la Sección V se presenta la experimentación realizada. Por último, en la Sección VI, se explican las conclusiones así como las líneas futuras del trabajo.

II. PRELIMINARES

En esta sección, recordamos nociones básicas y terminología necesarias para abordar el desarrollo del trabajo. Presentamos por un lado las definiciones teóricas básicas de función de agregación, medida difusa e integral de Choquet, y por otro la explicación del funcionamiento de una red neuronal recurrente de tipo LSTM.

A. Funciones de agregación e integral de Choquet

Consideremos el retículo $(L, \leq)$ donde $L = [0, 1]$ y sea $\leq$ el orden natural de los números reales. Denotamos $\mathbf{0} = (0, \dots, 0) \in L^n$ y $\mathbf{1} = (1, \dots, 1) \in L^n$.

Definición II.1. Sea m un entero positivo. Una función $M : L^m \rightarrow L$ es una función de agregación m -aria si satisface las siguientes propiedades:

- (i) $M(\mathbf{0}) = 0$ y $M(\mathbf{1}) = 1$
- (ii) es no-decreciente en cada variable, es decir, para todo $(x_1, \dots, x_m), (y_1, \dots, y_m) \in L^m$, $M(x_1, \dots, x_m) \leq M(y_1, \dots, y_m)$ si $x_1 \leq y_1, \dots, x_m \leq y_m$.

Denotamos el conjunto $\{1, \dots, m\}$ por $[m]$. Dos vectores $(x_1, \dots, x_m), (y_1, \dots, y_m) \in L^m$ son considerados comonótonos si y sólo si existe una permutación $\sigma : [m] \rightarrow [m]$ tal que $x_{\sigma(1)} \leq \dots \leq x_{\sigma(m)}$ e $y_{\sigma(1)} \leq \dots \leq y_{\sigma(m)}$.

Denotamos con letras negrita los elementos en L^n , esto es, $\mathbf{x} = (x_1, \dots, x_n) \in L^n$. Existe un orden parcial $\leq_P$ inducido por $\leq$ y dado de la siguiente manera:

$$\mathbf{x} \leq_P \mathbf{y} \text{ si y sólo si } x_i \leq y_i$$

para todo $i \in \{1, \dots, n\}$.

De hecho, podemos verificar que $(L^n, \leq)$ es un retículo donde el elemento mínimo es $\mathbf{0}$ y el máximo es $\mathbf{1}$. En este retículo el ínfimo y el supremo de dos elementos vienen dados por las siguientes operaciones:

$$\mathbf{x} \wedge \mathbf{y} = (\min(x_1, y_1), \dots, \min(x_n, y_n)) \quad (1)$$

$$\mathbf{x} \vee \mathbf{y} = (\max(x_1, y_1), \dots, \max(x_n, y_n)) \quad (2)$$

Con anterioridad a la definición del concepto de integral de Choquet, consideramos la definición de medida difusa.

Definición II.2. [18] Una medida difusa definida sobre $[m]$ es una función $\nu : 2^{[m]} \rightarrow L$ tal que:

- (i) $\nu(\emptyset) = 0$ y $\nu([m]) = 1$
- (ii) $\nu(\mathcal{A}) \leq \nu(\mathcal{B})$ para todo $\mathcal{A} \subseteq \mathcal{B} \subseteq [m]$

Una medida difusa $\nu : 2^{[m]} \rightarrow L$ se dice que es aditiva si $\nu(\mathcal{A} \cup \mathcal{B}) = \nu(\mathcal{A}) + \nu(\mathcal{B})$ para todo $\mathcal{A}, \mathcal{B} \subseteq [m]$ tales que $\mathcal{A} \cap \mathcal{B} = \emptyset$.

Ejemplo II.3. Un ejemplo de medida difusa considerada en este trabajo es la medida de potencia. Es definida para todo $\mathcal{A} \subseteq [m]$ por:

$$\nu_q(\mathcal{A}) = \left(\frac{|\mathcal{A}|}{m} \right)^q \quad (3)$$

donde $q > 0$.

Ejemplo II.4. La medida difusa $\nu_l : 2^{[m]} \rightarrow L$ más pequeña viene dada por

$$\nu_l(\mathcal{A}) = \begin{cases} 1 & \text{si } \mathcal{A} = [m] \\ 0 & \text{en otro caso} \end{cases} \quad (4)$$

La medida difusa $\nu_u : 2^{[m]} \rightarrow L$ más grande viene dada por

$$\nu_u(\mathcal{A}) = \begin{cases} 0 & \text{si } \mathcal{A} = \emptyset \\ 1 & \text{en otro caso} \end{cases} \quad (5)$$

Para cualquier medida difusa $\nu : 2^{[m]} \rightarrow L$ se cumple:

$$\nu_l(\mathcal{A}) \leq \nu(\mathcal{A}) \leq \nu_u(\mathcal{A})$$

para todo $\mathcal{A} \subseteq [m]$

Una vez introducida la medida difusa, presentamos la definición de la integral de Choquet discreta, la cual es un ejemplo de función de agregación presentada en la Definición II.1.

Definición II.5. [11] La integral de Choquet discreta en L con respecto a la medida difusa ν es definida como una aplicación $Ch_\nu : L^m \rightarrow L$

$$Ch_\nu(\mathbf{x}) = \sum_{i=1}^m (x_{\sigma(i)} - x_{\sigma(i-1)}) \nu(\mathcal{A}_{\sigma(i)}) \quad (6)$$

donde $\mathbf{x} = (x_1, \dots, x_m) \in L^m$, $\nu : 2^{[m]} \rightarrow L$ es una medida difusa en el conjunto $[m]$, $\sigma : [m] \rightarrow [m]$ es una permutación, con $x_{\sigma(1)} \leq \dots \leq x_{\sigma(m)}$ con la convención $x_{\sigma(0)} = 0$ y $\mathcal{A}_{\sigma(i)} := \{\sigma(i), \dots, \sigma(m)\}$ es el subconjunto de los índices correspondiente a los $m - i + 1$ mayores elementos de $\mathbf{x}$ para todo $i \in [m]$.

Por último, si bien definimos la integral de Choquet $Ch_\nu : L^m \rightarrow L$, en la literatura también podemos encontrar la definición [19] para un intervalo $\mathbb{I} = [a, b] \subset \mathbb{R}$. Dado que posteriormente en la aplicación (Sección IV) la utilizaremos, se puede definir la integral de Choquet discreta como una aplicación $Ch_\nu : \mathbb{I}^m \rightarrow \mathbb{I}$, mediante la misma expresión que en la Eq. 6. Donde $\mathbf{x} = (x_1, \dots, x_m) \in \mathbb{I}^m$ con la medida difusa $\nu : 2^{[m]} \rightarrow L$, manteniendo la convención $x_{\sigma(0)} = 0$ y $\mathcal{A}_{\sigma(i)} := \{\sigma(i), \dots, \sigma(m)\}$.

A continuación introducimos la definición de funciones de agregación n -dimensionales, como marco introductorio a la Sección III.

Definición II.6. [20] Sean n, m enteros positivos. Una aplicación $M : (L^n)^m \rightarrow L^n$ es una función de agregación m -aria n -dimensional si satisface las siguientes propiedades:

- (i) $M(\mathbf{0}, \dots, \mathbf{0}) = \mathbf{0}$ y $M(\mathbf{1}, \dots, \mathbf{1}) = \mathbf{1}$
- (ii) Para todo $\mathbf{x}_1, \dots, \mathbf{x}_m, \mathbf{y}_1, \dots, \mathbf{y}_m \in L^n$ tal que $\mathbf{x}_1 \leq \mathbf{y}_1, \dots, \mathbf{x}_m \leq \mathbf{y}_m$, entonces $M(\mathbf{x}_1, \dots, \mathbf{x}_m) \leq M(\mathbf{y}_1, \dots, \mathbf{y}_m)$.

B. Redes Neuronales Recurrentes: Memoria a corto y largo plazo (LSTM)

Las Redes Neuronales Recurrentes (RNN) nacen con la intención de modelar datos que tienen dependencia secuencial o de tiempo. Sin embargo, dado que los algoritmos de aprendizaje para redes neuronales suelen estar basados en el gradiente, en estas redes surge el problema llamado *vanishing gradient* [21], es decir, el decrecimiento recurrente del valor de una variable en la salida de la red neuronal. Esto es un problema especialmente grave cuando tratamos de entrenar redes con dependencias o secuencias temporales largas.

Las *Long Shot-Term Memory* (LSTM) surgen principalmente como respuesta a este problema y suponen un cambio radical [21] en el entrenamiento de las redes recurrentes ya que evitan el decrecimiento continuo de los parámetros. De esta manera, esta arquitectura de neurona artificial [17] genera un estado que permite la memorización de conocimiento que se utiliza en instantes temporales posteriores.

Las neuronas LSTM han tenido diversas modificaciones en la literatura, pero en este trabajo utilizamos una de las más extendidas [16]. En la Fig. 1 podemos observar el detalle del interior de una LSTM donde es importante recalcar las puertas [22] *forget gate* (f), *input gate* (i) y *output gate* (o) así como la celda candidata ($\tilde{c}$).

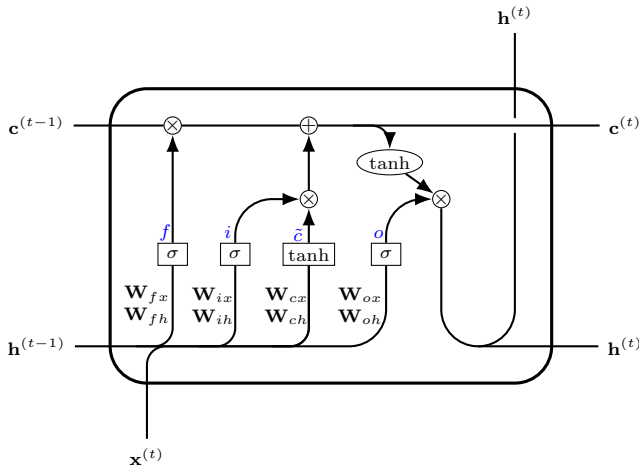


Fig. 1. Representación de una unidad LSTM.

A continuación, explicamos el funcionamiento de una unidad LSTM. Sea N la longitud de la secuencia de entrada, H el número de características que extraiga la celda y T el número de instantes de tiempo de la secuencia. Los que se detallan a continuación son las matrices y vectores asociados a cada una de las puertas y celda candidata:

- Matrices de pesos de entrada: $W_{fx}, W_{ix}, W_{cx}, W_{ox} \in \mathbb{R}^{H \times N}$
- Matrices de pesos recurrentes: $W_{fh}, W_{ih}, W_{ch}, W_{oh} \in \mathbb{R}^{H \times H}$
- Vectores de pesos del sesgo: $b_f, b_i, b_c, b_o \in \mathbb{R}^H$

El funcionamiento de forma descriptiva para cada instante de tiempo $t \in \{1, \dots, T\}$ es el siguiente:

- Los valores de entrada $x^{(t)}, h^{(t-1)}$ entran a las puertas f (Eq. 7), i (Eq. 8), $\tilde{c}$ (Eq. 9) y o (Eq. 11). En cada una de ellas, el valor $x^{(t)}$ se multiplica por cada una de las matrices de pesos de entrada, en función de la puerta. De forma análoga ocurre con los valores $h^{(t-1)}$ y las matrices de pesos recurrentes. Los vectores H -dimensionales obtenidos se suman junto con los vectores de sesgo correspondientes en cada una de ellas. Se utiliza una sigmoide logística ($\sigma(x) = \frac{1}{1+e^{-x}}$) como función de activación de las puertas y la tangente hiperbólica como función de activación de la celda candidata.
- El vector de la memoria a largo plazo del instante anterior ($c^{(t-1)}$) y el de la celda candidata ($\tilde{c}^{(t)}$) se combinan. Para ello se calcula el producto de Hadamard o elemento a elemento ($\circ$) entre el valor de la *forget gate* y la *input gate* respectivamente (Eq. 10). Ambos valores se suman obteniendo el valor de la celda en el instante actual ($c^{(t)}$).
- Por último, se calcula el vector de la memoria a corto plazo. Para ello, en primer lugar se pasa la información a corto plazo por una función de activación de salida. Para ello, utilizamos la tangente hiperbólica. Posteriormente se calcula el producto de Hadamard entre el valor de la *output gate* con la información obtenida de la última función de activación, obteniendo el valor de la memoria a corto plazo, $h^{(t)}$.

Las ecuaciones que describen el proceso son las siguientes (Eq. 7-12):

$$f^{(t)} = \sigma(W_{fx}x^{(t)} + W_{fh}h^{(t-1)} + b_f) \quad (7)$$

$$i^{(t)} = \sigma(W_{ix}x^{(t)} + W_{ih}h^{(t-1)} + b_i) \quad (8)$$

$$\tilde{c}^{(t)} = \tanh(W_{cx}x^{(t)} + W_{ch}h^{(t-1)} + b_c) \quad (9)$$

$$c^{(t)} = f^{(t)} \circ c^{(t-1)} + i^{(t)} \circ \tilde{c}^{(t)} \quad (10)$$

$$o^{(t)} = \sigma(W_{ox}x^{(t)} + W_{oh}h^{(t-1)} + b_o) \quad (11)$$

$$h^{(t)} = o^{(t)} \circ \tanh(c^{(t)}) \quad (12)$$

III. INTEGRAL DE CHOQUET MULTIDIMENSIONAL

A continuación se presenta una definición de la Integral de Choquet en retículos $(L^n, \leq)$, la cual es también un ejemplo de función de agregación n -dimensional, presentada en la Definición II.6.

Sean $x_1 = (x_{11}, \dots, x_{1n}), \dots, x_m = (x_{m1}, \dots, x_{mn})$ m vectores en L^n . Denotaremos por $x^1, \dots, x^n$ a los siguientes n vectores en L^m : $x^1 = (x_{11}, \dots, x_{m1}), \dots, x^n = (x_{1n}, \dots, x_{mn})$. Es decir, por ejemplo si $x_1, \dots, x_m$ son filas de una matriz de tipo $m \times n$, entonces $x^1, \dots, x^n$ son las columnas de dicha matriz.

Definición III.1. Sean n y m dos números enteros positivos. Sea $\nu = (\nu_1, \dots, \nu_n)$ una secuencia de medidas difusas en el conjunto $[m]$ y sean $Ch_{\nu_i} : L^m \rightarrow L$ las integrales de Choquet en L con respecto a la medida ν_i para todo $i \in [m]$. Una función $Ch_\nu^r : (L^n)^m \rightarrow L^n$ dada por:

$$Ch_\nu^r(\mathbf{x}_1, \dots, \mathbf{x}_m) = (Ch_{\nu_1}(\mathbf{x}^1), \dots, Ch_{\nu_n}(\mathbf{x}^n)) \quad (13)$$

para todo $\mathbf{x}_1, \dots, \mathbf{x}_m \in L^n$ es una Integral de Choquet discreta representable en L^n con respecto a ν y al orden $\leq$.

La integral de Choquet Ch_ν^r se denomina representable ya que es obtenida a través de la utilización de n integrales de Choquet en L de forma separada para cada componente:

$$Ch_\nu^r(\mathbf{x}_1, \dots, \mathbf{x}_m) = (Ch_{\nu_1}(x_{11}, \dots, x_{m1}), \dots, Ch_{\nu_n}(x_{1n}, \dots, x_{mn})) \quad (14)$$

Esta expresión es una generalización de la integral de Choquet estándar en L , ya que si todos los vectores de entrada son n -tuplas con las mismas coordenadas, p.e. $\mathbf{x} = (x, \dots, x)$ y la secuencia de medidas difusas es un vector con las mismas medidas, esto es, $\nu_1 = \dots = \nu_n = \nu$, la salida es una n -tupla con las mismas coordenadas iguales a Ch_ν .

Así mismo, de forma análoga a la que se explica en la Sección II, se puede extender la Definición III.1 a un intervalo $\mathbb{I} = [a, b] \subset \mathbb{R}$. De esta manera, podemos extender la función $Ch_\nu^r : (L^n)^m \rightarrow L^n$ a una aplicación $Ch_\nu^r : (\mathbb{I}^n)^m \rightarrow \mathbb{I}^n$ dada por:

$$Ch_\nu^r(\mathbf{x}_1, \dots, \mathbf{x}_m) = (Ch_{\nu_1}(\mathbf{x}^1), \dots, Ch_{\nu_n}(\mathbf{x}^n)) \quad (15)$$

para todo $\mathbf{x}_1, \dots, \mathbf{x}_m \in \mathbb{I}^n$ con respecto a ν y al orden $\leq$.

Proposición III.2. Bajo las condiciones de la definición anterior, sean $\nu_1 = \dots = \nu_n = \nu$. Entonces:

$$Ch_\nu^r(\mathbf{x}_1, \dots, \mathbf{x}_m) = (Ch_\nu(x_1, \dots, x_m), \dots, Ch_\nu(x_1, \dots, x_m)) \quad (16)$$

para todo $\mathbf{x}_i = (x_i, \dots, x_i) \in L^n$, $i \in [m]$.

IV. MODIFICACIÓN DE LA ARQUITECTURA DE UNA RED NEURONAL RECURRENTE BASADA EN LA DEFINICIÓN MULTIDIMENSIONAL DE LA INTEGRAL DE CHOQUET

En la presente sección se explica la introducción de las definiciones expuestas en las anteriores secciones en la arquitectura de una red neuronal recurrente, concretamente en una LSTM.

En este sentido, modificamos el operador de agregación de la red LSTM (suma de vectores) por la nueva definición de integral de Choquet discreta multidimensional.

En las nuevas ecuaciones generadas por la modificación de la función de agregación, esto se ve reflejado en el cálculo de la salida de la *forget gate* (Eq. 17), *input gate* (Eq. 18), *output gate* (Eq. 21) y celda candidata (Eq. 19), donde los vectores son fusionados mediante la integral de Choquet multidimensional.

El conjunto de las ecuaciones modificadas que describen el proceso son las siguientes:

$$\mathbf{f}^{(t)} = \sigma \left(Ch_\nu^r(\mathbf{W}_{fx}\mathbf{x}^{(t)}, \mathbf{W}_{fh}\mathbf{h}^{(t-1)}, \mathbf{b}_f) \right) \quad (17)$$

$$\mathbf{i}^{(t)} = \sigma \left(Ch_\nu^r(\mathbf{W}_{ix}\mathbf{x}^{(t)}, \mathbf{W}_{ih}\mathbf{h}^{(t-1)}, \mathbf{b}_i) \right) \quad (18)$$

$$\tilde{\mathbf{c}}^{(t)} = \tanh \left(Ch_\nu^r(\mathbf{W}_{cx}\mathbf{x}^{(t)}, \mathbf{W}_{ch}\mathbf{h}^{(t-1)}, \mathbf{b}_c) \right) \quad (19)$$

$$\mathbf{c}^{(t)} = \mathbf{f}^{(t)} \circ \mathbf{c}^{(t-1)} + \mathbf{i}^{(t)} \circ \tilde{\mathbf{c}}^{(t)} \quad (20)$$

$$\mathbf{o}^{(t)} = \sigma \left(Ch_\nu^r(\mathbf{W}_{ox}\mathbf{x}^{(t)}, \mathbf{W}_{oh}\mathbf{h}^{(t-1)}, \mathbf{b}_o) \right) \quad (21)$$

$$\mathbf{h}^{(t)} = \mathbf{o}^{(t)} \circ \tanh(\mathbf{c}^{(t)}) \quad (22)$$

Como hemos mostrado, la modificación principal del funcionamiento de la celda LSTM será la sustitución de la suma por la integral de Choquet multidimensional. No obstante, por completitud en el estudio, en la fusión de información vectorial en la unidad LSTM utilizaremos distintas funciones de agregación, así como las combinaciones lineales entre sí. Tomamos $\mathbf{z} = (\mathbf{W}_{gx}\mathbf{x}^{(t)}, \mathbf{W}_{gh}\mathbf{h}^{(t-1)}, \mathbf{b}_g) = (z_i)_{i=1}^3$ para $g \in \{f, i, c, o\}$, dado que son la misma expresión con matrices de pesos diferentes. A partir de ello, en la Tabla I mostramos las distintas funciones que utilizaremos.

TABLE I
DISTINTAS FUNCIONES DE AGREGACIÓN UTILIZADAS

M	Expresión función agregación
Max	$\max_{i=1}^3 z_i$
Ch ₂	$Ch_{\nu_2}^r(\mathbf{z})$
Ch _q	$Ch_{\nu_q}^r(\mathbf{z})$
Sum	$\sum_{i=1}^3 z_i$
Max + Sum	$\lambda_1 \max_{i=1}^3 z_i + \lambda_2 \sum_{i=1}^3 z_i$
Ch ₂ + Sum	$\lambda_1 Ch_{\nu_2}^r(\mathbf{z}) + \lambda_2 \sum_{i=1}^3 z_i$
Ch _q + Sum	$\lambda_1 Ch_{\nu_q}^r(\mathbf{z}) + \lambda_2 \sum_{i=1}^3 z_i$
Ch ₂ + Max	$\lambda_1 Ch_{\nu_2}^r(\mathbf{z}) + \lambda_2 \max_{i=1}^3 z_i$
Ch _q + Max	$\lambda_1 Ch_{\nu_q}^r(\mathbf{z}) + \lambda_2 \max_{i=1}^3 z_i$

Siendo $\nu_q(\mathcal{A})$ la expresión de la Eq. 3 donde q es un parámetro aprendido por la red neuronal. La medida $\nu_2(\mathcal{A})$ se corresponde a la misma expresión de la Eq. 3 evaluando $q = 2$. Los parámetros λ_1 y λ_2 son aprendidos por la propia red neuronal recurrente mediante el método de descenso por gradiente estocástico.

V. ESTUDIO EXPERIMENTAL

En la presente sección explicamos la arquitectura, conjunto de datos e hiperparámetros utilizados así como los resultados obtenidos y su posterior valoración.

A. Marco de trabajo experimental

1) *Conjunto de datos:* El conjunto de datos utilizado es el *Fashion-MNIST* [23], el cual consiste en un conjunto de entrenamiento de 60.000 imágenes de dimensiones 28×28 distribuidas en 10 clases, junto con un conjunto de test compuesto por 10.000 imágenes similares. Las imágenes corresponden a 10 artículos diferentes de categorías de ropa. En este caso, planteamos los datos que contiene una imagen como información secuencial [24].

2) *Arquitectura*: En este caso, como podemos observar en Fig. 2, establecemos una arquitectura en la que las imágenes son tomadas como datos secuenciales. En cada instante de tiempo $t \in \{1, \dots, T\}$ es tomada una fila de la imagen en forma de vector como dato de entrada $\mathbf{x}^{(t)} \in [0, 255]^N$. En el caso de este dataset concreto, $T = N = 28$.

La arquitectura consta de dos capas. La primera, una unidad de memoria LSTM con $H = 128$ nodos de capa oculta. En segundo lugar, una capa totalmente conectada (FC) que conecta los 128 nodos de la unidad LSTM con 10 nodos de FC, asignándoles un valor de probabilidad en $[0, 1]$ a cada uno de ellos. Se clasifica en el número de clase correspondiente al máximo valor de probabilidad del vector softmax. Esto es,

$$\hat{y} = \arg \max \frac{\exp(FC_j)}{\sum_{k=0}^9 \exp(FC_k)}.$$

En este experimento se han realizado 10 ejecuciones independientes de 40 epochs cada una. La tasa de aprendizaje fijada para el experimento es de $\alpha = 0.1$ y el método de optimización utilizado para el aprendizaje ha sido el descenso por gradiente estocástico (SGD).

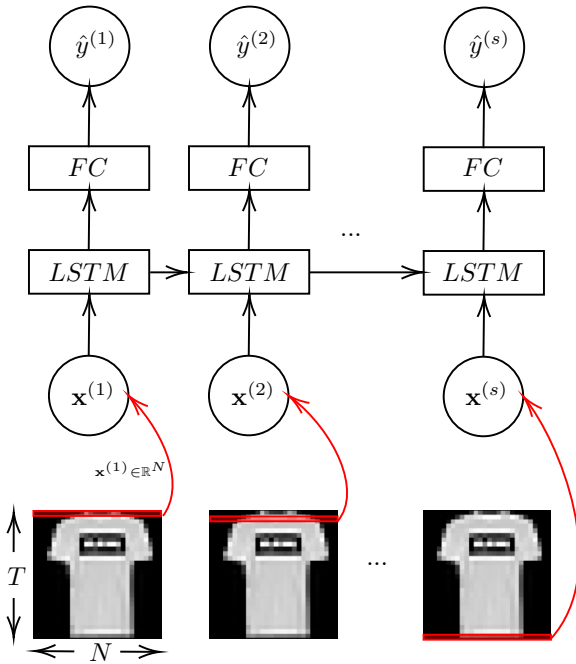


Fig. 2. Arquitectura de la red utilizada.

B. Resultados experimentales

A continuación mostramos los resultados obtenidos tras el cálculo de una media aritmética de diez ejecuciones independientes para cada una de las funciones de agregación así como para las combinaciones lineales de funciones. En las tablas II y III se recalca con negrita el resultado con mayor precisión obtenida.

En primer lugar, en la tabla II mostramos los resultados obtenidos con una sola función de agregación. El mejor resultado promedio lo obtenemos cuando realizamos la agregación

de los valores mediante la integral de Choquet, pero cuando el exponente $q > 0$ es aprendido por la propia red neuronal recurrente. Esto supone que el algoritmo modeliza mejor la interacción y posible coalición entre los datos. De esta manera, obtenemos una ponderación de los datos que permite una mejora de 1.01 puntos con respecto a la forma de agregación clásica en esta arquitectura, la suma.

TABLE II
PRECISIÓN DE DISTINTAS AGREGACIONES PARA EL DATASET
FASHION-MNIST (MEDIA DE 10 EJECUCIONES)

Agregación	Precisión
Max	86.45
Ch ₂	85.65
Ch _q	90.01
Sum	89.00

Dado que es posible mejorar los resultados mediante la combinación lineal de funciones de agregación, en la tabla III mostramos los resultados obtenidos a partir de la aplicación de estas últimas. A nivel general, en comparación con los resultados obtenidos con la aplicación de una única función de agregación, podemos observar que todos los resultados son mejores, ya que los valores $\lambda_1, \lambda_2 \in \mathbb{R}$ permiten modelar con mayor acierto los datos. Así mismo, con similitud a la tabla II, las agregaciones con las que mejor resultado obtenemos son aquellas las cuales uno de sus componentes es la integral de Choquet en la que se aprende la medida.

TABLE III
PRECISIÓN DE DISTINTAS COMBINACIONES DE AGREGACIONES PARA EL
DATASET FASHION-MNIST (MEDIA DE 10 EJECUCIONES)

Combinación	Precisión
Max + Sum	89.78
Ch ₂ + Sum	90.04
Ch _q + Sum	90.12
Ch ₂ + Max	89.71
Ch _q + Max	90.14

VI. CONCLUSIONES

En este trabajo hemos propuesto un nuevo método para la fusión de vectores multidimensionales, así como la utilización de dicha expresión para la fusión de información secuencial en redes neuronales recurrentes tipo LSTM. Así mismo, se ha corroborado una mejora en la precisión en el dataset *Fashion-MNIST*. Hemos observado que los mejores resultados se obtienen cuando sustituimos la suma por la integral de Choquet.

En cuanto a las líneas futuras, en el aspecto teórico nuestra intención es continuar investigando nuevas formas de fusión de vectores basadas en la integral de Choquet, como por ejemplo la generalización de expresiones o la utilización de órdenes admisibles. En la vertiente aplicada, las líneas futuras van en la dirección de la modificación de más arquitecturas y más complejas, así como la utilización de otros conjuntos de datos.

AGRADECIMIENTOS

Este trabajo ha sido financiado por la Agencia Estatal de Investigación (España) bajo el proyecto PID2019-108392GB-I00 (AEI/10.13039/ 501100011033).

REFERENCES

- [1] G. Beliakov, A. Pradera, and T. Calvo, "Aggregation functions: A guide for practitioners," in *Studies in Fuzziness and Soft Computing*, 2007.
- [2] L. De Miguel, M. Sesma-Sara, M. Elkano, M. Asiain, and H. Bustince, "An algorithm for group decision making using n-dimensional fuzzy sets, admissible orders and owa operators," *Information Fusion*, vol. 37, pp. 126–131, 2017.
- [3] M. Grabisch, J.-L. Marichal, R. Mesiar, and E. Pap, *Aggregation Functions*. Encyclopedia of Mathematics and its Applications, Cambridge University Press, 2009.
- [4] D. Paternain, J. Fernandez, 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. Theme: Aggregation operators.
- [5] G. Lucca, J. A. Sanz, G. P. Dimuro, B. Bedregal, M. J. Asiain, M. Elkano, and H. Bustince, "Cc-integrals: Choquet-like copula-based aggregation functions and its application in fuzzy rule-based classification systems," *Knowledge-Based Systems*, vol. 119, pp. 32–43, 2017.
- [6] C. A. Dias, J. C. S. Bueno, E. N. Borges, S. Botelho, G. Dimuro, G. Lucca, J. Fernández, H. Bustince, and P. Drews, "Using the choquet integral in the pooling layer in deep learning networks," in *NAFIPS*, 2018.
- [7] G. Beliakov, H. Bustince, and T. Calvo, "A practical guide to averaging functions," in *Studies in Fuzziness and Soft Computing*, 2016.
- [8] M. Grabisch, J.-L. Marichal, R. Mesiar, and E. Pap, "Aggregation functions: Means," *Inf. Sci.*, vol. 181, pp. 1–22, 2011.
- [9] M. Grabisch and C. Labreuche, "A decade of application of the choquet and sugeno integrals in multi-criteria decision aid," *Annals of Operations Research*, vol. 175, pp. 247–286, 2008.
- [10] G. Lucca, J. A. Sanz, G. P. Dimuro, E. N. Borges, H. Santos, and H. Bustince, "Analyzing the performance of different fuzzy measures with generalizations of the choquet integral in classification problems," in *2019 IEEE International Conference on Fuzzy Systems (FUZZ-IEEE)*, pp. 1–6, 2019.
- [11] G. Choquet, "Theory of capacities," *Annales de l'Institut Fourier*, vol. 5, pp. 131–295, 1954.
- [12] H. Bustince, R. Mesiar, J. Fernandez, M. Galar, D. Paternain, A. Altalhi, G. Dimuro, B. Bedregal, and Z. Takáč, "d-choquet integrals: Choquet integrals based on dissimilarities," *Fuzzy Sets and Systems*, vol. 414, pp. 1–27, 2021. Aggregation Functions.
- [13] G. P. Dimuro, G. Lucca, B. Bedregal, R. Mesiar, J. A. Sanz, C.-T. Lin, and H. Bustince, "Generalized cflf2-integrals: From choquet-like aggregation to ordered directionally monotone functions," *Fuzzy Sets and Systems*, vol. 378, pp. 44–67, 2020. Theme : Aggregation Operations.
- [14] G. Lucca, J. Antonio Sanz, G. P. Dimuro, B. Bedregal, H. Bustince, and R. Mesiar, "Cf-integrals: A new family of pre-aggregation functions with application to fuzzy rule-based classification systems," *Information Sciences*, vol. 435, pp. 94–110, 2018.
- [15] A. Graves, *Supervised Sequence Labelling with Recurrent Neural Networks*. Studies in computational intelligence, Berlin: Springer, 2012.
- [16] K. Greff, R. K. Srivastava, J. Koutník, B. R. Steunebrink, and J. Schmidhuber, "Lstm: A search space odyssey," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 28, no. 10, pp. 2222–2232, 2017.
- [17] S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural computation*, vol. 9, no. 8, pp. 1735–1780, 1997.
- [18] T. Murofushi, M. Sugeno, and M. Machida, "Non-monotonic fuzzy measures and the choquet integral," *Fuzzy Sets and Systems*, vol. 64, no. 1, pp. 73–86, 1994.
- [19] L. Jin, M. Kalina, R. Mesiar, and S. Borkotokey, "Discrete choquet integrals for riemann integrable inputs with some applications," *IEEE Transactions on Fuzzy Systems*, vol. 26, no. 5, pp. 3164–3169, 2018.
- [20] B. Bedregal, G. Beliakov, H. Bustince, T. Calvo, R. Mesiar, and D. Paternain, "A class of fuzzy multisets with a fixed number of memberships," *Information Sciences*, vol. 189, pp. 1–17, 2012.
- [21] S. Hochreiter, Y. Bengio, P. Frasconi, and J. Schmidhuber, "Gradient flow in recurrent nets: the difficulty of learning long-term dependencies," in *A Field Guide to Dynamical Recurrent Neural Networks* (S. C. Kremer and J. F. Kolen, eds.), IEEE Press, 2001.
- [22] F. Gers, J. Schmidhuber, and F. Cummins, "Learning to forget: continual prediction with lstm," in *1999 Ninth International Conference on Artificial Neural Networks ICANN 99. (Conf. Publ. No. 470)*, vol. 2, pp. 850–855 vol.2, 1999.
- [23] H. Xiao, K. Rasul, and R. Vollgraf, "Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms," *ArXiv*, vol. abs/1708.07747, 2017.
- [24] Q. V. Le, N. Jaitly, and G. E. Hinton, "A simple way to initialize recurrent networks of rectified linear units," *CoRR*, vol. abs/1504.00941, 2015.

Monotonicity in Non-deterministic Computable Aggregations

^{1st} Luis Magdalena

ETSI Informáticos
Universidad Politécnica de Madrid
Madrid, Spain
0000-0001-7639-8906

^{2nd} Daniel Gómez

Departamento de Estadística y Ciencia de Datos
Complutense University of Madrid
Madrid, Spain
0000-0001-9548-5781

^{3rd} Luis Garmendia

Facultad de Informática
Complutense University of Madrid
Madrid, Spain
0000-0003-3429-8434

^{4th} Javier Montero

Departamento de Estadística e Investigación Operativa
Complutense University of Madrid
Madrid, Spain
0000-0001-8333-2155

Abstract—Computable aggregation operators can be seen as a generalization of aggregation operators where the mathematical function is replaced by a program that performs the aggregation process. This extension allows the introduction of new aggregation processes not feasible under the classical framework. Particularly interesting are some non-deterministic processes widely considered to merge information. However, especially in non-deterministic processes, the extension of some of the well-known concepts for aggregation operators such as monotony, is needed. In this work, a new concept of monotonicity is proposed, from a probabilistic perspective, for non-deterministic computable aggregation operators. To be consistent, the concept coincides with the classical definition in the deterministic case. In addition, some cases of interest are analysed.

Index Terms—aggregation, computable aggregation, monotonicity.

I. INTRODUCTION

One of the most important processes for dealing with complex information is the aggregation process [1]–[6].

Aggregation is a key tool for most knowledge-based systems. In general, we can say that aggregation has the aim of using different pieces of information to come to a conclusion or a decision. Several research communities consider these tools, such as the multi-criteria community, the decision-sensor fusion community, the decision-making community, and the data mining community, among many others.

Aggregation functions have been associated in literature with aggregation processes. In this sense, the process to aggregate information has usually been modelled by a function or a family of functions.

Nevertheless, in the pioneering work of Montero [7] the concept of computable aggregations was defined. In that definition, an aggregation process is associated with a program not necessarily being expressed in terms of functions. This approach focuses on the way each aggregation is obtained, i.e., the procedure that produces the aggregated value. Relevant properties come from these procedures, and it is the specific procedure we apply what should be the main object of study.

Furthermore, the rupture between functions and aggregation operators allows to open the domain of aggregation processes to a field not yet analysed in this discipline: non-deterministic computable aggregations [8], which are those aggregations in which it cannot be guaranteed that the results of the information that is going to be aggregated coincide when replicating the process, so they are not functions. This type of aggregation procedure is very common in statistics, where, due to the volume of information that is processed, it is frequent to choose a representative sample on which to operate. Obviously, replicating the process does not imply that the sample coincides and therefore the result varies. This kind of computable aggregation needs to redefine some common properties of aggregation operators, as monotonicity, being no longer valid for non-deterministic computable aggregations.

Monotonicity has been always present in aggregation processes, from the initial ideas by Zadeh (corresponding to set operations as unions and intersections) to more recent proposals (directional monotonicity [6], arity-monotonic aggregation operators [9], ...). Monotonicity seems to be a desirable property associated with a certain robustness of the aggregation process (obviously not always needed). As an example, when aggregating information representing positive evidence on a fact, one expects aggregation being monotone. Even if the aggregation process is non-deterministic (a non-deterministic computable aggregation), one would expect some kind of monotonicity in this process as well.

The purpose of this paper is to revisit and redefine the monotony property for non-deterministic computable aggregations. Section II contains some preliminary knowledge useful to follow the rest of the paper. In Section III we address the key issue of how the output of a non-deterministic computable aggregation can be described. In Section IV we discuss the concept of empirical monotonicity, and population monotonicity, which generalizes monotonicity in the deterministic framework. Comparison between both concepts is developed in Subsection IV-C. Final Section includes an additional analysis

and stresses the interest of our results.

II. PRELIMINARIES

A. Non-deterministic Computable Aggregations

Aggregation operators were initially defined to deal with fuzzy sets [3], [5], [10]–[12], and as a consequence, they were associated with membership functions. So, this is the reason why they were defined as follows:

Definition 1. [2] An aggregation operator is a mapping $Ag : [0, 1]^n \rightarrow [0, 1]$ that satisfies:

- 1) $Ag(0, 0, \dots, 0) = 0$ and $Ag(1, 1, \dots, 1) = 1$.
- 2) Ag is monotonic.

It is even possible in some cases to define aggregation processes going beyond functions by considering methods that do not match with the concept of mapping. To analyse this option let us remind the concept of computable aggregation. The main contribution in [7] was to separate the strong association that existed between "aggregation processes" and explicit functions. But to do so let's first rewrite the previous definition in terms of lists as we will use a list notation all through this paper.

Definition 2. Let $\mathcal{L}$ be the set of non-empty and finite lists of degrees in $[0, 1]$. Then an aggregation operator is a mapping $Ag : \mathcal{L} \rightarrow [0, 1]$ that satisfies:

- 1) $Ag([0, 0, \dots, 0]) = 0$ and $Ag([1, 1, \dots, 1]) = 1$.
- 2) Ag is monotonic.

Definition 3. [7] (**Computable aggregation**). Let $L < T >$ be a non-empty and finite list of n elements with type T . A **computable aggregation** P is a program that transforms the list $L < T >$ into an element of T .

Remark 1. *This paper combines computational and mathematical aspects. From a computational point of view, the term list is used to refer to the type of data with which our program (Computable aggregation) works. However, from a mathematical approach, even considering that we will maintain the term list, these lists should be analysed as tuples.*

Given a program, there exist some situations in which it is not possible to build a function associated to it. Those are, for example, the non-deterministic programs that we introduce in the following definition.

Definition 4 (Deterministic program). A program is deterministic, or repeatable, if it produces the very same output when given the same input, no matter how many times it is run.

On the other hand, we have programs that can not be modelled by functions because of the intrinsic definition of function.

Definition 5 (Non-deterministic computable aggregation). A computable aggregation P over the set T is non deterministic if and only if the program implementing it is non deterministic.

Hereafter, we will assume $T = [0, 1]$, and consequently, our lists (L) will be defined in $\mathcal{L}$. The obtained results can be easily

extended to any lattice totally ordered and with maximum and minimum.

Let us denote by $\mathcal{P}_{\mathcal{D}}$ the class of deterministic computable aggregations and by $\mathcal{P}_{\mathcal{ND}}$ the set of non deterministic computable aggregations.

The non deterministic behaviour of a program can arise from different sources, but we will concentrate now in those aggregation processes involving random or probabilistic decisions. Some elements of $\mathcal{P}_{\mathcal{ND}}$ with this characteristic are the following [8]:

Definition 6. (Probability sampled computable aggregation). Given a value $p \in (0, 1]$, and given a family of aggregation operators $\{Ag_n : [0, 1]^n \rightarrow [0, 1], n \geq 1\}$, let us define the computable aggregation $P_{Ag,p}$ as the three steps program that for a given list $L_1 = [x_1, \dots, x_n] \in \mathcal{L}$ performs the following actions:

- **Step 1.** Reduce the list L_1 into another list L_2 of lower (or equal) dimension by randomly erasing the elements of the list with probability $1 - p$.
- **Step 2.** If L_2 is empty then return to Step 1.
- **Step 3.** Return the value $Ag_{|L_2|}(L_2)$ if $|L_2| \geq 1$.

Note that the computable aggregation $P_{Ag,p=1}$ is deterministic since $P_{Ag,1}(L_1) = Ag_{|L_1|}(L_1)$ (since in this case $L_2 = L_1$).

Another way to sample a list is by fixing a value k , and randomly selecting k elements from it. Given a list of m elements in $[0, 1]$, $L_1 = [x_1, \dots, x_m] \in \mathcal{L}$, let us denote by Sel_k a program that randomly chooses a sample (without re-sampling) of k elements of the list if $k < m$, and maintains the same list if $k \geq m$.

Definition 7. (k-sampled computable aggregation). Given a family of aggregation operators $\{Ag_n, n \geq 1\}$, the computable aggregation $P_{Ag,k}$, is defined as the program that for a given list L_1 of m elements, first applies the procedure $L_2 = Sel_k(L_1)$, and then computes the value $Ag_{|L_2|}(L_2)$.

It is important to notice that in the previous definitions the sampled list L_2 should be a sublist of L_1 , then maintaining the relative positions among elements.

B. Distribution functions

In our analysis based on populations we will use distribution functions and will consider some related concepts.

Definition 8. Given a random distribution X , the **cumulative distribution function** of X (denoted by F_X), is a function that for each value x_0 computes the probability induced by X of the set $\{X \leq x_0\}$. Formally the function $F_X : \mathbb{R} \rightarrow [0, 1]$, is defined as

$$F_X(x_0) = P(X \leq x_0).$$

We can simply use the term *distribution function* to refer to the *cumulative distribution function*.

Another concept that we need to introduce is the idea of empirical distribution. The empirical distribution function can be understood as an estimation of the cumulative distribution

function that is obtained from a sample of the global population.

Definition 9. Given a list $l = [a_1, \dots, a_m]$ with m values in $[0, 1]$, then the empirical distribution function associated to l and denoted by EF_l can be formally defined for all $x \in \mathbb{R}$ as

$$EF_l(x) = \frac{|\{i \mid a_i \leq x\}|}{m}.$$

Given a distribution X , if we obtain $l = [a_1, \dots, a_m]$ in an independent way from the global distribution X , there exist many statistical results that deal with the convergence of the empirical distribution function to the underlying cumulative distribution function. So, in general, when m increases, EF_l converges to F_X .

In addition, given two distributions we can compare them.

Definition 10. Given two distributions X_1 and X_2 with distribution functions F_{X_1} and F_{X_2} : $\mathbb{R} \rightarrow [0, 1]$. We will say that $X_1 \leq_D X_2$ when for all $x \in \mathbb{R}$, $F_{X_1}(x) \geq F_{X_2}(x)$.

III. CHARACTERISING THE OUTPUT OF A NON DETERMINISTIC COMPUTABLE AGGREGATION

Taking into account that the core properties of aggregation operators are monotonicity and boundary conditions, it is obvious that defining a new kind of aggregation, as the non deterministic computable aggregations, requires the analysis of this two properties.

The boundary conditions are quite simple concepts that have a straightforward counterpart in this kind of non deterministic computable aggregations. Any sampling of a list of 0s will be a new list of 0s (similarly for a list of 1s), so, as the underlying aggregation satisfies the boundary conditions, the computable aggregation also does. But monotonicity is a quite different question. The analysis of monotonicity requires the comparison of the outputs of two aggregations. The problem is that, differently from other aggregation processes, the output of a non-deterministic aggregation cannot be described in terms of a known and fixed value. This situation rises a previous question prior to consider monotonicity. How can we describe the output of a non-deterministic computable aggregation?

Given a computable aggregation P aggregating a list $L_1 \in \mathcal{L}$, we could describe its output by running the program several times and compiling the outputs. Clearly, if considering a deterministic computable aggregation, all outputs should be the same. However, a non-deterministic computable aggregation would probably produce several different values.

From now on we will denote by $LP_{L_1}^m \in \mathcal{L}$ the list of length m obtained after m executions of the program P over the list L_1 . These m -realizations could be characterized in terms of a distribution [13].

Definition 11. (Empirical distribution of a computable aggregation). Given a computable aggregation P and given a list $L_1 \in \mathcal{L}$, the distribution of results obtained after m executions of the program P over the list L_1 , will be referred as the *empirical distribution with size m of the program P over the list L_1 in $\mathcal{L}$* , represented by $DP_{L_1}^m$.

It is also possible that, by simply analysing the program or the underlying algorithm, we could determine what should theoretically be the distribution of outputs for a given input.

Given a computable aggregation P that aggregates a list $L_1 \in \mathcal{L}$, let us denote by $\mathcal{P}(L_1)$ the theoretical distribution after all possible realizations of the program P over the fixed list L_1 . Obviously, if $P \in \mathcal{P}_D$, the associate $\mathcal{P}(L_1)$ for any list will be a single value. For non deterministic programs, we will have here a probability distribution $\mathcal{P}(L_1)$ for each fixed value of L_1 . However there will be some cases where we would be not able to obtain that theoretical distribution.

In general, given a non deterministic computable aggregation P , it is not possible to know the theoretical distribution $\mathcal{P}(L_1)$. Nevertheless, we could try to approximate it by making many realizations of $P(L_1)$.

After defining $LP_{L_1}^m$, $DP_{L_1}^m$ and $\mathcal{P}(L_1)$ we need to consider its relations. In order to show the differences between the list of executions, the empirical distribution, and the theoretical distribution, let's consider the following example.

Example 1. Let $L_1 = [0.1, 0.2, 0.6, 0.7]$ be a list of four elements in $[0, 1]$, and let P be a deterministic program that calculates the average of the elements of the list. Let $m = 5$ be the number of executions. In one hand, it is easy to see that $LP_{L_1}^5 = [0.4, 0.4, 0.4, 0.4, 0.4]$ (i.e. a list of 5 elements each one of them taking the value 0.4), since the program is deterministic and the program will give always the same result. On the other hand, the empirical and theoretical distributions $DP_{L_1}^5$ and $\mathcal{P}(L_1)$ coincide with the degenerate distribution in the point 0.4 (i.e. $DP_{L_1}^5 = \mathcal{P}(L_1) = \{0.4\}$) that can be described by the following cumulative distribution function.

$$F(x) = \begin{cases} 0 & \text{If } x < 0.4 \\ 1 & \text{If } x \geq 0.4. \end{cases}$$

We have obtained the list of executions, the empirical distribution and the theoretical distribution.

Once analysed the previous example considering a deterministic computable aggregation, let's consider now a generic approach including the non deterministic case. Taking into account the concept of distribution function presented in previous section, if we want to estimate the distribution of $\mathcal{P}(L_1)$, we can execute m times $P(L_1)$, obtaining the list $LP_{L_1}^m$. From this list, it is possible to obtain what we have denoted by $DP_{L_1}^m$ that is the empirical distribution of the values of the list $LP_{L_1}^m$. The result will be:

$$F_{DP_{L_1}^m} = EF_{LP_{L_1}^m}. \quad (1)$$

Proposition 1. [14] Given a deterministic computable aggregation P , the following holds:

$$DP_{L_1}^m = \mathcal{P}(L_1).$$

Starting from a list $L_1 \in \mathcal{L}$ of elements to be aggregated, and a non deterministic aggregation process, we have defined two distributions describing the result of the aggregation process: the theoretical distribution $\mathcal{P}(L_1)$ and the empirical distribution ($DP_{L_1}^m$). In addition we have the list $LP_{L_1}^m$

obtained after m executions of the aggregation process. The interactions and relations among these elements will describe and characterize a non deterministic computable aggregation.

IV. MONOTONICITY IN NON-DETERMINISTIC COMPUTABLE AGGREGATION

In a first approach to the monotonicity condition, we will assume that the output for a certain input list $L_1 \in \mathcal{L}$ is described in terms of m -realizations $LP_{L_1}^m \in \mathcal{L}$, being also a list. In this situation any monotonicity analysis will require a method to compare/order elements of $\mathcal{L}$.

Let $L_1 = [a_1, \dots, a_n]$ and $L_2 = [b_1, \dots, b_n]$ be two lists of $\mathcal{L}$ with the same length $n > 0$. We will represent by $s(L_1)$ the list obtained by sorting L_1 in increasing order. Accordingly, $s(L_2)$ will correspond to the sorted version of L_2 .

Definition 12. (Sorted 1to1 preorder). A list L_1 in $\mathcal{L}$ is S1to1 lower or equal than a list L_2 in $\mathcal{L}$, if and only if the k^{th} value of $s(L_1)$ is lower or equal to the k^{th} value of $s(L_2)$, for all k from 1 to n , that is:

$$L_1 \leq_{\text{S1to1}} L_2 \text{ if and only if } s(L_1)_k \leq s(L_2)_k, \forall k = 1, \dots, n.$$

Definition 13. (1to1 partial order). A list $L_1 = [a_1, \dots, a_n]$ in $\mathcal{L}$ is 1to1 lower or equal than a list $L_2 = [b_1, \dots, b_n]$ in $\mathcal{L}$, if and only if the k^{th} element of L_1 is lower or equal than the k^{th} degree of L_2 . It is:

$$L_1 \leq_{\text{1to1}} L_2 \text{ if and only if } a_k \leq b_k, \forall k = 1, \dots, n.$$

This partial order is equivalent to the one obtained by considering the lists (of length n) in $\mathcal{L}$ as elements in $\mathbb{R}^n$ with its usual order.

A. Empirical monotonicity for non-deterministic computable aggregations

How to generalize the idea of monotonicity from classical aggregation operators to non-deterministic computable aggregations is not a trivial task. In general terms, monotonicity implies that if we have two input lists $L_1 \leq L_2$, their outputs should maintain that order. Consequently, the usual concept of monotonicity of aggregation operators is not valid for non-deterministic computable aggregations, as these aggregations do not produce the same output when receiving the same input. To cope with this situation, the concept of list associated to m executions for a given list L_i ($LP_{L_i}^m$) has been previously introduced. This concept describes the result obtained (a list $L_o \in \mathcal{L}$) after m executions of the program P over the list L_i . When m is large enough it allows an empirical analysis of the computable aggregation.

The following properties try to answer the question of how to define monotonicity from this point of view, applying the different methods previously considered for the comparison of lists. It is important to notice that we have lists as inputs (L_i), and we describe the outputs also in terms of lists ($LP_{L_i}^m$).

In what concerns the lists of inputs (L_i), from the point of view of an aggregation processes these lists behave as vectors (the n inputs of an aggregation process are an element of

$[0, 1]^n$), and should be ordered considering the usual order of $[0, 1]^n$, that is equivalent to $\leq_{\text{1to1}}$ previously defined. Consequently we will use this notation in the definitions.

Definition 14. (Strong $\leq$ -monotonicity of non-deterministic computable aggregations). Let P be a non-deterministic computable aggregation. Let $L_1 = [a_1, \dots, a_n]$ and $L_2 = [b_1, \dots, b_n]$ be two lists of degrees in $\mathcal{L}$. Let $\leq$ be a (partial) order on the set $\mathcal{L}$. A non-deterministic computable aggregation P is strong $\leq$ -monotone if and only if, when $L_1 \leq_{\text{1to1}} L_2$ then $(LP_{L_1}^m) \leq (LP_{L_2}^m)$ for any $m \geq 1$.

Definition 15. (Asymptotic $\leq$ -monotonicity of non-deterministic computable aggregations). Let P be a non-deterministic computable aggregation. Let $L_1 = [a_1, \dots, a_n]$ and $L_2 = [b_1, \dots, b_n]$ be two lists of degrees in $\mathcal{L}$. Let $\leq$ be a partial order on the set $\mathcal{L}$. A non-deterministic computable aggregation P is asymptotic $\leq$ -monotone if and only if, there exists a natural number n_0 such that If $L_1 \leq_{\text{1to1}} L_2$ then $(LP_{L_1}^m) \leq (LP_{L_2}^m)$ for any $m > n_0$.

B. Population monotonicity for non-deterministic computable aggregations

Once defined an order in the output-space of non deterministic computable aggregations, let us consider monotonicity.

Definition 16 (Population monotonicity of non-deterministic computable aggregations). Let P be a non-deterministic computable aggregation, and let L_1 and L_2 be two lists $\in \mathcal{L}$ ordered with the classical vectorial order definition (equivalent to the previously defined 1to1 order). A non-deterministic computable aggregation P is **population monotone** if and only if, when $L_1 \leq L_2$ then $\mathcal{P}(L_1) \leq_D \mathcal{P}(L_2)$.

Remark 2. Let us observe that previous definition generalizes the classical definition of monotonicity for deterministic computable aggregations. In the following proposition we will see in detail.

Proposition 2. Let P be a deterministic computable aggregation with associated function Ag . Then, the following holds:

P is population monotone if and only if Ag is monotone.

Proof. • From left to right. Given $L_1 \leq L_2$, will prove that $Ag(L_1) \leq Ag(L_2)$. First, since P is deterministic with associated aggregation function Ag , the following holds for the list L_1 and L_2 .

$$F_{\mathcal{P}(L_1)}(x) = \begin{cases} 0 & \text{If } x < Ag(L_1) \\ 1 & \text{If } x \geq Ag(L_1) \end{cases}$$

and

$$F_{\mathcal{P}(L_2)}(x) = \begin{cases} 0 & \text{If } x < Ag(L_2) \\ 1 & \text{If } x \geq Ag(L_2) \end{cases}$$

Then, since P is monotone, $\mathcal{P}(L_1) \leq_D \mathcal{P}(L_2)$ which implies that for all x $F_{\mathcal{P}(L_1)}(x) \geq F_{\mathcal{P}(L_2)}(x)$. Finally, it is very easy to see that these inequalities hold if and only if $Ag(L_1) \leq Ag(L_2)$ so the result is proved.

- From right to left. Given $L_1 \leq L_2$, we are going to prove that $\mathcal{P}(L_1) \leq_D \mathcal{P}(L_2)$. Since P is deterministic, $\mathcal{P}(L_1)$ and $\mathcal{P}(L_2)$ are degenerated distributions in the values $Ag(L_1)$ and $Ag(L_2)$, respectively. Since Ag is monotone, $Ag(L_1) \leq Ag(L_2)$, and thus it is very easy to see that $F_{\mathcal{P}(L_2)}(x) \leq F_{\mathcal{P}(L_1)}(x)$ for any value of x . $\square$

For practical reasons, we can approximate the theoretical distribution empirically if the number of executions is enough to ensure that empirical distribution is close to theoretical distribution. Taking into account this consideration, we present the following definition.

Definition 17 (Population asymptotic monotonicity of non-deterministic computable aggregations). Let P be a non-deterministic computable aggregation, and let L_1 and L_2 be two lists in $\mathcal{L}$ ordered with the classical vectorial order definition. A non-deterministic computable aggregation P is **population asymptotic monotone** if and only if when $L_1 \leq L_2$ then, $\forall x \in R$,

$$\lim_{m \rightarrow \infty} (F_{DP_{L_1}^m}(x) - F_{DP_{L_2}^m}(x)) \geq 0.$$

Proposition 3. Let P be a non deterministic computable aggregation, if P is **population monotone** then P is **population asymptotic monotone**.

Proof. First of all, let us note that $F_{DP_{L_1}^m}$ and $F_{DP_{L_2}^m}$ are bounded functions in $[0, 1]$. Also, by Glivenko-Cantelli theorem, there exist punctual convergence of the empirical distribution function ($F_{DP_L^m}$) into the theoretical one (denoted as $F_{P(L)}$) for $L \in \{L_1, L_2\}$. Consequently, the limits of $F_{DP_{L_1}^m}(x)$ and $F_{DP_{L_2}^m}(x)$ exist for all $x \in R$, and the limit of the difference coincides with the difference of limits. So, the following holds:

$$\lim_{m \rightarrow \infty} (F_{DP_{L_1}^m}(x) - F_{DP_{L_2}^m}(x)) = \lim_{m \rightarrow \infty} F_{DP_{L_1}^m}(x) - \lim_{m \rightarrow \infty} F_{DP_{L_2}^m}(x)$$

Now, due to the punctual convergence previously mentioned,

$$\lim_{m \rightarrow \infty} F_{DP_{L_1}^m}(x) - \lim_{m \rightarrow \infty} F_{DP_{L_2}^m}(x) = F_{P_{L_1}}(x) - F_{P_{L_2}}(x).$$

Finally, since P is population monotone, $P(L_1) \leq_D P(L_2)$ so the previous difference is greater than or equal to zero and the result is proved. $\square$

C. Empirical and population monotonicity

Finally, in this subsection we will see the relationship between the defined concepts of monotonicity for lists and distributions. In particular, let us analyse the relations between population monotonicity, asymptotic population monotonicity, asymptotic $\leq$ monotonicity and strong $\leq$ monotonicity.

Let us first consider the relation between the order $\leq_D$ for distribution functions and the order $\leq_{S1to1}$ for lists.

Proposition 4. Given a computable aggregation P , two lists $L_1, L_2 \in \mathcal{L}$, two m -realizations of the computable aggregation on these lists $LP_{L_1}^m, LP_{L_2}^m \in \mathcal{L}$, and the corresponding empirical distributions $DP_{L_1}^m$ and $DP_{L_2}^m$, then $DP_{L_1}^m \leq_D DP_{L_2}^m$ if and only if $LP_{L_1}^m \leq_{S1to1} LP_{L_2}^m$.

It is relevant to mention that the strong monotonicity defined for lists will not have an equivalence in the context of distributions since the requirements are too strong, so we will focus on the equivalence in the asymptotic cases for lists and distributions. Our first result is a natural consequence of Proposition 4, which establishes an equivalence between the orders in distributions after m realizations and the associated lists.

Proposition 5. P is population asymptotic monotone if P is $\leq_{S1to1}$ -asymptotic monotone.

Proof. This is direct by Proposition 4 in which we have that $DP_{L_1}^m \leq_D DP_{L_2}^m$ if and only if $LP_{L_1}^m \leq_{S1to1} LP_{L_2}^m$. So if P is $\leq_{S1to1}$ -asymptotic monotone, then $DP_{L_1}^m \leq_D DP_{L_2}^m$ for a given value of $m \geq m_0$, and as a consequence $\lim_{m \rightarrow \infty} (F_{DP_{L_1}^m}(x) - F_{DP_{L_2}^m}(x)) \geq 0$. $\square$

Now taking into account the previous implications between list of orders, distributions, and different classes of monotonicity, the following implications trivially holds.

Corollary 1. Let P be a computable aggregation.

- If P is strong $\leq_{S1to1}$ monotone, then P is population asymptotic monotone.
- If P is strong $\leq_{S1to1}$ monotone, then P is population monotone.
- If P is a deterministic computable aggregation, then P is strong $\leq$ monotone, for any order.

Prior to conclude the paper with a preliminary analysis on sampling aggregations, it is important to remark the reason to introduce the analysis based on populations, once we have the empirical approach. The main reason relies in the fact that once we move from lists to distributions, the dimension of the list does not affect. We can compare distributions related to lists of different dimension. That is, we can compare $DP_{L_1}^m$ and $DP_{L_2}^n$, no matter if $n \neq m$.

V. FINAL ANALYSIS AND CONCLUSIONS

Once introduced the different definitions and properties, it is possible to analyse monotonicity for a family of non deterministic computable aggregations that includes the most famous non deterministic process in statistics: sampling. We will consider now the k -sampled and probability sampled computable aggregations.

It is possible to check that both families of non deterministic computable aggregations are population monotone, for any classical aggregation operator function Ag . To do so, a possible guideline of the proof is to use the concept of empirical

distribution function associated to a list (EF_l , see definition 9) and the distribution associated to a list l denoted as D_l , and described (according to Eq. 1) as $F_{D_l} = EF_l$. Note that this is a generalization of the relation defined between a list of m -realizations ($LP_{L_1}^m$) and the corresponding empirical distribution ($DP_{L_1}^m$).

The idea beyond this, is to compare the associated empirical distribution with the population distributions associated with lists that are generated after we aggregate all possible scenarios of these two computable aggregations.

To conclude this section, let us note that in this work we provide an approach to monotonicity of non deterministic computable aggregations. This definition includes the classical notion in the deterministic case and offers two different approaches. The first approach relies on orders (and preorders) on lists. The second represents a probabilistic conception appearing in a natural way since we model the outputs of a computable aggregation (that in the deterministic case are single points), as probability distributions. As a consequence, we have had to answer the question of how to order lists and/or probability distributions. In this work, we provide a natural way of comparing distributions following the idea of stochastic dominance between distribution functions.

To conclude we make two additional comments. On the one hand, it is clear that other definitions for ordering lists and probability distributions could be considered and would give different versions of monotonicity. How to establish other orders between lists and between probability distributions is a question that we believe is worth studying in the future. On the other hand, we have also interpreted the output of a computable aggregation after a sufficient number of iterations, simply as a list of values. In that framework, we can also explore other order relations (now between lists) that would produce different notions of monotonicity.

ACKNOWLEDGMENT

This research has been partially supported by the Government of Spain (grant PGC2018-096509-B-I00), Comunidad de Madrid (grant S2013/ICCE-2845 and Convenio Plurianual con la Universidad Politécnica de Madrid en la línea de actuación Programa de Excelencia para el Profesorado Universitario), Complutense University (UCM Research Group 910149), and Universidad Politécnica de Madrid.

REFERENCES

- [1] B. Bouchon-Meunier, *Aggregation and fusion of imperfect information*. Physica, 2013, vol. 12.
- [2] G. Beliakov, A. Pradera, and T. Calvo, *Aggregations Functions: A guide for Practitioners*. Springer, 2007.
- [3] G. Beliakov, D. Gómez, S. James, J. Montero, and J. Rodríguez, "Approaches to learning strictly-stable weights for data missing values," *Fuzzy Sets and Systems*, vol. 325, pp. 97–113, 2017.
- [4] D. Gómez and J. Montero, "A discussion on aggregations operators," *Kybernetika*, vol. 40, pp. 107–120, 2004.
- [5] D. Gómez, K. Rojas, J. Montero, J. Rodríguez, and G. Beliakov, "Consistency and stability in aggregation operators, an application to missing data problems," *International Journal of Computational Intelligence Systems*, vol. 7, pp. 595–604, 2014.
- [6] M. Sesma-Sara, R. Mesiar, and H. Bustince, "Weak and directional monotonicity of functions on riesz spaces to fuse uncertain data," *Fuzzy Sets and Systems*, vol. 386, pp. 145–160, 2020, aggregation Operations.
- [7] J. Montero, R. G. del Campo, L. Garmendia, D. Gómez, and J. Rodríguez, "Computable aggregations," *Information Sciences*, vol. 460, pp. 439–449, 2018.
- [8] L. Garmendia, D. Gómez, L. Magdalena, and J. Montero, "Analyzing non-deterministic computable aggregations," in *Information Processing and Management of Uncertainty in Knowledge-Based Systems*, M.-J. Lesot, S. Vieira, M. Z. Reformat, J. P. Carvalho, A. Wilbik, B. Bouchon-Meunier, and R. R. Yager, Eds. Cham: Springer International Publishing, 2020, pp. 551–564.
- [9] M. Gagolewski and P. Grzegorzewski, "Arity-monotonic extended aggregation operators," in *Information Processing and Management of Uncertainty in Knowledge-Based Systems. Theory and Methods*, E. Hüllermeier, R. Kruse, and F. Hoffmann, Eds. Springer Berlin Heidelberg, 2010, pp. 693–702.
- [10] T. Calvo, A. Kolesárová, M. Komorníková, and R. Mesiar, "Aggregation operators: properties, classes and constructions methods," in *Aggregation Operators*, ser. Studies in Fuzziness and Soft Computing, C. T., M. G., and M. R., Eds. Springer, 2002, vol. 97, pp. 3–104.
- [11] K. Rojas, D. Gómez, J. Montero, and J. Rodríguez, "Strictly stable families of aggregation operators," *Fuzzy Sets and Systems*, vol. 228, pp. 44–63, 2013.
- [12] P. Olaso, K. Rojas, D. Gómez, and J. Montero, "A generalization of stability for families of aggregation operators," *Fuzzy Sets and Systems*, 2019.
- [13] L. Garmendia, D. Gómez, L. Magdalena, and J. Montero, "Empirical monotonicity of non-deterministic computable aggregations," in *The 19th World Congress of the International Fuzzy Systems Association (IFSA-2021)*, 2021.
- [14] L. Magdalena, D. Gómez, L. Garmendia, and J. Montero, "Population monotonicity of non-deterministic computable aggregations," in *2021 IEEE International Conference on Fuzzy Systems (FUZZ-IEEE)*, 2021.

Optimizando Desviaciones Moderadas Ponderadas para Interfaces Cerebro Ordenador

1st Javier Fumanal-Idocin

Estadística, Informática y Matemáticas
Universidad Pública de Navarra

Pamplona, Spain

javier.fumanal@unavarra.es

4th Asier Urio

Estadística, Informática y Matemáticas
Universidad Pública de Navarra

Pamplona, Spain

asier.uriio@unavarra.es

2nd Carmen Vidaurre

Estadística, Informática y Matemáticas
Universidad Pública de Navarra

Pamplona, Spain

carmen.vidaurre@unavarra.es

5th Graçaliz Pereira Dimuro

Informática e Matemática Aplicada
Universidade do Rio Grande do Norte

Natal, Brazil

gracalizdimuro@furg.br

3rd Marisol Gomez

Estadística, Informática y Matemáticas
Universidad Pública de Navarra

Pamplona, Spain

marisol@unavarra.es

6th Humberto Bustince

Estadística, Informática y Matemáticas
Universidad Pública de Navarra

Pamplona, Spain

bustince@unavarra.es

Abstract—Las interfaces cerebro-ordenador (BCI) basadas en el análisis de Electroencefalografía (EEG) están compuestas por varios elementos para procesar y clasificar las señales de entrada del cerebro. Una fase relevante de estos sistemas es el módulo de toma de decisiones, en el que la salida de diferentes clasificadores se fusiona en uno solo. En este trabajo proponemos el uso de funciones basadas en desviaciones moderadas con ponderaciones para la fase de toma de decisiones del sistema de BCI de fusión multimodal mejorado (EMF). Las funciones de agregación basadas en desviación moderada (MD) nos permiten elegir el mejor valor para agregar un vector de puntos utilizando una función de desviación moderada. Usando una MD ponderada, también podemos tener en cuenta la importancia relativa de cada dimensión en los datos multidimensionales que estamos agregando. Utilizando estas funciones en el EMF, podemos ponderar cada una de las diferentes señales cerebrales según su importancia, y utilizando la diferenciación automática, también podemos optimizarlas para el problema concreto a solucionar.

I. INTRODUCCIÓN

Las interfaces cerebro-ordenador (BCI) tienen como objetivo decodificar patrones de señales cerebrales para controlar diferentes mecanismos del cerebro [1], [2]. Se pueden utilizar diferentes propiedades de una señal para controlar dispositivos mediante señales cerebrales, [3], [4]. Una estrategia popular para modular las señales cerebrales de modo que se puedan interpretar en comandos es el entrenamiento con imágenes motoras (MI). Durante el proceso de MI, una persona imagina el movimiento de una parte del cuerpo, como por ejemplo la mano izquierda o derecha, los pies o la lengua [5]. Durante la imaginación de los movimientos, la potencia eléctrica de las áreas sensoriomotoras, contralaterales al lado del movimiento, se reduce, en un efecto conocido como desincronización relacionada con el evento (ERD), [6]–[8] y, a menudo, también se da un aumento de la potencia, o sincronización relacionada con el evento (ERS) en el lado ipsilateral. La identificación correcta de ERD/ERS influye en gran medida en el rendimiento de un sistema BCI basado en MI. Algunos algoritmos muy populares para identificar y clasificar los

cambios de potencia de las señales MI son el patrón espacial común (CSP), las máquinas vectoriales de soporte o las redes neuronales de aprendizaje profundo, [6], [9]–[14].

Los sistemas BCI se diseñan utilizando una amplia gama de técnicas para extraer características de los datos originales en bruto antes de extraer patrones cerebrales clasificables. Algunos procedimientos comunes incluyen la estimación de la potencia de banda estrecha en el dominio de tiempo [15], [16], en el dominio de frecuencia [17]–[19] o en el dominio de tiempo-frecuencia utilizando, por ejemplo, ondas de Meyer [20], [21]. Posteriormente, la clasificación se realiza generalmente empleando clasificadores lineales como el Análisis Discriminante Lineal (LDA), pero también QDA o SVMs [22] son procedimientos de clasificación populares. A veces, especialmente en el caso de paradigmas multiclase, están involucradas muchas características diferentes o se combinan diferentes características. En esos casos, el módulo de reconocimiento de patrones puede estar compuesto por un conjunto de clasificadores. La estrategia más común para combinar resultados de clasificación es probablemente la votación por mayoría [23].

En [24], los autores propusieron una nueva forma de realizar la toma de decisiones en un marco BCI en dos fases, la toma de decisiones multimodal. En [25] los autores propusieron una actualización de este mismo proceso, lo que resultó en el marco BCI de fusión multimodal mejorado (EMF). Este trabajo también presentó un estudio detallado de diferentes funciones de agregación, que mostró diferencias significativas en cuanto a la precisión de la clasificación final.

Sin embargo, el problema de elegir la mejor función de agregación sigue estando abierto. La literatura sobre teoría de la agregación sugiere el uso de agregaciones basadas en desviaciones moderadas (MD) como una opción para resolver el problema de escoger la agregación más adecuada para un conjunto de datos [26]. La idea de las MD es utilizar una desviación moderada para calcular la similitud entre el posible valor de salida y el conjunto de datos de entrada. Luego,

buscamos el valor que maximiza esa similitud con respecto a los datos de entrada.

En este trabajo, usamos la MD desarrollada en [27] para tomar en cuenta las diferentes importancias de cada característica de entrada usando pesos, y para resolver el problema de elegir la mejor agregación. Mostramos cómo se aplican estas funciones al marco BCI publicado en [25] y un procedimiento para aprender los pesos de acuerdo con los datos de entrada.

El resto del artículo es el siguiente: en la sección II explicamos las funciones basadas en la desviación moderada multivariante ponderada (II-A) y el marco de fusión multimodal mejorado (II-B) y como aplicar estas funciones en dicho marco. En la Sección III describimos nuestros experimentos y mostramos los resultados obtenidos, y en la Sección IV damos nuestras observaciones finales para este trabajo y explicamos nuestras pautas futuras.

II. MÉTODOS

En esta sección explicamos los conceptos asociados con las funciones basadas en la MD y el EMF para la clasificación de señales MI-BCI.

A. Desviaciones moderadas

Sea I un intervalo cerrado de números reales y sea $a = \min I$ y $b = \max I$. Una aplicación $M: I^n \rightarrow I$ se dice que es una función de agregación (n -dimensional) si M es creciente y $M(a, \dots, a) = a$, $M(b, \dots, b) = b$. Además, se dice que una función de agregación M es *promedia* [28] si solo si

$$\min(x_1, \dots, x_n) \leq M(x_1, \dots, x_n) \leq \max(x_1, \dots, x_n)$$

por cada $x_1, \dots, x_n \in I$.

Se puede encontrar una lista completa de varios ejemplos de funciones de agregación en [29]. Existen muchas familias de funciones de agregación, utilizadas según las circunstancias de cada problema. Entre ellas, el concepto de agregación basada en desviación resulta muy eficaz en muchas circunstancias. Su procedimiento es sencillo: se mide la diferencia entre los valores a agregar y un posible valor de salida y se escoge el valor que sea más similar a todos ellos.

Para medir la disimilitud en estos pares de valores, se puede utilizar una función de desviación o una función de desviación moderada. Sea I un intervalo cerrado de números reales. Se dice que un mapeo $D: I^2 \rightarrow \mathbb{R}$ es una *función de desviación* si y solo si $D(x, \cdot): I \rightarrow \mathbb{R}$ es continuo y estrictamente creciente, y $D(x, x) = 0$ para todo $x \in I$. Además, sea n un número natural, un mapeo $M_D: I^n \rightarrow \mathbb{R}$ se dice que es una *media de Daróczy* si y solo si M_D viene dado por

$$M_D(\mathbf{x}) = y,$$

donde y es la solución de la ecuación

$$\sum_{i=1}^n D(x_i, y) = 0.$$

Una media de Daróczy nos permite determinar qué tan diferente es cada entrada x_i de y , y el resultado final de la función M_D es ese y para el cual la suma de todas las diferencias es tan pequeña como sea posible. Sin embargo, debido a algunos problemas de esta definición, se propusieron las desviaciones moderadas:

Sea I un intervalo cerrado de números reales. Se dice que un mapeo $D: I^2 \rightarrow \mathbb{R}$ es una *función de desviación moderada* si y solo si

- 1) Por cada $x \in I$, $D(x, \cdot): I \rightarrow \mathbb{R}$ no es decreciente;
- 2) Por cada $y \in I$, $D(\cdot, y): I \rightarrow \mathbb{R}$ no es creciente;
- 3) $D(x, y) = 0$ si y solo si $x = y$, $x, y \in I$.

Una función dada por $D(x, y) = y - x$ para $x, y \in \mathbb{R}$, representa el ejemplo prototípico de función de desviación moderada.

Ahora procedemos a presentar un método para construir desviaciones moderadas ponderadas para datos multi-variable, que luego usaremos para agregar preferencias de distintos expertos.

Proposition II.1. Sea $f: \mathbb{R} \rightarrow \mathbb{R}$ una función no decreciente tal que $f(x) = 0$ si y solo si $x = 0$. Sea $s: [0, 1] \rightarrow \mathbb{R}$ una función estrictamente creciente. Luego, un mapeo $D_{f,s}: [0, 1]^2 \rightarrow \mathbb{R}$ dado por

$$D_{f,s}(x, y) = f(s(y) - s(x))$$

es una función de desviación moderada.

Sea el conjunto de datos multidimensionales una matriz $p \times q$ $\mathbb{A}$ de n vectores de manera que tanto p como q sean divisibles por un entero positivo r , $r \neq 1$. La matriz $\mathbb{A}$ puede considerarse compuesta de n submatrices $\mathbb{A}_1, \mathbb{A}_2, \dots, \mathbb{A}_n$ con entradas reales, donde $\mathbb{A}_j$ es una “sección transversal” de $\mathbb{A}$ a través de todos sus j -ésimos componentes. Por tanto, la matriz $\mathbb{A}$ puede considerarse como el vector $(\mathbb{A}_1, \mathbb{A}_2, \dots, \mathbb{A}_n)$ de matrices.

Descompongamos la matriz $\mathbb{A}$ en $(p/r) \cdot (q/r)$ mutuamente disjuntas submatrices $r \times r$ $\mathbb{B}^{\alpha\beta}$, $\alpha \in \{1, 2, \dots, p/r\}$, $\beta \in \{1, 2, \dots, q/r\}$. Al aplicar el método de fusión en cuestión, reemplazaremos cada submatriz de n -tuplas por un representante apropiado que es nuevamente una n -tupla, lo que da como resultado un $(p/r) \times (q/r)$ matriz $\mathbb{C}$ de n reales -tuplas.

Paso #1 (Descomposición de la matriz A): sea $\mathcal{I} = [p/r]$ y $\mathcal{J} = [q/r]$. Denotando el elemento en la posición (r, s) de $\mathbb{A}$ por a_{rs} . Se define para cada $\alpha \in \mathcal{I}$ y $\beta \in \mathcal{J}$ una $r \times r$ matriz $\mathbb{B}^{\alpha\beta}$ de n -tuplas como

$$b_{i,j}^{\alpha\beta} = a_{(\alpha-1)r+i, (\beta-1)r+j} \quad \text{para } i, j \in [r]. \quad (1)$$

Entonces, la matriz $\mathbb{A}$ puede verse como una unión disjunta de matrices $\mathbb{B}^{\alpha\beta}$ sobre todos los índices $\alpha \in \mathcal{I}$, $\beta \in \mathcal{J}$.

Sea α un índice arbitrario pero fijo de $\mathcal{I}$ y β de $\mathcal{J}$, respectivamente. Entonces la matriz $\mathbb{B}^{\alpha\beta}$ consta de n de $r \times r$

matrices cuadradas

$$\mathbb{B}_k^{\alpha\beta} = \begin{pmatrix} b_{11k}^{\alpha\beta} & b_{12k}^{\alpha\beta} & \cdots & b_{1rk}^{\alpha\beta} \\ b_{21k}^{\alpha\beta} & b_{22k}^{\alpha\beta} & \cdots & b_{2rk}^{\alpha\beta} \\ \vdots & \vdots & \ddots & \vdots \\ b_{r1k}^{\alpha\beta} & b_{r2k}^{\alpha\beta} & \cdots & b_{rrk}^{\alpha\beta} \end{pmatrix} \text{ para cada } k \in [n]. \quad (2)$$

Nota: Dentro del espacio euclidiano n -dimensional $\mathbf{E}^n$ alcance, los elementos de $\mathbb{B}_k^{\alpha\beta}$ corresponden a las proyecciones sobre el k -eje.

Paso #2 (Fijando un dominio): teniendo en cuenta las posibles aplicaciones, es necesario fijar un intervalo real cerrado para ser utilizado como dominio de una función de agregación basada en desviaciones. Por esta razón, designamos

$$\wedge_k^{\alpha\beta} = \min_{i,j \in [r]} \{b_{ijk}^{\alpha\beta}\} \quad \text{y} \quad \vee_k^{\alpha\beta} = \max_{i,j \in [r]} \{b_{ijk}^{\alpha\beta}\}, \quad k \in n.$$

Llamando $I_k^{\alpha\beta}$ al intervalo real cerrado $[\wedge_k^{\alpha\beta}, \vee_k^{\alpha\beta}]$, i.e.

$$I_k^{\alpha\beta} = \{x \in \mathbb{R} \mid \wedge_k^{\alpha\beta} \leq x \leq \vee_k^{\alpha\beta}\}. \quad (3)$$

Servirá como el dominio de una función de agregación basada en la desviación que se aplicará como la agregación de $\mathbb{B}_k^{\alpha\beta}$ elementos.

Paso #3 (Implementación de la agregación basada en desviaciones):

Theorem II.2. Sea I un intervalo de números reales. Sea $\mathbb{X}$ una matriz $s \times t$ tal que $x_{ij} \in I$ por cada $i \in [s]$, $j \in [t]$. Denotamos por $I_{\mathbb{X}}$ un intervalo real cerrado

$$\left[\min_{i \in [s], j \in [t]} \{x_{ij}\}, \max_{i \in [s], j \in [t]} \{x_{ij}\} \right].$$

Sea $D: I_{\mathbb{X}} \times I_{\mathbb{X}} \rightarrow \mathbb{R}$ una función de desviación moderada. Se define una función $M_D: I_{\mathbb{X}}^{s \times t} \rightarrow \mathbb{R}$ por

$$M_D(\mathbb{X}) = \frac{1}{2} \left(\sup \{y \in I \mid \sum_{i=1}^s \sum_{j=1}^t D(x_{ij}, y) < 0\} + \inf \{y \in I \mid \sum_{i=1}^s \sum_{j=1}^t D(x_{ij}, y) > 0\} \right). \quad (\text{MDX})$$

entonces M_D es una función de agregación idempotente y simétrica.

Paso #4 (Ponderando las matrices): Sea I un intervalo de números reales. Sea $D: I^2 \rightarrow \mathbb{R}$ una función de desviación moderada. Sea $\mathbb{W}$ una matriz de ponderación no negativa $s \times t$ tal que $w_{ij} \in [0, \infty)$ por cada $i \in [s]$, $j \in [t]$, y $\mathbb{X}$ será una matriz $s \times t$ de números reales tal que $x_{ij} \in I$ por cada $i \in [s]$, $j \in [t]$. Se dice que el mapeo $M_{D, \mathbb{W}}: I^{p \cdot q} \rightarrow I$ es una función

de agregación matricial basada en la desviación ponderada si y solo si

$$M_{D, \mathbb{W}}(\mathbb{X}) = \frac{1}{2} \left(\sup \{y \in I \mid \sum_{i=1}^s \sum_{j=1}^t w_{ij} D(x_{ij}, y) < 0\} + \inf \{y \in I \mid \sum_{i=1}^s \sum_{j=1}^t w_{ij} D(x_{ij}, y) > 0\} \right). \quad (\text{MDWX})$$

La imagen $M_{D, \mathbb{W}}(\mathbb{X})$ de $\mathbb{X}$ se denomina agregación de matrices basada en desviaciones moderadas ponderada.

Proposition II.3. Siendo $k \in [n]$ un número entero escogido arbitrariamente. Sea $\mathbb{B}_k^{\alpha\beta}$ una matriz $r \times r$ definida como Eq. (2). Siendo $\mathbf{w} = (w_1, w_2, \dots, w_n) \in [0, \infty)^n$ un vector de pesos no negativos. Siendo $M_D: (I_k^{\alpha\beta})^{r \cdot r} \rightarrow \mathbb{R}$ una agregación de matrices basada en desviaciones moderadas. Entonces, para cada $k \in [n]$, la función $M_{D, \mathbf{w}}: (I_k^{\alpha\beta})^{r \cdot r} \rightarrow \mathbb{R}$ definida como:

$$M_{D, \mathbf{w}_k}(\mathbb{B}_k^{\alpha\beta}) = M_D(w_k \mathbb{B}_k^{\alpha\beta}) = M_D \begin{pmatrix} w_k \cdot b_{11k}^{\alpha\beta} & \cdots & w_k \cdot b_{1rk}^{\alpha\beta} \\ w_k \cdot b_{21k}^{\alpha\beta} & \cdots & w_k \cdot b_{2rk}^{\alpha\beta} \\ \vdots & \ddots & \vdots \\ w_k \cdot b_{r1k}^{\alpha\beta} & \cdots & w_k \cdot b_{rrk}^{\alpha\beta} \end{pmatrix} \quad (\text{MDw})$$

es una agregación de matrices basada en desviaciones moderadas ponderada.

Example II.4. Siendo I un intervalo de números reales cerrado. Siendo $\mathbb{A}$ una matriz de dimensión $12\,000 \times 800$ de cuádruplas reales a_{ij} de modo que

$$a_{ij} = (a_{ij1}, a_{ij2}, a_{ij3}, a_{ij4}), \quad i \in [10\,000], \quad j \in [800],$$

y

$$(a_{ij1}, a_{ij2}, a_{ij3}, a_{ij4}) \in I^5.$$

Siendo $r = 2$. Entonces, siguiendo (1), $\alpha \in [10\,000/2] = [5\,000]$ y $\beta \in [800/2] = [400]$.

Para $\varepsilon \geq 1$, se define $D_\varepsilon: I^2 \rightarrow \mathbb{R}$ como

$$D_\varepsilon(x, y) = (x + \varepsilon)(y - x).$$

Pues, D_ε es una función de desviación, y la correspondiente función de agregación de dos variables $M_{D_\varepsilon}: I^2 \rightarrow \mathbb{R}$ es dada por:

$$M_{D_\varepsilon}(u, v) = \frac{u(u + \varepsilon) + v(v + \varepsilon)}{u + v + 2\varepsilon} \quad (4)$$

para todo $u, v \in I$ siendo $u + v + 2\varepsilon \neq 0$.

Siendo $\mathbf{w} = (w_1, w_2, w_3, w_4) \in [0, \infty)$ un vector de pesos. Para todo $k \in [4]$, $\alpha \in [5\,000]$, $\beta \in [400]$, se define, de

acuerdo con Eq. ((2)), una matriz $2 \times 2 \mathbb{B}_k^{\alpha\beta}$. Entonces la función $M_{D_\varepsilon, \mathbf{w}}: I^4 \rightarrow \mathbb{R}$:

$$\begin{aligned} M_{D_\varepsilon, \mathbf{w}_k}(\mathbb{B}_k^{\alpha\beta}) &= M_{D_\varepsilon}(w_k \mathbb{B}_k^{\alpha\beta}) \\ &= M_{D_\varepsilon} \begin{pmatrix} w_k b_{11k}^{\alpha\beta} & w_k b_{12k}^{\alpha\beta} \\ w_k b_{21k}^{\alpha\beta} & w_k b_{22k}^{\alpha\beta} \end{pmatrix} \\ &= \frac{\sum_{i=1}^2 \sum_{j=1}^2 w_k b_{ijk}^{\alpha\beta} (b_{ijk}^{\alpha\beta} + \varepsilon)}{4w_k \varepsilon + \sum_{i=1}^2 \sum_{j=1}^2 w_k b_{ijk}^{\alpha\beta}} \end{aligned} \quad (5)$$

es la agregación de matrices basada en desviaciones moderadas ponderada. Finalmente, la resultante matriz de $5000 \times 400 \mathbb{C}$ de $5000 \cdot 400 = 2000000$ cuadrúplas

$$\mathbf{y}_w^{\alpha\beta} = (M_{D_\varepsilon, w_1}(\mathbb{B}_1^{\alpha\beta}), \dots, M_{D_\varepsilon, w_4}(\mathbb{B}_4^{\alpha\beta})),$$

$\alpha \in [5000]$, $\beta \in [400]$, es el resultado $\mathbb{A}$ con respecto a D_ε .

B. Enhanced-Multimodal Fusion BCI Framework

El EMF es el marco MI BCI que clasifica señales EEG [25]. El EMF consta de 5 fases diferentes:

- 1) Calcular la transformada rápida de Fourier de la señal de EEG para transformar los datos de potencia de EEG en el dominio de la frecuencia. Luego, se realiza una diferenciación de la salida FFT en cada frecuencia.
- 2) Se dividen los datos en cinco bandas de ondas diferentes: δ 1 – 4 Hz, θ 4 – 8 Hz, α 8 – 14 Hz, β 14 – 30 Hz y γ 1 – 30 Hz (Todas).
- 3) Se calcula el CSP en cada banda de onda para extraer características con separación espacial máxima [30].
- 4) Se entrena un conjunto de clasificadores para cada banda de onda: un análisis discriminante lineal (LDA), análisis discriminante cuadrático (QDA), máquina de vectores de soporte (SVM), K-vecinos cercanos (KNN) y proceso gaussiano (GP). De esta manera tenemos un clasificador de cada tipo para cada banda de onda estudiada.
- 5) Se realiza la decisión multimodal utilizando dos funciones de agregación. Los mejores resultados en [25] se obtuvieron utilizando una integral de Choquet / Sugeno en la primera fase de la agregación y una función de overlap generalizada [31].

Se puede encontrar un esquema visual del EMF en la Figura 1.

C. Aprendizaje de ponderaciones en una agregación basada en desviación moderada ponderada

Como se detalla en la Sección II-A, es posible usar una MD para agregar un vector de entradas, dando más importancia a algunos canales. Sin embargo, este proceso de ponderación no es sencillo, ya que qué canales deberían ser más importantes que otros depende en gran medida de la tarea a resolver. Además, en el caso de la decisión multimodal, los pesos aprendidos para la fase 1 pueden influir mucho en los pesos óptimos para la fase 2.

En el caso de una tarea de aprendizaje supervisado, podemos plantear este problema como uno de optimización. Usando los logits de los clasificadores como datos de entrenamiento, los agregamos usando la MD ponderada y luego clasificamos cada muestra. La función de coste para la función de optimización es la entropía cruzada para las clases C y las muestras de M :

$$- \sum_{m=1}^M \sum_{c=1}^C y_{c,m} * \log(p_{c,m}) \quad (6)$$

donde $y_{c,m}$ es un valor binario que es 1 solo si c es la etiqueta real para la muestra m , y $p_{c,m}$ es la probabilidad predicha para la muestra m a ser de clase c .

Procedemos entonces a solucionar este problema mediante la autodiferenciación. El algoritmo de backpropagation para resolver este problema está presente en todas las bibliotecas de cálculo de tensores, que se utilizan comúnmente para Deep Learning [32].

III. EXPERIMENTOS Y RESULTADOS

Para nuestros experimentos, hemos utilizado el conjunto de datos BCI Competition IV 2a [33]. Este conjunto de datos consta de cuatro clases de tareas: lengua, pie, mano izquierda y mano derecha realizadas por 9 voluntarios. Para cada tarea, se recolectaron 22 canales de EEG. Hay un total de 288 ensayos para cada participante, distribuidos equitativamente entre las 4 clases. Para nuestra configuración experimental, hemos utilizado 4 de los 22 canales disponibles (8, 12, 14, 18), siguiendo los procedimientos de [24], [25].

De cada sujeto, hemos generado veinte particiones de las 288 pruebas que constan de 50% entrenamiento (144 pruebas) y 50% test (144 pruebas) elegidas al azar. Dado que tenemos 9 sujetos, esto produce un total de 90 conjuntos de datos diferentes.

Estudiamos tanto la clasificación binaria de mano izquierda/derecha como la tarea completa de cuatro clases.

A. Resultados de la clasificación binaria

En esta sección hemos estudiado el rendimiento del EMF utilizando una MD ponderada para la clasificación binaria. Hemos estudiado dos versiones de esta MD ponderada, estableciendo todos los pesos en 1 y aprendiéndolos usando la propagación hacia atrás.

En la Tabla I mostramos los resultados de la MD ponderada en ambos casos y los comparamos con otras agregaciones. Obtuvimos un mejor resultado aprendiendo los pesos que fijándolos al mismo valor, lo que mostró que el algoritmo de entrenamiento funcionó bien. Funcionó de manera similar a la media aritmética y la integral de Choquet, y mejor que la integral de Sugeno. Creemos que esta ventaja de rendimiento se debe a la flexibilidad que se obtiene al utilizar diferentes funciones de agregación en ambas fases.

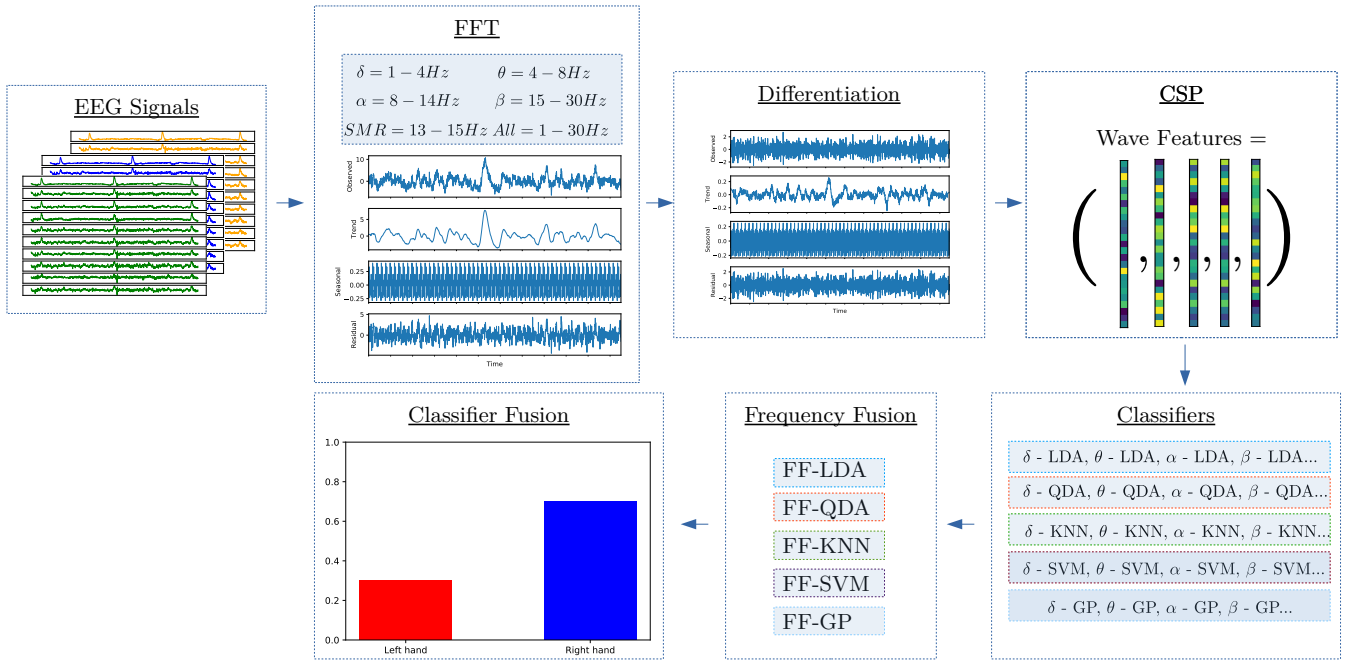


Fig. 1. Esquema visual para el marco BCI de fusión multimodal mejorado.

TABLE I
RESULTADOS DE LA CLASIFICACIÓN BINARIA UTILIZANDO EL EMF CON
DIFERENTES CONFIGURACIONES.

Agregación	Tasa de acierto
MD-Ponderada aprendida	86.97% $\pm$ 3.98
MD-Ponderada fijada	85.80% $\pm$ 4.08
Media Aritmética	85.80% $\pm$ 4.04
Integral de Choquet	86.39% $\pm$ 4.20
Integral de Sugeno	81.39% $\pm$ 4.39

TABLE II
RESULTADOS DE LA CLASIFICACIÓN DE LAS CUATRO CLASES UTILIZANDO
EL EMF CON DIFERENTES CONFIGURACIONES.

Agregación	Tasa de acierto
MD-Ponderada aprendida	65.91% $\pm$ 13.15
MD-Ponderada fijada	72.93% $\pm$ 2.29
Media Aritmética	72.22% $\pm$ 2.31
Integral de Choquet	72.93% $\pm$ 1.85
Integral de Sugeno	64.45% $\pm$ 2.66

B. Resultados de la tarea completa

En esta sección, hemos estudiado el rendimiento del sistema BCI EMF utilizando la MD ponderada para la tarea de clasificación de cuatro clases. Hemos estudiado la versión de pesos fijos de la MD, donde todos los canales son igualmente importantes, y la versión de peso aprendido usando la propagación hacia atrás, al igual que en la tarea de clasificación binaria.

Hemos mostrado nuestros resultados en la Tabla II. En esta Tabla también encontramos los resultados para la media aritmética y las integrales de Sugeno y Choquet. Descubrimos que la MD ponderada funciona igual o mejor que las otras integrales. Aprender el peso del canal también nos dio un peor resultado que fijar todos los pesos a 1, probablemente por la falta de suficientes datos. El mejor resultado se obtuvo usando la MD ponderado fija y la integral de Choquet, que resultaron en una precisión de 72,93%.

IV. CONCLUSIONES

En este trabajo hemos presentado las agregaciones ponderadas basadas en desviaciones moderadas aplicadas a un marco de interfaz cerebro-computadora de imágenes motoras, el marco de fusión multimodal mejorado. Hemos mostrado en qué parte del marco BCI se aplica y cómo se pueden optimizar los pesos de esta función mediante la diferenciación automática.

Hemos comparado los resultados utilizando pesos optimizados y fijos, obteniendo mejores resultados al aprenderlos. Asimismo, hemos comparado las agregaciones basadas en Desviación Moderada con el promedio aritmético y las integrales de Sugeno y Choquet, obteniendo resultados favorables a nuestras soluciones.

Nuestra investigación futura tendrá como objetivo mejorar el proceso de aprendizaje para requerir menos muestras de entrenamiento y mejorar aún más la fase de salida del clasificador explorando más funciones de desviación moderada.

AGRADECIMIENTOS

La investigación de Javier Fumanal Idocin, A. Urio y Humberto Bustince ha sido apoyada por el proyecto PID2019-108392GB-I00 (AEI / 10.13039 / 501100011033).

La investigación de Carmen Vidaurre ha sido financiada por el proyecto RyC-2014-15671.

REFERENCES

- [1] C. Vidaurre, C. Klauer, T. Schauer, A. Ramos-Murguialday, and K.-R. Müller, "Eeg-based bci for the linear control of an upper-limb neuroprosthesis," *Medical engineering & physics*, vol. 38, pp. 1195–1204, 2016.
- [2] B. Blankertz, L. Acqualagna, S. Dähne, S. Haufe, M. Schultze-Kraft, I. Sturm, M. Uscumlic, M. Wenzel, G. Curio, and K. Müller, "The berlin brain-computer interface: Progress beyond communication and control," *Front Neurosci*, vol. 10, p. 530, 2016.
- [3] B. Blankertz, S. Lemm, M. Treder, S. Haufe, and K.-R. Müller, "Single-trial analysis and classification of ERP components—a tutorial," *NeuroImage*, vol. 56, no. 2, pp. 814–825, 2011.
- [4] C. Sannelli, C. Vidaurre, K. Müller, and B. Blankertz, "A large scale screening study with a smr-based bci: Categorization of bci users and differences in their smr activity," *PLOS ONE*, vol. 14, p. e0207351, 2019.
- [5] R. Scherer and C. Vidaurre, "Chapter 8 - motor imagery based brain-computer interfaces," in *Smart Wheelchairs and Brain-Computer Interfaces*, P. Diez, Ed. Academic Press, 2018, pp. 171–195.
- [6] B. J. Lance, J. Touryan, Y.-K. Wang, S.-W. Lu, C.-H. Chuang, P. Khooshabeh, P. Sajda, A. Marathe, T.-P. Jung, C.-T. Lin *et al.*, "Towards serious games for improved bci," *Handbook of Digital Games and Entertainment Technologies*. Singapore: Springer, pp. 1–28, 2015.
- [7] C. Vidaurre, S. Haufe, T. Jorajuriá, K.-R. Müller, and V. V. Nikulin, "Sensorimotor functional connectivity: A neurophysiological factor related to bci performance," *Frontiers in Neuroscience*, vol. 14, p. 1278, 2020.
- [8] C. Vidaurre, R. Scherer, R. Cabeza, A. Schlögl, and G. Pfurtscheller, "Study of discriminant analysis applied to motor imagery bipolar data," *Medical & biological engineering & computing*, vol. 45, no. 1, p. 61, 2007.
- [9] Y. R. Tabar and U. Halici, "A novel deep learning approach for classification of eeg motor imagery signals," *Journal of neural engineering*, vol. 14, no. 1, p. 016003, 2016.
- [10] Y.-K. Wang, T.-P. Jung, and C.-T. Lin, "Eeg-based attention tracking during distracted driving," *IEEE transactions on neural systems and rehabilitation engineering*, vol. 23, no. 6, pp. 1085–1094, 2015.
- [11] Y. Zhang, C. S. Nam, G. Zhou, J. Jin, X. Wang, and A. Cichocki, "Temporally constrained sparse group spatial patterns for motor imagery bci," *IEEE Transactions on Cybernetics*, vol. 49, no. 9, pp. 3322–3332, Sep. 2019.
- [12] B. Blankertz, R. Tomioka, S. Lemm, M. Kawanabe, and K.-R. Müller, "Optimizing spatial filters for robust eeg single-trial analysis," *IEEE Signal processing magazine*, vol. 25, no. 1, pp. 41–56, 2008.
- [13] C. Sannelli, C. Vidaurre, K.-R. Müller, and B. Blankertz, "CSP patches: an ensemble of optimized spatial filters. an evaluation study," *Journal of Neural Engineering*, vol. 8, no. 2, p. 025012, 2011.
- [14] K.-R. Müller, C. W. Anderson, and G. E. Birch, "Linear and Non-Linear Methods for Brain-Computer Interfaces," *IEEE Transactions on Neural Systems and Rehabilitation Engineering*, vol. 11, no. 2, pp. 165–169, 2003.
- [15] T. Nierhaus, C. Vidaurre, C. Sannelli, K.-R. Müller, and A. Villringer, "Immediate brain plasticity after one hour of brain-computer interface (bci)," *The Journal of Physiology*, vol. n/a, no. n/a, 2019.
- [16] C. Vidaurre, A. R. Murguialday, S. Haufe, M. Gómez, K.-R. Müller, and V. Nikulin, "Enhancing sensorimotor bci performance with assistive afferent activity: An online evaluation," *NeuroImage*, vol. 199, pp. 375–386, 2019.
- [17] M. Mironovova and J. Bíla, "Fast fourier transform for feature extraction and neural network for classification of electrocardiogram signals," in *2015 Fourth International Conference on Future Generation Communication Technology (FGCT)*, 2015, pp. 1–6.
- [18] W. Zheng, W. Liu, Y. Lu, B. Lu, and A. Cichocki, "Emotionmeter: A multimodal framework for recognizing human emotions," *IEEE Transactions on Cybernetics*, vol. 49, no. 3, pp. 1110–1122, Mar. 2019.
- [19] M. Murugappan, S. Murugappan, C. Gerard *et al.*, "Wireless eeg signals based neuromarketing system using fast fourier transform (fft)," in *2014 IEEE 10th International Colloquium on Signal Processing and its Applications*. IEEE, 2014, pp. 25–30.
- [20] D. Iacoviello, A. Petracca, M. Spezialetti, and G. Placidi, "A classification algorithm for electroencephalography signals by self-induced emotional stimuli," *IEEE Transactions on Cybernetics*, vol. 46, no. 12, pp. 3171–3180, Dec. 2016.
- [21] L. Xie, Z. Deng, P. Xu, K. Choi, and S. Wang, "Generalized hidden-mapping transductive transfer learning for recognition of epileptic electroencephalogram signals," *IEEE Transactions on Cybernetics*, vol. 49, no. 6, pp. 2200–2214, Jun. 2019.
- [22] S. Bhattacharyya, A. Khasnobish, S. Chatterjee, A. Konar, and D. N. Tibrewala, "Performance analysis of lda, qda and knn algorithms in left-right limb movement classification from eeg data," in *2010 International Conference on Systems in Medicine and Biology*, Dec. 2010, pp. 126–131.
- [23] G. Dornhege, B. Blankertz, G. Curio, and K. Müller, "Boosting bit rates in noninvasive eeg single-trial classifications by feature combination and multiclass paradigms," *IEEE Transactions on Biomedical Engineering*, vol. 51, no. 6, pp. 993–1002, 2004.
- [24] L.-W. Ko, Y.-C. Lu, H. Bustince, Y.-C. Chang, Y. Chang, J. Fernandez, Y.-K. Wang, J. A. Sanz, G. P. Dimuro, and C.-T. Lin, "Multimodal fuzzy fusion for enhancing the motor-imagery-based brain computer interface," *IEEE Computational Intelligence Magazine*, vol. 14, no. 1, pp. 96–106, 2019.
- [25] J. Fumanal-Idocin, Y.-K. Wang, C.-T. Lin, J. Fernández, J. Sanz, and H. Bustince, "Motor-imagery-based brain computer interface using signal derivation and aggregation functions," *arXiv preprint arXiv:2101.06968*, 2021.
- [26] A. H. Altalhi, J. I. Forcén, M. Pagola, E. Barrenechea, H. Bustince, and Z. Takac, "Moderate deviation and restricted equivalence functions for measuring similarity between data," *Information Sciences*, vol. 501, pp. 19–29, 2019.
- [27] M. Papčo, I. Rodríguez-Martínez, J. Fumanal-Idocin, A. H., and H. Bustince, "A fusion method for multi-valued data," *Information Fusion*, vol. 71, pp. 1–10, 2021.
- [28] G. Beliakov, H. Bustince, and T. C. Sánchez, *A practical guide to averaging functions*. Springer, 2016, vol. 329.
- [29] G. Beliakov, A. Pradera, T. Calvo *et al.*, *Aggregation functions: A guide for practitioners*. Springer, 2007, vol. 221.
- [30] Y. Wang, S. Gao, and X. Gao, "Common spatial pattern method for channel selection in motor imagery based brain-computer interface," in *2005 IEEE Engineering in Medicine and Biology 27th Annual Conference*, 2005, pp. 5392–5395.
- [31] L. D. Miguel, D. Gómez, J. T. Rodríguez, J. Montero, H. Bustince, G. P. Dimuro, and J. A. Sanz, "General overlap functions," *Fuzzy Sets and Systems*, vol. 372, pp. 81–96, 2019, theme: Aggregation Operations.
- [32] A. Paszke, S. Gross, S. Chintala, G. Chanan, E. Yang, Z. DeVito, Z. Lin, A. Desmaison, L. Antiga, and A. Lerer, "Automatic differentiation in pytorch," 2017.
- [33] M. Tangermann, K.-R. Müller, A. Aertsen, N. Birbaumer, C. Braun, C. Brunner, R. Leeb, C. Mehring, K. J. Miller, G. Mueller-Putz *et al.*, "Review of the bci competition iv," *Frontiers in neuroscience*, vol. 6, p. 55, 2012.

Operador de comparación de elementos multivaluados basado en Funciones de Equivalencia Restringida

1º Aitor Castillo-López^a, 2º Carlos Lopez-Molina^{a,b,c}, 3º Javier Fernández^a,
4º Humberto Bustince^a y 5º Sesma-Sara^a

^aDepartamento de Estadística, Informática y Matemática, Universidad Pública de Navarra, (31006) Pamplona, Spain

^bKERMIT, Dept. Data Analysis and Mathematical Modelling, Ghent University, (9000) Ghent, Belgium

^cNavarraBiomed, Complejo Hospitalario de Navarra, (31008) Pamplona, Spain

aitor.castillo@unavarra.es

Abstract—En este trabajo proponemos un nuevo enfoque del algoritmo de clustering gravitacional basado en lo que Einstein consideró su “mayor error”: la constante cosmológica. De manera similar al algoritmo de clustering gravitacional, nuestro enfoque está inspirado en principios y leyes del cosmos, y al igual que ocurre con la teoría de la relatividad de Einstein y la teoría de la gravedad de Newton, nuestro enfoque puede considerarse una generalización del agrupamiento gravitacional, donde, el algoritmo de clustering gravitacional se recupera como caso límite. Además, se desarrollan e implementan algunas mejoras que tienen como objetivo optimizar la cantidad de iteraciones finales, y de esta forma, se reduce el tiempo de ejecución tanto para el algoritmo original como para nuestra versión.

Index Terms—cosmos, clustering, no supervisado, simulación, fuerza gravitacional

I. INTRODUCCIÓN

LOS OPERADORES DE COMPARACIÓN han sido objeto de estudio en el ámbito del procesamiento de la información. De hecho, la comparación (cuantitativa) de información es una de las tres operaciones más básicas sobre datos, junto a las operaciones de igualdad y ordenación. De manera general, la investigación se ha centrado en simular el comportamiento humano al realizar este tipo de operaciones.

Gran parte de la literatura sobre operadores de comparación está dedicada a las métricas, así como a clases de operadores estrechamente relacionados (pseudométricas [1], cuasimétricas [2], etc.). Una de las principales críticas a las métricas como operadores de comparación es el hecho de que se imponga la desigualdad triangular [3]. Si bien la desigualdad triangular es matemáticamente conveniente en una amplia gama de escenarios, no está claro si los humanos realmente se comportan de acuerdo con esta propiedad, y se pueden encontrar muchos contraejemplos diferentes en contextos específicos [4]. Por esta razón, los investigadores han intentado construir paradigmas de comparación que no se basen ni se inspiren en métricas.

Dentro de la teoría de conjuntos difusos, la comparación se ha abordado de diferentes maneras. Una parte importante se ha dedicado a la idea de métricas difusas [5], [6] o pseudométricas [7], [8], y las Funciones de Equivalencia

Restringida (REF por sus siglas en inglés), son en este ámbito de vital relevancia. Estas fueron presentadas en [9] para la comparación de grados de pertenencia en el intervalo $[0, 1]$ adaptando los axiomas originales propuestos por J. Fodor y M. Roubens [10]. Desde su introducción, el concepto de REF se ha adaptado a ámbitos en los que los valores a comparar se encuentran dentro del intervalo $[0, 1]$. Ejemplos relevantes son las REF intervalo-valoradas (IV-REF), diseñadas para comparar grados de pertenencia con intervalos [11], o las REF radiales (RREF), adaptadas para datos escalares en configuraciones radiales [12]. Una necesidad crítica en la adaptación de las REF a escenarios distintos al original es en relación al modelado del orden de crecimiento, que se utiliza críticamente en la definición axiomática de REF.

En este trabajo presentamos una adaptación de las REF a datos multivaluados, que denotamos como L^n . Para lograr este objetivo, presentamos la idea de L^n -REF, y desarrollamos un conjunto de axiomas que estos operadores deben cumplir. Además, introducimos métodos de construcción para L^n -REF capaces de acomodar diferentes interpretaciones en el ordenamiento multivaluado. Nuestras propuestas, en términos de ordenación de datos multivaluados, se circunscriben dentro de la taxonomía de Barnett [13]. Tengase en cuenta que imponer algún orden para los datos multivaluados es necesariamente arbitrario y dependiente del contexto, ya que *no hay un orden natural para los datos multivaluados* [13]. A modo de ejemplo, utilizamos las L^n -REF en la comparación de colores, ya que el color se representa de forma habitual como un dato multivaluado (independientemente del espacio de color específico). Además sirve como ejemplo de aplicación de las L^n -REF en visión artificial.

El resto de este trabajo está organizado de la siguiente manera. La Sección II resume algunos conceptos de uso para las próximas secciones. La Sección III presenta nuestra propuesta, mientras que diferentes ejemplos ilustrativos presentan una prueba de concepto de nuestros operadores en el contexto de la comparación de colores. Finalmente, en la Sección IV se recogen las conclusiones de nuestro trabajo y posibles líneas futuras.

II. PRELIMINARES

A. REF en L

Considérese $(L, \leq)$ donde $L = [0, 1]$ y $\leq$ es el orden natural de los números. Ahora bien, considérense en letra mayúscula los elementos pertenecientes a L^n , es decir $X = (x_1, \dots, x_n) \in L^n$ donde $n \in \mathbb{N}$. Existe un orden parcial $\leq_P$ inducido por $\leq$ dado de la siguiente manera: $X \leq_P Y$ iff $x_i \leq y_i$ para todo $i \in \{1, \dots, n\}$ ¹.

Denótese $\mathbf{0} = (0, \dots, 0) \in L^n$ y $\mathbf{1} = (1, \dots, 1) \in L^n$. Una función de agregación n -aria M de elementos multivaluados en L^n es una función $M : L^n \rightarrow L$ creciente para cada una de las variables y debe satisfacer $M(\mathbf{0}) = 0$, $M(\mathbf{1}) = 1$ [14]–[16]. Las siguientes propiedades para las funciones de agregación $M : L^n \rightarrow L$ son de utilidad para las siguientes secciones:

(P1) $M(x_1, \dots, x_n) = 0$ iff $x_1 = \dots = x_n = 0$.

(P2) $M(x_1, \dots, x_n) = 1$ iff $x_1 = \dots = x_n = 1$.

Una Media Aritmética Ponderada n -aria (MAP) en L con pesos normalizados $w_1, \dots, w_n \in L$ y $X = (x_1, \dots, x_n)$ es una función $\omega : L^n \rightarrow L$ definida como $\omega(X) = w_1 x_1 + \dots + w_n x_n$ tal que $w_1 + \dots + w_n = 1$.

Definición II.1. Un automorfismo de L es una función continua estrictamente creciente $\varphi : L \rightarrow L$ tal que $\varphi(0) = 0$ y $\varphi(1) = 1$. Además, la identidad en L se indica con Id .

Defínase una REF en L construida por automorfismos de la siguiente manera.

Definición II.2. [9] Una función $R : [0, 1]^2 \rightarrow [0, 1]$ es llamada Función de Equivalencia Restringida si cumple:

(R1) $R(x, y) = 1$ iff $x = y$;

(R2) $R(x, y) = 0$ iff $\{x, y\} = \{0, 1\}$;

(R3) $R(x, y) = R(y, x)$ para todo $x, y \in [0, 1]$;

(R4) Si $x \leq y \leq z$, entonces $R(x, z) \leq R(x, y)$ y $R(x, z) \leq R(y, z)$ para todo $x, y, z \in [0, 1]$.

En [17] se introduce un método para construir REFs en términos de automorfismos.

Proposición II.3. [17] Si φ_1, φ_2 son dos automorfismos de L , entonces la función $R : L^2 \rightarrow L$ definida como

$$R(x, y) = \varphi_1^{-1} \left(1 - |\varphi_2(x) - \varphi_2(y)| \right),$$

es una REF.

Definición II.4. Una función $f : (L^n)^m \rightarrow L^n$ es llamada representable si existen $f_1, \dots, f_n : L^m \rightarrow L$ tales que

$$f(X_1, \dots, X_m) = (f_1(x_{11}, \dots, x_{m1}), \dots, f_n(x_{1n}, \dots, x_{mn})) \quad (1)$$

para todo $X_1, \dots, X_m \in L^n$ with $X_i = (x_{i1}, \dots, x_{in})$ para todo $i \in \{1, \dots, m\}$.

¹A lo largo de este trabajo, la definición de i será la misma para abreviar. De lo contrario, se redefinirá explícitamente para algunas excepciones si así fuera necesario.

B. Espacios de color y operadores lineales

Como se expone en [18], para reproducir una imagen a color, es necesario generar nuevos vectores en el espacio espectral n a partir de los obtenidos por un sensor multi-espectral dado, mientras que los dispositivos de salida, que pueden caracterizarse como aditivos o sustractivos, debe poder reproducir colores de estos vectores. Dado que el ojo humano se puede representar como un sensor $n = 3$, lo más común es usar espacios tridimensionales para guardar información de color. Entre todos los diferentes sistemas de color posibles para dispositivos de salida aditiva, el modelo RGB [19] es el más extendido debido a la evolución de la codificación, que trata las imágenes a color como tres bandas monocromáticas independientes [20]. A pesar de que existen otros espacios de color como CIE $L^*a^*b^*$ que intentan reproducir un espacio de color uniforme basado en la percepción humana del color, donde el significado de cada valor L^* , a^* y b^* es bastante diferente. Por el contrario, el espacio RGB es un espacio simple en forma de cubo y el significado de cada valor que compone un *tripleto* es siempre el mismo, siendo la cantidad de rojo, verde o azul que compone un color respectivamente. Por simplicidad e idoneidad, tomaremos el modelo RGB como espacio de color en el que trabajar. Incluso si se puede elegir cualquier otro espacio de color, los resultados y las interpretaciones podrían ser bastante diferentes.

En RGB es habitual representar los tres componentes de un color en una escala de 0 a 255 y sólo se corresponde un único color para cada triplete en el espacio $[0, 255]^3$. Consideremos ahora el espacio RGB, pero refactorizado al espacio $[0, 1]^3$, es decir, L^3 . Entonces, los colores que corresponden a las esquinas del cubo RGB son:

■ Rojo: C_R (1, 0, 0)	■ Verde: C_G (0, 1, 0)	■ Azul: C_B (0, 0, 1)
■ Cian: C_C (0, 1, 1)	■ Magenta: C_M (1, 0, 1)	■ Amarillo: C_Y (1, 1, 0)
■ Negro: C_K (0, 0, 0) o $\mathbf{0}$	□ Blanco: C_W (1, 1, 1) o $\mathbf{1}$	

Estos son los únicos colores compuestos usando solo 0s y 1s y por esta razón los llamamos *tripletes crisp* o *colores crisp*.

Todos los colores posibles en el espacio RGB se pueden describir, de acuerdo con la terminología común, en los siguientes términos de apariencia de color [21]: *brillo*, *tonalidad* o *matiz* y *saturación*. Un color pierde brillo si disminuye el promedio de los valores del triplete que describe su posición en el espacio RGB. Entonces, todos los colores dentro de un plano perpendicular a la diagonal entre $\mathbf{0}$ y $\mathbf{1}$ tienen brillo idéntico. El tono del color es otra propiedad principal y está relacionada con la longitud de onda dominante percibida; en otras palabras, si un color es rojo, azul, amarillo, naranja... Los colores dentro de la diagonal entre $\mathbf{0}$ y $\mathbf{1}$ son los únicos colores sin matiz, y además, en esta diagonal se encuentran todos los colores grises posibles; por eso dicha recta se llama *diagonal de grises*. Finalmente, la saturación describe cuán cerca está

un triplete de la diagonal de grises; cuanto más cerca(lejos) de la diagonal de grises, menor(mayor) saturación. A mayor saturación, más fácil es percibir el tono de un color dado.

Denotamos *colores complementarios* a los pares de colores que, cuando se combinan, producen un determinado color en escala de grises [22]. Estos pares de colores se consideran complementarios dependiendo de la teoría de color que se utilice: la teoría de color moderna utiliza el modelo de color aditivo RGB o el modelo CMY para los sustractivos. Es por ello que en este trabajo consideramos *colores complementarios crisp* los pares $C_R - C_C$, $C_G - C_M$, y $C_B - C_Y$. El par de colores $C_K - C_W$ es común a todas las teorías del color.

A lo largo de todo este trabajo, se pueden aplicar diferentes $f : (L^n)^m \rightarrow L^n$ (como se presenta en Def. II.4) a un punto $X = (x_1, \dots, x_n) \in L^n$, es decir, $f(X \times m)$. Además, algunas funciones f se construyen con $m = n$ MAPs. En estos casos, la función f toma la forma

$$f(X \times n) = (\omega_1(X), \dots, \omega_n(X)) ,$$

que se puede reescribir como una transformación lineal de un punto $X \in L^n$ a otro punto $X' \in L^n$ mediante una multiplicación matricial;

$$\begin{aligned} (f(X \times n))^T &= \\ \begin{pmatrix} \omega_1(X) \\ \vdots \\ \omega_n(X) \end{pmatrix} &= \begin{pmatrix} w_{11}x_1 + \dots + w_{1n}x_n \\ \vdots \\ w_{n1}x_1 + \dots + w_{nn}x_n \end{pmatrix} = \\ \begin{pmatrix} w_{11} & \dots & w_{1n} \\ \vdots & & \vdots \\ w_{n1} & \dots & w_{nn} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} &= WX^T = X'^T , \end{aligned} \quad (2)$$

donde la fila i -ésima de la matriz W se compone de los pesos de la i -ésima MAP, ω_i .

Esta forma de representación nos permite tratar muchas de las futuras operaciones en este trabajo como operaciones lineales algebraicas en el espacio RGB y se utilizarán para:

- 1) Visualizar de una manera sencilla lo que sucede con el espacio de color, mientras,
- 2) La notación algebraica reduce la complejidad de las expresiones obtenidas, y además,
- 3) Simplifica la implementación del algoritmo presentado.

Nótese que para cualquier n el cubo L^n se transforma en un n -hiperparalelepípedo más pequeño dentro del cubo original L^n con dos de sus vértices opuestos en $\mathbf{0}$ y $\mathbf{1}$ debido a las propiedades (P1) y (P2). Además, sea $X_d = (x, \dots, x)$ cualquier punto en la diagonal de L^n , i.e., X_d está en la línea entre $\mathbf{0}$ y $\mathbf{1}$. Es fácil ver que X_d permanecerá inmutable bajo la transformación W presentada en la Eq. 2 debido a que cada una de las filas de W suma 1, entonces,

$$X_d'^T = WX_d^T = X_d^T. \quad (3)$$

En términos de color, la Eq. 3 implica que cualquier color gris es inmutable bajo cualquier W . Con respecto a los colores de fuera de la diagonal de grises, estos serán diferentes bajo una transformación W . De hecho, los colores que no son grises

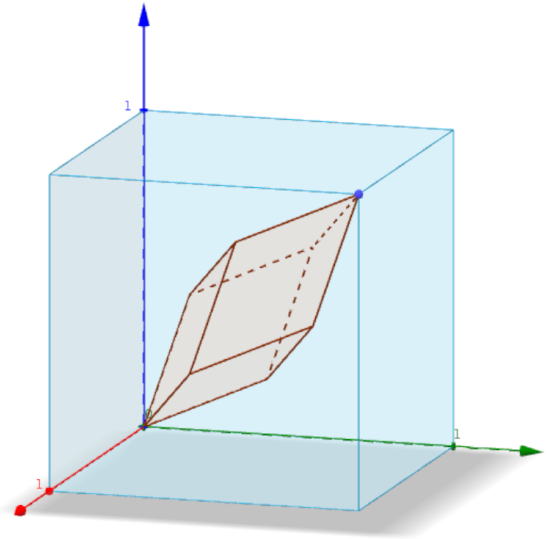


Fig. 1: Transformación del cubo L^3 para los valores del Ejemplo II.5. Es evidente que las propiedades del espacio resultante se pueden deducir de las propiedades de la matriz W como, por ejemplo, sus simetrías, el determinante, etc.

se aproximarán a la diagonal de la escala de grises cuando $\det(W) < 1$, por lo que los colores perderán saturación. En muchas otras transformaciones, estos colores también pueden cambiar su tono y brillo.

Ejemplo II.5. Tomemos la siguiente matriz W para un ejemplo en L^3 ,

$$W = \begin{pmatrix} \frac{1}{3} & \frac{1}{3} & \frac{1}{3} \\ \frac{1}{2} & \frac{1}{4} & \frac{1}{4} \\ \frac{1}{4} & \frac{1}{2} & \frac{1}{4} \end{pmatrix}.$$

Teniendo en cuenta las columnas de W , es fácil ver que los tripletes C_R , C_G y C_B , que son la base de nuestro espacio RGB, se transforman en $\blacksquare = (\frac{1}{3}, \frac{1}{2}, \frac{1}{4})$, $\blacksquare = (\frac{1}{3}, \frac{1}{4}, \frac{1}{2})$ y $\blacksquare = (\frac{1}{3}, \frac{1}{4}, \frac{1}{4})$ respectivamente. La representación visual de la transformación del espacio RGB se muestra en la Figura 1. Como se deduce de la Eq. 3, cualquier color de la escala de grises permanecerá invariante bajo transformación. También se puede apreciar cómo los colores de las esquinas C_R , C_G y C_B han perdido saturación, y además, han cambiado de tono debido a la rotación inducida por W . Los colores C_R y C_G han ganado brillo, mientras C_B lo ha perdido.

Como resumen, si aplicamos esta W a los colores que forman la base del espacio RGB, C_R , C_G y C_B , se transformarán en

$$\text{RGB} \rightarrow \text{transformed colors}.$$

III. L^n -REF: NUEVO OPERADOR PARA COMPARAR DOS COLORES OBTENIENDO UN TERCERO COMO RESULTADO

El objetivo de esta sección es construir un operador que sea capaz de devolver una medida de equivalencia en L^n como resultado de comparar dos elementos multivaluados pertenecientes a dicho espacio. Para ello, mientras la teoría es desarrollada para cualquier n , los ejemplos ilustrativos se darán para $n = 3$, y en concreto, se presentará un algoritmo de comparación de imágenes. Este, dadas dos imágenes, debe devolver una imagen de salida donde sus píxeles pueden ser interpretados como un mapa, donde sea posible interpretar las regiones donde las imágenes de entrada son similares en color o, en su defecto, diferentes. Para lograr este objetivo, introducimos un algoritmo para comparar pixel a pixel todos los píxeles de una imagen con los píxeles correspondientes en una segunda imagen de entrada, siendo necesario que, ambas tengan las mismas dimensiones.

Definamos el pixel ij -ésimo de una imagen A como $P_{A_{ij}}$. En el Algoritmo 1 presentamos la estructura computacional de nuestra propuesta.

Algoritmo 1: Comparación pixel a pixel

Entrada: Dos imágenes A y B del mismo tamaño.

Resultado: Una imagen a color de tamaño idéntico.

Escoger un método de comparación de píxeles ;

para cada posición ij **hacer**

 Comparar los píxeles $P_{A_{ij}}$ y $P_{B_{ij}}$;

 Asignar al ij -ésimo pixel de la imagen resultado el valor de la comparación entre píxeles anterior;

fin

Por lo tanto, para la acción de comparación del Alg. 1 es necesario construir un método para comparar píxeles, es decir, encontrar un operador que dados dos tripletes de entrada su salida sea otro triplete que guarde la información de la comparación.

Llegados a este punto, podemos considerar dos filosofías diferentes teniendo en cuenta cuáles de los tripletes crisp complementarios son antagonistas para el caso del color. Es posible construir dos métodos diferentes, uno para cada una de las siguientes filosofías:

- 1) Un método que trata el par de colores complementarios $C_K - C_W$ de manera diferente a los otros pares crisp complementarios al compararlos, siendo $C_K - C_W$ el único par que presenta el mayor antagonismo.
- 2) Un método que trata todos los pares crisp complementarios por igual al compararlos.

Para comparar dos píxeles de color, debemos definir la noción de REF sobre L^3 (en general sobre L^n), de tal forma que el valor que obtengamos sea nuevamente un elemento de $L^3(L^n)$. En la siguiente subsección definimos y presentamos una versión de REF en L^n con este propósito. De las dos filosofías posibles, en este trabajo solamente se ha desarrollado la primera planteada donde el par $C_K - C_W$ se trata de diferente manera que los restantes.

A. L^n -REF basadas en la filosofía donde el par $0 - 1$ es tratado de diferente manera

Esta versión de REF trata el par $C_K - C_W$ como el único que representa la menor equivalencia posible.

Definición III.1. Sea n un número entero positivo. Una función $R_{L^n} : L^n \times L^n \rightarrow L^n$ se llama función de equivalencia restringida en L^n (R_{L^n}), si satisface:

(RL1) $R_{L^n}(X, Y) = \mathbf{1}$ iff $X = Y$;

(RL2) $R_{L^n}(X, Y) = \mathbf{0}$ iff $\{X, Y\} = \{0, 1\}$;

(RL3) $R_{L^n}(X, Y) = R_{L^n}(Y, X)$ para todo $X, Y \in L^n$;

(RL4) Si $X \leq_P Y \leq_P Z$, entonces $R_{L^n}(X, Z) \leq_P R_{L^n}(X, Y)$ y $R_{L^n}(X, Z) \leq_P R_{L^n}(Y, Z)$ para todo $X, Y, Z \in L^n$.

Teniendo en cuenta Def. II.2, la justificación del axioma (RL1) es natural; comparar tripletes equivalentes debe devolver el valor más alto en L^n , i.e., $\mathbf{1}$, como medida de equivalencia dado que estos son equivalentes. En cuanto al axioma (RL2), se trata de comparar los tripletes $\mathbf{0}$ y $\mathbf{1}$. En este caso, estamos teniendo en cuenta que los tripletes C_K y C_W son los tripletes menos equivalentes entre todos los pares de colores posibles según la filosofía escogida, entonces, siendo este nuestro punto de partida, se justifica que son los tripletes únicos que al compararse, deben devolver $\mathbf{0}$, es decir, el valor más bajo posible de todo L^n . La justificación del axioma (RL3) es que se exige a la REF que la comparación entre dos píxeles debe cumplir con la simetría. Esta propiedad puede no ser necesaria en otras aplicaciones donde existe una dependencia entre imágenes comparadas (como imágenes de vídeo donde existe una relación temporal entre fotogramas), en nuestro caso, no se considera dependencia del tiempo, por lo que (RL3) está justificado. Finalmente, la justificación del axioma (RL4) es que la equivalencia resultante entre comparar dos tripletes similares debe ser mayor que la equivalencia de comparar dos tripletes que son, al menos, más diferentes que los anteriores.

Ahora damos un método de construcción para R_{L^n} en L^n de acuerdo con Def. III.1.

Teorema III.2. Sea $\omega_i : L^n \rightarrow L$ una MAP n -aria como vector de pesos normalizados $(w_{i1}, \dots, w_{in})$ tal que los vectores son linealmente independientes y existe $k \in \{1, \dots, n\}$ con $w_{kj} \neq 0$ para todo $j \in \{1, \dots, n\}$. Sea $\mathbf{R} = (R_1, \dots, R_n)$ una secuencia de REFs en L . Entonces la función $R_{L^n} : L^n \times L^n \rightarrow L^n$ dada por,

$$R_{L^n}(X, Y) = \left(R_1(\omega_1(X), \omega_1(Y)), \dots, R_n(\omega_n(X), \omega_n(Y)) \right), \quad (4)$$

para todo $X, Y \in L^n$, es una R_{L^n} en L^n .

Proof. (RL1) La suficiencia se deriva de Eq. (4). Con respecto a la necesidad, sea $R_{L^n}(X, Y) = \mathbf{1}$, entonces, $R_i(\omega_i(X), \omega_i(Y)) = 1$, por lo tanto $\omega_i(X) = \omega_i(Y)$ y finalmente $X = Y$.

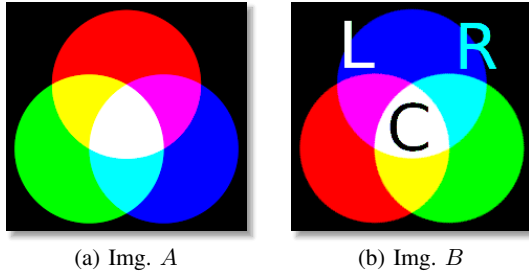


Fig. 2: Las imágenes sintéticas A (Fig. 2a) y B (Fig. 2b) son el input para el Algoritmo 1 en el Ejemplo III.4.

(RL2) Sea $R_{L^n}(X, Y) = \mathbf{0}$. Por lo tanto para todo i , se tiene que $R_i(\omega_i(X), \omega_i(Y)) = 0$, de modo que $\{\omega_i(X), \omega_i(Y)\} = \{0, 1\}$. Ya que existe k tal que ω_k satisface (P1) y (P2), se sigue que $\{X, Y\} = \{0, 1\}$.

(RL3) La prueba es directa teniendo en cuenta (R3).

(RL4) De $x_i \leq y_i \leq z_i$ se obtiene $\omega_j(X) \leq \omega_j(Y) \leq \omega_j(Z)$ para todo j y consecuentemente $R_j(\omega_j(X), \omega_j(Z)) \leq R_j(\omega_j(X), \omega_j(Y))$, de donde se sigue que $R_{L^n}(X, Z) \leq_P R_{L^n}(X, Y)$. La prueba para $R_{L^n}(X, Z) \leq_P R_{L^n}(Y, Z)$ es similar. $\square$

Si aplicamos el método de construcción de las REFs $R_1, \dots, R_n$ en el Teorema. III.2 en términos de automorfismos se obtiene que al seguir la construcción de R_{L^n} .

Corolario III.3. Si se asume el Teorema III.2, sea φ_{ij} para $j = 1, 2$, un automorfismo de L . Luego la función $R_{L^n} : L^n \times L^n \rightarrow L^n$ dada por,

$$R_{L^n}(X, Y) = \left(\varphi_{11}^{-1} \left(1 - |\varphi_{12}(\omega_1(X)) - \varphi_{12}(\omega_1(Y))| \right), \dots, \varphi_{n1}^{-1} \left(1 - |\varphi_{n2}(\omega_n(X)) - \varphi_{n2}(\omega_n(Y))| \right) \right), \quad (5)$$

para todo $X, Y \in L^n$, es una R_{L^n} en L^n .

Nótese que es posible reescribir la anterior expresión en términos de Eq. 2 dado el caso en el que ω_i son MAPs,

$$R_{L^n}(X, Y) = \mathbf{R}((WX^T)^T, (WY^T)^T), \quad (6)$$

donde W es la matriz construida con los pesos de las MAP y $\mathbf{R}$ es una secuencia de REFs en L .

Tomemos los colores crisp presentados en la Subsección II-B y compongamos las imágenes sintéticas A y B (Figuras 2a y 2b respectivamente). Ambas son la misma imagen pero una rotada con respecto a la otra, y además, la segunda contiene tres caracteres adicionales: el carácter blanco "L" a la izquierda, el carácter cian "R" a la derecha y un tercer carácter negro "C" en el centro.

Ejemplo III.4. Las imágenes sintéticas Fig. 2a y Fig. 2b nos permiten comparar muchos de los colores crisp entre sí al

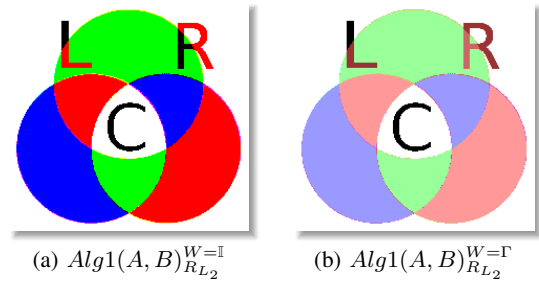


Fig. 3: Mapas resultantes para $W = \mathbb{I}$ (Fig. 3a), y para $W = \Gamma$ (Fig. 3b).

usar el Alg. 1. En este ejemplo, el resultado se calcula dos veces; uno para $W = \mathbb{I}$ y otro para $W = \Gamma$, siendo

$$\Gamma = \begin{pmatrix} 0.6 & 0.2 & 0.2 \\ 0.2 & 0.6 & 0.2 \\ 0.2 & 0.2 & 0.6 \end{pmatrix},$$

donde, en ambos casos la secuencia de REFs elegida es $\mathbf{R} = (1 - |x - y|, 1 - |x - y|, 1 - |x - y|)$.

Los mapas de características resultantes se muestran en la Fig. 3. Se observa cómo al comparar C_R y C_B devuelve el mismo tono verde que al comparar C_C y C_Y para diferentes valores de W sin rotación. Esta simetría se debe a que el canal que permanece inalterado es el segundo. Véase como, tomando $C_C = (0, 1, 1)$ y $C_Y = (1, 1, 0)$ el canal con valor equivalente sigue siendo el segundo. Entonces, esto se puede interpretar como; los tonos de color del mapa resultante identifican los canales que son más equivalentes entre los pixeles comparados, haciendo coincidir el tono del pixel resultante con sus canales correspondientes cuando no hay rotación inducida por W .

Ejemplo de (RL1) es que el fondo de ambos mapas es C_W porque se compara el mismo color; en ambos casos, el color C_K está como fondo. Además, el área circundante del carácter "C" en los mapas es C_W porque en ambas imágenes de entrada se compara el color C_W . Por el contrario, teniendo en cuenta (RL2) C_K se obtiene en el mapa resultante para el carácter "C" y la parte superior del carácter "L" al comparar $\mathbf{0}$ con $\mathbf{1}$.

El color en la parte izquierda de "R" es el resultado de comparar dos pares de colores crisp complementarios ($C_R - C_C$) y depende de W . Para $W = \mathbb{I}$ (en la Fig. 3a), se obtiene C_K y lo mismo sucede si se comparan cualquiera de los colores crisp complementarios en el cubo RGB; $C_K - C_W$, $C_R - C_C$, $C_G - C_M$ o $C_B - C_Y$. Esto implica que todos los pares crisp complementarios tienen el mismo comportamiento con respecto a la REF que va en contra de la filosofía establecida. Entonces, debido a esta razón, W no puede ser la matriz de identidad para R_{L^n} . En otras palabras, esto significa que C_K se puede obtener en más casos que comparando exclusivamente C_K con C_W , entonces, el axioma RL2 no se cumple para $W = \mathbb{I}$ (o, en cualquier caso, cuando $\det(W) = 1$).

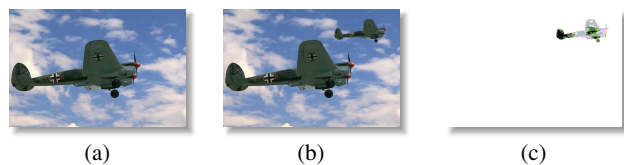


Fig. 4: La imagen 3063 original (Fig. 4a) y su versión modificada (Fig. 4b) para usar como input del Alg. 1 en el Ejemplo III.5. El mapa resultante se muestra en la Fig. 4c.

Para concluir este trabajo, mostramos un ejemplo donde aplicamos el Alg. 1 usando R_{L^n} . Este Ejemplo III.5 es una posible aplicación del Algoritmo 1 donde se detecta un objeto en función de la diferencia de color entre las imágenes de entrada. Con este propósito, hemos tomado del Berkeley Segmentation Dataset [23] la imagen indexada como 3063 (Fig. 4a) y la hemos modificado para agregar un segundo avión en la esquina superior derecha como se muestra en la Fig. 4b.

Ejemplo III.5. La R_{L^n} tal y como aparece en la Eq. 6 es usada para comparar mediante el Alg. 1 las imágenes de la Fig. 4 tomando como secuencia de REFs $\mathbf{R} = (1 - |x - y|, 1 - |x - y|, 1 - |x - y|)$. Por simplicidad W es la misma que en el Ejemplo II.5. El mapa resultante se muestra en la Fig. 4c.

IV. LÍNEAS FUTURAS

Una evidente línea futura es generar una nueva clase de operadores L^n -REF que se basen en la filosofía en la cual todos los posibles complementarios crisp sean tratados de la misma manera, por ejemplo, deberán tomar como igualmente antagónicos todos los elementos multivaluados que sean complementarios crisp y no acentuar la diferencia para el único par $0 - 1$.

En concreto, la creación de este tipo de operador nos llevaría a tener que analizar qué diferencias se podrían apreciar entre la aplicación de los operadores construidos en base a las diferentes filosofías en el procesamiento de imagen a color.

Otra posible línea futura es extender las aplicaciones a otro tipo de formato de información como pueden ser las imágenes multi e hiperespectrales.

AGRADECIMIENTOS

Este trabajo ha sido respaldado por los proyectos PID2019-108392GB-I00 (AEI/10.13039/ 501100011033) de la Agencia Estatal de Investigación.

REFERENCES

- [1] F. Van Breugel, J. Worrell, A behavioural pseudometric for probabilistic transition systems, *Theoretical Computer Science* 331 (1) (2005) 115–142.
- [2] C. Lopez-Molina, S. Iglesias-Rey, H. Bustince, B. De Baets, On the role of distance transformations in baddeley's delta metric, *Information Sciences* (2021).
- [3] A. Tversky, I. Gati, Similarity, separability, and the triangle inequality., *Psychological review* 89 (2) (1982) 123.
- [4] A. Tversky, D. H. Krantz, The dimensional representation and the metric structure of similarity data, *Journal of mathematical psychology* 7 (3) (1970) 572–596.
- [5] I. Kramosil, J. Michálek, Fuzzy metrics and statistical metric spaces, *Kybernetika* 11 (5) (1975) 336–344.
- [6] O. Kaleva, S. Seikkala, On fuzzy metric spaces, *Fuzzy sets and systems* 12 (3) (1984) 215–229.
- [7] Z. Deng, Fuzzy pseudo-metric spaces, *Journal of Mathematical Analysis and Applications* 86 (1) (1982) 74–95.
- [8] Y. Yue, F.-G. Shi, On fuzzy pseudo-metric spaces, *Fuzzy Sets and Systems* 161 (8) (2010) 1105–1116.
- [9] H. Bustince, E. Barrenechea, M. Pagola, Restricted equivalence functions, *Fuzzy Sets and Systems* 157 (17) (2006) 2333–2346.
- [10] J. C. Fodor, M. Roubens, *Fuzzy preference modelling and multicriteria decision support*, Vol. 14, Springer Science & Business Media, 1994.
- [11] A. Jurio, M. Pagola, D. Paternain, C. Lopez-Molina, P. Melo-Pinto, Interval-valued restricted equivalence functions applied on clustering techniques, in: *Proc. of the IFSA-EUSFLAT*, 2009, pp. 831–836.
- [12] C. Marco-Detchart, J. Cerron, L. De Miguel, C. Lopez-Molina, H. Bustince, M. Galar, A framework for radial data comparison and its application to fingerprint analysis, *Applied Soft Computing* 46 (2016) 246–259.
- [13] V. Barnett, The ordering of multivariate data, *Journal of the Royal Statistical Society: Series A* 139 (1976) 318–355.
- [14] H. Bustince, J. Montero, E. Barrenechea, M. Pagola, Semiautoduality in a restricted family of aggregation operators, *Fuzzy Sets and Systems* 158 (12) (2007) 1360–1377.
- [15] D. Dubois, H. Prade, A review of fuzzy set aggregation connectives, *Information sciences* 36 (1-2) (1985) 85–121.
- [16] G. J. Klir, *Fuzzy sets, Uncertainty and Information* (1988).
- [17] H. Bustince, M. Pagola, E. Barrenechea, Construction of fuzzy indices from fuzzy di-subsethood measures: application to the global comparison of images, *Information Sciences* 177 (3) (2007) 906–929.
- [18] H. J. Trussell, E. Saber, M. Vrhel, Color image processing: Basics and special issue overview, *IEEE signal processing magazine* 22 (1) (2005).
- [19] T. Mélangé, *Fuzzy techniques for noise removal in image sequences and interval-valued fuzzy mathematical morphology*, Ph.D. thesis, Ghent University (2010).
- [20] G. Sharma, H. J. Trussell, Digital color imaging, *IEEE transactions on image processing* 6 (7) (1997) 901–932.
- [21] G. Wyszecki, W. S. Stiles, *Color science: concepts and methods, quantitative data and formulas* (1982).
- [22] J. Cohen, *Visual color and color mixture: The fundamental color space*, University of Illinois Press, 2001.
- [23] P. Arbelaez, M. Maire, C. Fowlkes, J. Malik, Contour detection and hierarchical image segmentation, *IEEE Trans. on Pattern Analysis and Machine Intelligence* 33 (2011) 898–916.

XX Congreso Español sobre Tecnologías y Lógica Fuzzy

TRABAJOS DESTACADOS



Computing with Comparative Linguistic Expressions and Symbolic Translation for Decision Making: ELICIT Information

1st Álvaro Labella

Department of Computer Science
University of Jaén
Jaén, Spain
alabella@ujaen.es

2nd Rosa M. Rodríguez

Department of Computer Science
University of Jaén
Jaén, Spain
rmrodrig@ujaen.es

3rd 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 introduces a new fuzzy linguistic representation model for comparative linguistic expressions that takes advantage of the goodness of the 2-tuple linguistic representation model and improves the readability and accuracy of the results in computing with words processes, resulting the so-called extended comparative linguistic expressions with symbolic translation.

Index Terms—Decision making, computing with words, symbolic translation, comparative linguistic expressions

I. INTRODUCTION

Nowadays, *Decision Making* (DM) problems are defined in changing contexts in which uncertainty and vagueness are quite common. The *fuzzy linguistic approach* [2] has been used successfully to model such uncertainty by means of linguistic information, giving rise to the so-called *Linguistic Decision Making*. Consequently, the use of linguistic information implies to perform computations with it in order to solve decision problems under uncertainty. *Computing with Words* (CW) [3] methodology carries out processes “where words and not numbers are used for computing” and mimics the human beings’ reasoning process in which, from linguistic premises, provides linguistic results. In this way, CW methodology guarantees interpretable results.

There are several proposal that try to follow a CW approach. One of the most widely used in DM is the 2-tuple linguistic model [4] which presents interpretable and precise linguistic results, the latter thanks to use of the symbolic translation concept. However, such results are represented by single linguistic terms that are inadequate to represent experts’ hesitancy. To overcome the latter limitation, Rodríguez et al. [5] introduced the *Hesitant Fuzzy Linguistic Term Set* (HFLTS) that, together with the use of context-free grammars, allow to generate *Comparative Linguistic Expressions* (CLEs) close to the natural language used by human beings. However, the existing computational model that makes use of CLEs does not provide interpretable results.

Therefore, in order to overcome previous drawbacks, we propose a new fuzzy linguistic representation model that

extends the CLEs by using the concept of symbolic translation introduced by the 2-tuple linguistic model resulting the so-called *Extended Comparative Linguistic Expressions with Symbolic Translation* (ELICIT) information. These expressions extend the representation of CLEs generated by a context-free grammar into a continuous domain to perform CW processes without any kind of approximation. The proposed context-free grammar to generate ELICIT information is described below:

Definition 1: [1] Let G_H be a context-free grammar and $S = \{s_0, \dots, s_g\}$ a linguistic term set. The elements of $G_H = (V_N, V_T, I, P)$ are defined as follows.

$V_N = \{(continuous\ primary\ term), (composite\ term), (unary\ relation), (binary\ relation), (conjunction)\}$

$V_T = \{at\ least, at\ most, between, and, (s_0, \alpha)^\gamma, (s_1, \alpha)^\gamma, \dots, (s_g, \alpha)^\gamma\}$

$I \in V_N$

$P = \{I ::= (continuous\ primary\ term) | (composite\ term)$

$(composite\ term) ::= (unary\ relation)$

$(continuous\ primary\ term) |$

$(binary\ relation)(continuous\ primary\ term)$

$(conjunction)(continuous\ primary\ term)$

$(continuous\ primary\ term) ::= (s_0, \alpha)^\gamma |$

$(s_1, \alpha)^\gamma | \dots | (s_g, \alpha)^\gamma$

$(unary\ relation) ::= at\ least | at\ most$

$(binary\ relation) ::= between$

$(conjunction) ::= and\}$

Thus, the possible ELICIT expressions generated according to the new definition of the context-free grammar are: “at least $(s_i, \alpha)^\gamma$ ”, “at most $(s_i, \alpha)^\gamma$ ” and “between $(s_i, \alpha_1)^{\gamma_1}$ and $(s_j, \alpha_2)^{\gamma_2}$ ”

A CW approach for ELICIT information has been also proposed in [1]. Such a CW approach obtains linguistic results

modeled by ELICIT information from linguistic inputs represented by CLEs and ELICIT information. To carry out CW processes, linguistic inputs are transformed into *trapezoidal fuzzy numbers* (TrFNs), which are manipulated by means of fuzzy parametric operations. Whereas CLEs are transformed into TrFNs by means of the fuzzy envelope, such a transformation for ELICIT information is carried out by means of the function ζ^{-1} .

Definition 2: [1] Let x_{el} be an ELICIT expression and $T(a, b, c, d)$ a trapezoidal fuzzy number. The function ζ^{-1} is defined as:

$$\zeta^{-1} : x_{el} \rightarrow T(a, b, c, d) \quad (1)$$

The manipulation of the fuzzy envelopes through fuzzy parametric operations provides new fuzzy numbers noted as β . Now, it is necessary a retranslation process to obtain ELICIT information from TrFNs. This process consists of different steps:

- 1) *Identify the relation:* the relation is determined by the fuzzy number β and the ζ function, defined as follows:

Definition 3: Let $S = \{s_0, \dots, s_g\}$ be a set of linguistic terms and β a fuzzy number. The function ζ is given by

$$\zeta(\beta) = x_{el}, \text{ where } \begin{cases} x_{el} = \text{at least } (s_i, \alpha)^\gamma \text{ if } \beta = T(a, b, 1, 1) \\ x_{el} = \text{at most } (s_i, \alpha)^\gamma \text{ if } \beta = T(0, 0, c, d) \\ x_{el} = \text{between } (s_i, \alpha_1)^{\gamma_1} \text{ and } (s_j, \alpha_2)^{\gamma_2} \\ \text{if } \beta = T(a, b, c, d) \end{cases}$$

Henceforth, for sake of space, it is assumed that the ELICIT expression is “*between* $(s_i, \alpha_1)^{\gamma_1}$ *and* $(s_j, \alpha_2)^{\gamma_2}$ ”.

- 2) *2-tuple linguistic terms computation:* the process of obtaining the two continuous primary terms $(s_i, \alpha_1)^{\gamma_1}$ and $(s_j, \alpha_2)^{\gamma_2}$ is divided into different steps:

- a) *Compute linguistic terms:* select the linguistic terms s_i and $s_j \in S$, $i, j \in \{0, \dots, g\}$, whose distance between the coordinates x of their respective centroids, $\bar{x}_i$ and $\bar{x}_j$, and the points b and c belonging to β is minimal.

$$\begin{aligned} i &= \arg \min_h |b - \bar{x}_h|, \quad h \in \{0, \dots, g\} \\ j &= \arg \min_h |c - \bar{x}_h|, \quad h \in \{0, \dots, g\} \end{aligned} \quad (2)$$

When this step finishes, the ELICIT expression so far is “*between* $(s_i, ?)^?$ *and* $(s_j, ?)^?$ ”.

- b) *Compute symbolic translations:* according to [4], $1/2g$ represents the distance equivalent to a symbolic translation equal to 0.5 in S , where $g + 1$ is the cardinality of S :

$$\begin{aligned} \alpha_1 &= g \cdot (b - \bar{x}_i) \quad \alpha_1 \in [-0.5, 0.5] \\ \alpha_2 &= g \cdot (c - \bar{x}_j) \quad \alpha_2 \in [-0.5, 0.5] \end{aligned} \quad (3)$$

When this step finishes, the ELICIT expression so far is “*between* $(s_i, \alpha_1)^?$ *and* $(s_j, \alpha_2)^?$ ”.

- 3) *Compute adjustments:* the *adjustment* is an additional parameter included in the ELICIT expression, which allows to keep information related to the fuzzy number

β . This parameter will be used to obtain the fuzzy number β from an ELICIT expression by using its inverse function, ζ^{-1} . The steps to compute the adjustments for the ELICIT expression are:

- a) *Compute HFLTS:* the HFLTS of an ELICIT expression whose relation is *between* would be composed by:

$$E_{ELICIT}(\text{between } (s_i, \alpha) \text{ and } (s_j, \alpha)) = \{s_k \mid (s_i, \alpha) \text{ and } (s_j, \alpha), s_i < s_k < s_j \text{ where } s_k \in S\}$$

- b) *Compute fuzzy envelope:* the fuzzy envelope, $T_{ELICIT} = T(a', b', c', d')$, of the former HFLTS is computed.

- c) *Compute adjustments γ_1 and γ_2 :* the adjustments γ_1 and γ_2 are determined by the subtraction between the points a and d of $\beta = T(a, b, c, d)$ and the points a' and d' of $T(a', b', c', d')$, so that:

$$\begin{aligned} \gamma_1 &= a - a' \quad \gamma_1 \in [-1, 1] \\ \gamma_2 &= d - d' \quad \gamma_2 \in [-1, 1] \end{aligned} \quad (4)$$

When this step finishes, the ELICIT expression is completed “*between* $(s_i, \alpha_1)^{\gamma_1}$ *and* $(s_j, \alpha_2)^{\gamma_2}$ ”.

II. CONCLUSIONS

The need of a new fuzzy linguistic representation model that overcomes the existing limitations in previous linguistic models either from the point of view of interpretability or/and accuracy has resulted in the ELICIT representation model and its CW approach. This new linguistic model makes use of the ELICIT information, CLEs extended to a continuous domain by means of the symbolic translation concept. In this way, it is possible to carry out CW processes with high accuracy and interpretability.

As future works, we will study the definition of new aggregation operators for ELICIT information. Another aim is to apply this new type of information to consensus reaching processes.

ACKNOWLEDGMENT

This work is partially supported by the Spanish Ministry of Economy and Competitiveness through the Spanish National Research Project PGC2018-099402-B-I00 and the Postdoctoral fellow Ramón y Cajal (RYC-2017-21978).

REFERENCES

- [1] Á. Labella, R. M. Rodríguez, and L. Martínez, “Computing with comparative linguistic expressions and symbolic translation for decision making: Elicit information,” *IEEE Transactions on Fuzzy Systems*, vol. 28, no. 10, pp. 2510–2522, 2019.
- [2] L. A. Zadeh, “The concept of a linguistic variable and its application to approximate reasoning (i),” *Information Sciences*, vol. 8, no. 3, pp. 199 – 249, 1975.
- [3] —, “Fuzzy logic = computing with words,” *IEEE Transactions on Fuzzy Systems*, vol. 4, no. 2, pp. 103–111, 1996.
- [4] L. Martínez, R. M. Rodríguez, and F. Herrera, *The 2-tuple Linguistic Model*. Springer International Publishing, 2015.
- [5] R. M. Rodríguez, L. Martínez, and F. Herrera, “A group decision making model dealing with comparative linguistic expressions based on hesitant fuzzy linguistic term sets,” *Information Sciences*, vol. 241, no. 1, pp. 28–42, 2013.

Aggregation of fuzzy quasi-metrics

Tatiana Pedraza

Instituto Universitario de Matemática Pura y Aplicada
Universitat Politècnica de València
Valencia, Spain
tapedraz@mat.upv.es

Jesús Rodríguez-López

Instituto Universitario de Matemática Pura y Aplicada
Universitat Politècnica de València
Valencia, Spain
jrlopez@mat.upv.es

Óscar Valero

Departament de Ciències Matemàtiques i Informàtica
Universitat de les Illes Balears
Palma, Spain
o.valero@uib.es

Abstract—In this work we will present some results about the aggregation of fuzzy (quasi-)metrics that appear in [9]. Concretely, we provide a characterization of functions that merge a collection of fuzzy (quasi-)metrics into a single one in terms of $*$ -triangular triplets, isotonicity and $*$ -supmultiplicativity, where $*$ is a t -norm. We also show that, in contrast to the crisp case, this characterization does not depend on the symmetry of the fuzzy quasi-metrics. We also stress that these results are not only interesting from the aggregation theory viewpoint but also because they allow to generate examples of fuzzy (quasi-)metric spaces that are not easy to obtain. Moreover, from our results we can infer others about the aggregation of fuzzy preorders and indistinguishability operators.

I. INTRODUCTION

The problem of aggregating several structures of the same type into a single one has received a lot of attention in the last years. In this way, we can find results about the aggregation of: metrics [1]; quasi-metrics [7]; norms [4]; asymmetric norms [5]; fuzzy binary relations [11], [2]; indistinguishability operators [6]; etc. In the following, we will present some results about the aggregation of an important fuzzy structure: the fuzzy quasi-metrics. This concept has its origins in the probabilistic metric spaces introduced by Menger [8] in 1942 who gave a probabilistic interpretation of the concept of distances and proposed to associate a distribution function with a pair of elements, instead of associating a number.

In some sense, fuzzy binary relations can be considered as a particular class of fuzzy (quasi-)metrics. We recall its definition as well as one of its most important classes.

Definition 1 ([10]). A fuzzy binary relation on a nonempty set X is a map $E : X \times X \rightarrow [0, 1]$.

If a fuzzy binary relation E on X satisfies for all $x, y, z \in X$:

- $E(x, x) = 1$ (reflexivity)
- $E(x, y) = E(y, x)$ (symmetry)
- $E(x, y) * E(y, z) \leq E(x, z)$ ($*$ -transitivity)

where $*$ is a triangular norm, then it is called an $*$ -indistinguishability operator.

In the literature we can find several papers [2], [3], [6], [11] studying functions which preserve $*$ -transitivity of fuzzy binary relations in the following two, a priori, different senses:

Definition 2 (cf. [6], [11]). Let I be a set of indices and $F : [0, 1]^I \rightarrow [0, 1]$ be a function. We say that:

- F preserves $*$ -transitivity of fuzzy binary relations on products if whenever $\{(X_i, E_i) : i \in I\}$ is a family of nonempty sets X_i endowed with $*$ -transitive fuzzy binary relations E_i for all $i \in I$, then $F \circ \mathbf{E}$ is an $*$ -transitive fuzzy binary relation on $\prod_{i \in I} X_i$ where $\mathbf{E} : (\prod_{i \in I} X_i)^2 \rightarrow [0, 1]^I$ is given by

$$\tilde{\mathbf{E}}(a, b)_i = E_i(a_i, b_i) \quad \text{for all } i \in I.$$

- F preserves $*$ -transitivity of fuzzy binary relations on sets if whenever $\{E_i : i \in I\}$ is a family of $*$ -transitive fuzzy binary relations E_i on a fixed nonempty set X for all $i \in I$, then $F \circ \mathbf{E}$ is an $*$ -transitive fuzzy binary relation on X where $\mathbf{E} : X^2 \rightarrow [0, 1]^I$ is given by

$$\mathbf{E}(a, b)_i = E_i(a, b) \quad \text{for all } i \in I.$$

For studying functions which aggregate fuzzy quasi-metrics it is important to take into account the characterization of those functions which preserve $*$ -transitivity of fuzzy binary relations, since fuzzy (quasi-)metrics satisfy a property near to $*$ -transitivity. Surprisingly, we can prove that there is no difference between those functions which preserve $*$ -transitivity of fuzzy binary relations on products and those which preserve $*$ -transitivity of fuzzy binary relations on sets.

Definition 3. Let $*$ be a t -norm and I be a set of indices. A triplet $(a, b, c) \in ([0, 1]^I)^3$ is said to be asymmetric $*$ -triangular if $a_i * b_i \leq c_i$ for all $i \in I$.

Definition 4. Let $*$ be a t -norm and I be a set of indices. A function $F : [0, 1]^I \rightarrow [0, 1]$ preserves asymmetric $*$ -triangular triplets if $(F(a), F(b), F(c))$ is an asymmetric $*$ -triangular triplet whenever (a, b, c) so is, where $a, b, c \in [0, 1]^I$.

Proposition 5 ([9]). Let $F : [0, 1]^I \rightarrow [0, 1]$ be a function and $*$ be a t -norm. The following statements are equivalent:

- 1) F preserves $*$ -transitivity of fuzzy binary relations on products;
- 2) F preserves $*$ -transitivity of fuzzy binary relations on sets;
- 3) F preserves asymmetric $*$ -triangular triplets.

II. AGGREGATION OF FUZZY QUASI-METRICS

Definition 6. A fuzzy quasi-metric (in the sense of Kramosil and Michalek) on a nonempty set X is a pair $(M, *)$ such that $*$ is a t -norm and M is a fuzzy set in $X \times X \times [0, +\infty)$ such that for every $x, y, z \in X$ and $t, s > 0$ it verifies:

- $M(x, y, 0) = 0$;
- $M(x, y, t) = M(y, x, t) = 1$ for all $t > 0$ if and only if $x = y$;
- $M(x, y, t) * M(y, z, s) \leq M(x, z, t + s)$;
- $M(x, y, \cdot) : [0, \infty) \rightarrow [0, 1]$ is left-continuous.

If a fuzzy quasi-metric $(M, *)$ also satisfies

- $M(x, y, t) = M(y, x, t)$

for all $x, y \in X$ and all $t \geq 0$ then $(M, *)$ is said to be a fuzzy metric on X .

A fuzzy (quasi-)metric space is a triple $(X, M, *)$ such that X is a nonempty set and $(M, *)$ is a fuzzy (quasi-)metric on X .

Definition 7. A function $F : [0, 1]^I \rightarrow [0, 1]$ is said to be:

- a fuzzy (quasi-)metric aggregation function on products if whenever $*$ is a t -norm and $\{(X_i, M_i, *) : i \in I\}$ is a family of fuzzy (quasi-)metric spaces then $(\widetilde{F} \circ \widetilde{M}, *)$ is a fuzzy (quasi-)metric on $\prod_{i \in I} X_i$ where $\widetilde{M} : (\prod_{i \in I} X_i)^2 \times [0, +\infty) \rightarrow [0, 1]^I$ is given by

$$(\widetilde{M}(x, y, t))_i = M_i(x_i, y_i, t)$$

for every $x, y \in \prod_{i \in I} X_i$ and $t \geq 0$.

If F only satisfies the above condition for a fixed t -norm $*$ then it is said to be an $*$ -fuzzy (quasi-)metric aggregation function on products.

- a fuzzy (quasi-)metric aggregation function on sets if whenever $*$ is a t -norm and $\{(M_i, *) : i \in I\}$ is a family of fuzzy (quasi-)metrics on the same set X then $(F \circ \widetilde{M}, *)$ is a fuzzy (quasi-)metric on X where $\widetilde{M} : X^2 \times [0, +\infty) \rightarrow [0, 1]^I$ is given by

$$(\widetilde{M}(x, y, t))_i = M_i(x, y, t)$$

for every $x, y \in X$ and $t \geq 0$.

If F only satisfies the above condition for a fixed t -norm $*$ then it is said to be an $*$ -fuzzy (quasi-)metric aggregation function on sets.

The next result characterizes those functions which merge an arbitrary family of fuzzy quasi-metrics into a single one. Surprisingly, these functions are exactly those which merge a family of fuzzy metrics into a fuzzy metric, which is not true in the crisp case [7].

Definition 8 (cf. [11]). A function $F : [0, 1]^I \rightarrow [0, 1]$ is said to be $*$ -supmultiplicative for a t -norm $*$ if for all $x, y \in [0, 1]^I$ then

$$F(x) * F(y) \leq F(x *^I y)$$

where $x *^I y \in [0, 1]^I$ is given by $(x *^I y)_i = x_i * y_i$ for all $i \in I$.

Theorem 9 ([9]). Let $F : [0, 1]^I \rightarrow [0, 1]$ be a function and $*$ be a t -norm. The following statements are equivalent:

- 1) F is a $(*)$ -fuzzy quasi-metric aggregation function on products;
- 2) F is a $(*)$ -fuzzy metric aggregation function on products;
- 3) F is isotone, $(*)$ -supmultiplicative, left-continuous, $F(0) = 0$ and $F^{-1}(1) = 1$;
- 4) $F(0) = 0$, F is left-continuous, $F^{-1}(1) = 1$ and F preserves asymmetric $(*)$ -triangular triplets.

Theorem 10 ([9]). Let $F : [0, 1]^I \rightarrow [0, 1]$ be a function and $*$ be a t -norm. The following statements are equivalent:

- 1) F is a $(*)$ -fuzzy quasi-metric aggregation function on sets;
- 2) F is a $(*)$ -fuzzy metric aggregation function on sets;
- 3) F is isotone, $(*)$ -supmultiplicative, left-continuous, $F(0) = 0$ and if $\{x_n : n \in \mathbb{N}\} \subseteq F^{-1}(1)$ there exists $i \in I$ such that $(x_n)_i = 1$ for all $n \in \mathbb{N}$;
- 4) $F(0) = 0$, $F(1) = 1$, F is left-continuous, F preserves asymmetric $(*)$ -triangular triplets and if $\{x_n : n \in \mathbb{N}\} \subseteq F^{-1}(1)$ there exists $i \in I$ such that $(x_n)_i = 1$ for all $n \in \mathbb{N}$.

More results about this topic can be consulted in [9].

ACKNOWLEDGMENT

The two last authors acknowledge financial support from FEDER/Ministerio de Ciencia, Innovación y Universidades-Agencia Estatal de Investigación-Proyecto PGC2018-095709-B-C21.

REFERENCES

- [1] J. Doboš, "Metric preserving functions," Štrobek, Košice, 1998.
- [2] J. Drewniak and U. Dudziak, "Aggregation in classes of fuzzy relations," *Stud. Math.*, vol. 5, pp. 33–43, 2006.
- [3] U. Dudziak, "Preservation of t -norm and t -conorm based properties of fuzzy relations during aggregation process," *Proc. of 8th Conference of the European Society for Fuzzy Logic and Technology (EUSFLAT 2013)*, 2013, pp. 376–383.
- [4] I. Herbut and M. Moszyńska, "On metric products," *Colloq. Math.*, vol. 62, pp. 121–133, 1991.
- [5] J. Martín, G. Mayor, and O. Valero, "On aggregation of normed structures," *Math. Comput. Modelling*, vol. 54, pp. 815–827, 2011.
- [6] G. Mayor and J. Recasens, "Preserving T-transitivity," in *Artificial Intelligence Research and Development*, 2016, vol. 288, pp. 79–87.
- [7] G. Mayor and O. Valero, "Aggregation of asymmetric distances in Computer Science," *Information Sci.*, vol. 180, pp. 803–812, 2010.
- [8] K. Menger, "Statistical metrics," *Proc. Natl. Acad. Sci. USA*, vol. 28, pp. 535–537, 1942.
- [9] T. Pedraza, J. Rodríguez-López and O. Valero, "Aggregation of fuzzy quasi-metrics," *Information Sci.*, to appear.
- [10] J. Recasens, *Indistinguishability operators. Modelling fuzzy equalities and fuzzy equivalence relations*. Springer-Verlag, 2010.
- [11] S. Saminger, R. Mesiar, and U. Bodenhofer, "Domination of aggregation operators and preservation of transitivity," *Internat. J. Uncertain. Fuzziness Knowledge-Based Systems*, vol. 10, pp. 11–35, 2002.

A graded notion of quasi-intents

Manuel Ojeda-Hernández
Dept. Applied Mathematics
University of Mlaga
Mlaga, Spain
manuojeda@uma.es

Inma P. Cabrera
Dept. Applied Mathematics
University of Mlaga
Mlaga, Spain
ipcabrera@uma.es

Pablo Cordero
Dept. Applied Mathematics
University of Mlaga
Mlaga, Spain
pcordero@uma.es

Abstract—The notion of quasi-closed element is extended to fuzzy posets in two stages: First, in the crisp style, in which each element in a given universe either is quasi-closed or it is not. Second, in the graded style by defining a degree to which an element is quasi-closed. We discuss the different possible definitions and compare them with each other. Finally, we show that the most general one has good properties that can be used when we have a complete fuzzy lattice as a frame.

Index Terms—Quasi-closed element, fuzzy poset, closure operator.

I. INTRODUCTION

In this work, the research in [1] is revisited. The goal is to discuss which is the appropriate generalization of the notion of quasi-closedness when working on fuzzy posets. In the classical case, this notion is key to knowledge representation ensuring non-redundancy. However, obtaining an adequate generalization for fuzzy environments that guarantees similar properties, remains an open problem.

A wide variety of generalizations to the fuzzy framework of the notion of implication (and logics for reasoning about them) can be found in the literature, see for example [2]. In [3], the authors include a general framework for these generalizations. All the results on pseudo-closed elements for the fuzzy case have been obtained by using a recursive definition of pseudo-closedness. In the classical case, there exists an equivalent definition based on the notion of quasi-closed element. In this paper, we aim to generalize such notion to the fuzzy framework, which, in the short term, may provide with an alternative definition of pseudo-closed element, as a starting point for a new approach in the study of bases in a fuzzy environment.

As stated in [4], “clearly, in the graded setting, the topics related to non-redundancy and minimality of bases are considerably more involved than in the classic setting and further investigation focused on theory, algorithms, and experiments is needed”.

II. PRELIMINARIES

Throughout this paper, let $\mathbb{L} = (L, \wedge, \vee, \otimes, \rightarrow, 0, 1)$ be a complete residuated lattice. A non-empty set A with a binary $\mathbb{L}$ -relation ρ on A , is said to be a *fuzzy poset* if ρ is a *fuzzy order*, i.e. if ρ is reflexive, antisymmetric and transitive.

The notions of lower (resp. upper) bound and infimum (resp. supremum) used in this work are the ones presented in [5].

Theorem 1: Let $\mathbb{A} = (A, \rho)$ be a fuzzy poset and $X \in L^A$. An element $a \in A$ is supremum (resp. infimum) of X if and only if

$$\rho(a, x) = X^\rho(x) \quad (\text{resp. } \rho(x, a) = X_\rho(x)).$$

It is not difficult to see that, if a supremum (resp. infimum) of X exists, it is unique. We will denote it by $\bigsqcup X$ (resp. $\bigsqcap X$).

Definition 2 ([6]): We say that a fuzzy poset (A, ρ) is a complete fuzzy lattice if every fuzzy subset $X \in L^A$ has supremum and infimum.

We conclude this section with the usual definition of closure operator on a fuzzy poset.

Definition 3: Given a fuzzy poset $\mathbb{A} = (A, \rho)$, a mapping $c: A \rightarrow A$ is said to be a *closure operator* on $\mathbb{A}$ if the following conditions hold:

- 1) $\rho(a, b) \leq \rho(c(a), c(b))$, for all $a, b \in A$ (isotony)
- 2) $\rho(a, c(a)) = 1$, for all $a \in A$ (inflationarity)
- 3) $\rho(c(c(a)), c(a)) = 1$, for all $a \in A$ (idempotency)

Definition 4: Let $c: A \rightarrow A$ be a closure operator on a fuzzy poset (A, ρ) and X be an $\mathbb{L}$ -subset of A . The closure of X wrt c is the $\mathbb{L}$ -set defined by

$$c(X)(a) = \bigvee_{x \in c^{-1}(a)} X(x), \quad \text{for all } a \in A.$$

III. GENERALIZING THE NOTION OF QUASI-CLOSED ELEMENT TO FUZZY POSETS

The aim of this section is to analyse possible generalisations to fuzzy posets of the classical notion of quasi-closed element. We begin by recalling the definition in the case of crisp posets [7]. Throughout this section $\mathbb{A} = (A, \rho)$ is a fuzzy poset and c is a closure operator on $\mathbb{A}$.

There are four equivalent properties that define quasi-closedness in the crisp case. The direct extensions of each one these statements to a fuzzy setting are the following.

- (I) $\rho(a, q) \leq \rho(c(a), q) \vee (c(a) \approx c(q))$, for all $a \in A$.
- (II) $\rho(a, q) \leq \rho(c(a), q) \vee \rho(c(q), c(a))$, for all $a \in A$.
- (III) $\rho(a, q) \otimes \neg \rho(c(q), c(a)) \leq \rho(c(a), q)$, for all $a \in A$.
- (IV) $\rho(a, q) \otimes \neg \rho(q, a) \otimes \rho(c(a), c(q)) \otimes \neg \rho(c(q), c(a)) \leq \rho(c(a), q)$, for all $a \in A$.

These relation among these statements are in the following proposition.

Proposition 5:

- (I) implies (II).
- (II) implies (III).
- (III) implies (IV).

None of the converses hold.

Proposition 6: Given a closure operator c on a (crisp) poset, any quasi-closed element with respect to c satisfies condition (IV).

We adopt the most general of the statements as the definition of quasi-closed element wrt a closure operator.

Definition 7: Given a closure operator c on a fuzzy poset (A, ρ) , an element $q \in A$ is said to be *quasi-closed* (with respect to c) if it satisfies statement (IV).

Theorem 8:

- 1) Every closed element is quasi-closed.
- 2) There exist quasi-closed elements which are not closed.

IV. QUASI-CLOSED ELEMENTS IN GRADED SETTING

In the previous section, an element either was quasi-closed or not. In this section, we define the degree to which an element is quasi-closed. First, to ease the reading of the definitions and properties, we introduce the following notation.

Notation 1: Given a closure operator c on a fuzzy poset (A, ρ) and $q \in A$, we use X_q to denote the $\mathbb{L}$ -set with membership function defined as follows:

$$X_q(a) = \rho(a, q) \otimes \neg \rho(q, a) \otimes \rho(c(a), c(q)) \otimes \neg \rho(c(q), c(a)).$$

With this notation, an element q is quasi-closed iff $X_q(a) \leq \rho(c(a), q)$, for all $a \in A$.

Definition 9: Given a closure operator c on a fuzzy poset (A, ρ) , for any $q \in A$, we define the *degree in which q is quasi-closed* as follows

$$QC(q) = \bigwedge_{x \in A} [X_q(x) \rightarrow \rho(c(x), q)].$$

Theorem 10: Let c be a closure operator on a fuzzy poset (A, ρ) and $q \in A$. Then, $QC(q) = 1$ if and only if q is quasi-closed.

In the classical setting, there is an if-and-only-if condition for a set to be quasi-closed based on an operator that is usually denoted by $^\circ$. We can do an analogous characterization here.

Definition 11: Let c be a closure operator on a complete fuzzy lattice (A, ρ) and $q \in A$. We define the element q° as follows:

$$q^\circ = q \sqcup \bigsqcup c(X_q).$$

This is not a quasi-closed element in general, not even in the crisp case.

Theorem 12: Let c be a closure operator on a complete fuzzy lattice (A, ρ) . Then, $QC(q) = \rho(q^\circ, q)$, for all $q \in A$.

Corollary 13: Let c be a closure operator on a complete fuzzy lattice (A, ρ) . An element $q \in A$ is quasi-closed wrt c if and only if $q^\circ = q$.

V. CONCLUSIONS AND FURTHER WORK

We have presented a fuzzy definition of quasi-closed elements in the frame of fuzzy posets and checked its properties. On the other hand, we have extended the definition to graded setting and we have proved that it extends the classical results that are necessary for its effective use in the search for bases of implications or if-then rules in the fuzzy frame.

In a next step, we will study aspects related to the computability of quasi-closed elements looking for necessary and sufficient conditions to ensure that they can be calculated efficiently. As further work, we will generalize the notion of pseudo-closed element and compare the definition obtained with the recursive one proposed in [4]. In addition, we will study whether the bases of implications with pseudo-closed premises have the desired properties of completeness, non-redundancy and minimality. We will also consider whether this work can be extended to the multi-adjoint concept lattice framework [8].

ACKNOWLEDGMENT

This work has been partially supported by the FPU19/01467 internship (Spanish MCIU) and the Spanish research projects with reference TIN2017-89023-P (MCIU/AEI/FEDER), PGC2018-095869-B-I00 and UMA2018FEDERJA001 (JA/UMA/FEDER).

REFERENCES

- [1] M. Ojeda-Hernández, I. P. Cabrera, and P. Cordero, "Quasi-closed elements in fuzzy posets," *Journal of Computational and Applied Mathematics*, vol. In Press, pp. 113–390, 2021.
- [2] L. Jezková, P. Cordero, and M. Enciso, "Fuzzy functional dependencies: A comparative survey," *Fuzzy Sets Syst.*, vol. 317, pp. 88–120, 2017.
- [3] P. Cordero, M. Enciso, A. Mora, and V. Vychodil, "Parameterized simplification logic: Reasoning with implications and classes of closure operators," *Int. J. General Systems*, vol. In press, pp. 1–18, 2020.
- [4] V. Vychodil, "Computing sets of graded attribute implications with witnessed non-redundancy," *Information Sciences*, vol. 351, pp. 90 – 100, 2016. [Online]. Available: <http://www.sciencedirect.com/science/article/pii/S0020025516301451>
- [5] J. Konecny and M. Krupka, "Complete relations on fuzzy complete lattices," *Fuzzy Sets and Systems*, vol. 320, pp. 64 – 80, 2017. [Online]. Available: <http://www.sciencedirect.com/science/article/pii/S0165011416302718>
- [6] R. Bělohlávek, "Concept lattices and order in fuzzy logic," *Annals of Pure and Applied Logic*, vol. 128, no. 1, pp. 277 – 298, 2004.
- [7] A. Day, "The lattice theory of functional dependencies and normal decompositions," *Int. J. Algebra and Computation*, vol. 2, pp. 409–431, 1992.
- [8] M. E. Cornejo, J. Medina, and E. Ramírez-Poussa, "Attribute and size reduction mechanisms in multi-adjoint concept lattices," *Journal of Computational and Applied Mathematics*, vol. 318, pp. 388 – 402, 2017.

Community Detection Problem Based on Polarization Measures: An Application to Twitter. The COVID-19 Case in Spain. A Review

Inmaculada Gutiérrez^{*†}, Juan Antonio Guevara^{*†}, Daniel Gómez^{*‡}, Javier Castro^{*¶} and Rosa Espínola^{*||}

^{*}Faculty of Statistics

Complutense University Puerta de Hierro, s/n, 28040 Madrid, Spain

[†] Email: inmaguti@ucm.es and juanguelv@ucm.es

[§]Instituto de Evaluación Sanitaria, Complutense University, 28040 Madrid, Spain

Abstract—In this paper, we address one of the most important topics in the field of Social Networks Analysis: the community detection problem with additional information. That additional information is modeled by a fuzzy measure that represents the risk of polarization. Adding this type of information to the community detection problem makes it more realistic, as a community is more likely to be defined if the corresponding elements are willing to maintain a peaceful dialogue. Hereafter, we work in a real case obtained from Twitter, concerning the political position against the Spanish government. We analyze how the partitions obtained change when some additional information related to how polarized that society is, is added to the problem.

Index Terms—networks; community detection; extended fuzzy graphs; polarization; fuzzy sets; ordinal variation

I. INTRODUCTION

In this document, we review our paper [1] in which we addressed one of the hottest problems in the Social Networks Analysis (SNA) field: the community detection problem. We go beyond the “classic” approach of this problem, based on the crisp connections between the elements which are defined by the edges of the graph. Our main contribution is the incorporation of some additional information modeled by a fuzzy measure.

In particular, our background is related to Polarization. In broader terms, Polarization can be understood as the split of a given society into two different and opposite groups with significant size along an attitudinal axis [2]. In this work, we recall the concept of Polarization based on fuzzy sets developed in [3], where Guevara et al. introduced a Polarization measure based on the fuzzy set approach, the JDJ_{pol} . It uses the membership degree of each individual to the poles and some aggregation operators to measure the risk of polarization of a group. We consider as an important matter to include the concept of Polarization into the field of SNA and to community detection problems due that Polarization and communication are strongly related. The way in which a community is detected in a social context should fit reality not only based on the interactions between their nodes but in its attitudinal or ideology coherence as well.

On the basis of the JDJ_{pol} , we work with (non) polarization fuzzy measures [4], which allow us to measure the ability for peaceful dialog of a society.

Then, we introduce the non polarization extended fuzzy graph (non-polarization EFG). It is characterized on the basis of the extended fuzzy graph (EFG) [5]. Given a crisp graph $G = (V, E)$ and a fuzzy measure $\mu : 2^V \rightarrow [0, 1]$ defined on the set of nodes, the triplet $\tilde{G} = (V, E, \mu)$ is said to be an extended fuzzy graph. As it is pointed in [6], fuzzy/capacity measures are fundamental in modeling dependencies among the inputs. Then, with the combination of the ability of the graph to model connections between elements, and the ability of the fuzzy measures to handle the capacity related to any set of elements, we can represent situations in which more than two nodes are implied, independent of the way they are connected through the graph by using an EFG.

Finally, we work in a specific application of the non-polarization EFG: the community detection problem. Once we develop a methodology to find groups in a graph paying attention to how polarized the society is, we work in a real case obtained from Twitter and related to the hottest topic in the last 2020: the Covid-19 pandemic. In particular, we analyze the position of people against the Spanish government during this period of crisis.

II. POLARIZATION FUZZY MEASURES AND POLARIZATION EXTENDED FUZZY GRAPH

Our first goal is the characterization of a fuzzy measure obtained from the JDJ_{pol}, μ_{P-} . We emphasize on the fact the μ_{P-} can be re-formulated as a summation which involves the elements of a matrix P^- which is symmetric, non-negative, 1-normalized and whose main diagonal is null. Because of the interpretation of JDJ_{pol}, P_{ij}^- represents the risk of conflict concerning the elements i and j . So that, μ_{P-} represents the capacity of the elements to argue, to trigger conflict and arguments. Hence, it is a recommended model to properly represent the discrepancy or distance between individuals.

Because of the interpretation of the measure JDJ_{pol} , its negation, $\widetilde{JDJ}_{pol}$ can be understood as the minimum risk of polarization for a given population or community. Then, from

$\widetilde{JD}J_{pol}$, we define the matrix P^+ , non negative, symmetric, 1-normalized and with main diagonal null. It can be somehow understood as an “affinity” matrix [5], from which we can define a measure which represents the capacity of the elements of a set to peacefully dialogue without risk of Polarization, μ_{P^+} . Because of the properties of P^+ , we can affirm that μ_{P^+} is a 2-additive non-polarization fuzzy measure [7].

Then, we define a new representation model: the non polarization extended fuzzy graph. It is a crisp graph together with a non-polarization fuzzy measure, the triplet $G = (V, E, \mu_{P^+})$.

III. COMMUNITY DETECTION PROBLEM BASED ON POLARIZATION MEASURES

We approach the community detection problem based on fuzzy measures including this information inspired by the idea developed in the Additional Louvain algorithm (see in [7]), based on the Louvain algorithm [8]. The key point is to distinguish two different roles within the input parameters: one of them, to establish the neighbor relations, and the other, to calculate the variation of the modularity. The first role will be played by the adjacency matrix of the graph, A , so that only those nodes that are connected in G can be in the same group. On the other hand, we suggest to consider a combination of the two components of the non-polarization extended fuzzy graph $\tilde{G}$ as basis to calculate the variation of modularity, in order to incorporate the additional information. Then, having a crisp graph, the two membership functions, and a grouping, an overlapping and a negation operators, it can be obtained a non-polarization EFG, $\tilde{G} = (V, E, \mu_P)$, where $\mu_P = \mu_+$. In order to simplify the management of the synergies between elements, it is characterized the weighted graph associated with a fuzzy measure, G_{μ_P} , particularly, the one associated with μ_P . Being ξ an aggregation operator, $Sh_i(\mu_P)$ and $Sh_i^j(\mu_P)$ the Shapley index of i when it is in a coalition with all the elements of V and $V \setminus \{j\}$ respectively, the graph G_{μ_P} is that whose adjacency matrix is F , where

$$F_{ij} = \xi \left(Sh_i(\mu_P) - Sh_i^j(\mu_P), Sh_j(\mu_P) - Sh_j^i(\mu_P) \right) \quad (1)$$

In our specific proposal, we suggest summarizing the non-polarization fuzzy measure μ_P into the matrix F , with adjacency of its associated weighted graph defining the Polarization Louvain algorithm to detect communities. We combine the matrices A and F by means of a linear combination ($\theta(A, F) = \gamma A + (1 - \gamma)F$) using the parameter γ to assign a weight or importance to each component of the $\tilde{G}$. Note that when $\gamma = 1$ the additional information is not considered.

IV. A REAL CASE: THE IMPACT OF THE COVID-19 PANDEMIC IN THE ORGANIZATION OF THE PEOPLE

The nodes and theirs relations considered in this work have been obtained from the social network Twitter, particularly from some posts recorded along the state of alarm imposed by the central government in Spain. All data downloaded relate to the COVID-19 pandemic. The experiment design and all

the process are detailed in ¹. A random sample of tweets were labeled by an expert in order to label them as supporters or detractors towards the Spanish Government. Then, machine learning algorithms were applied to generalize the knowledge to all tweets. Support Vector Machines showed the best performance so we used that probability of being a supporter or detractor as membership degree values to compute JDJ_{pol} . Due that we labeled tweets but not users we computed the average probability of all messages posted by a given user in order to obtain the probability to be a supporter and detractor of that user. We apply the Polarization Louvain algorithm to find communities in the non-polarization extended fuzzy graph $\tilde{G} = (V, E, \mu_P)$. We vary the parameter γ which allow us to control the importance of that extra knowledge (Polarization values between two individuals). The partition which shows the best modularity value is considered as optimal.

To compare the Polarization Louvain algorithm with the traditional Louvain algorithm we show an example of how two pairs of nodes which should belong to the same communities, respectively, are split into four different communities with the Louvain algorithm. On one hand, we have nodes “38” and “115”, both left-wing political parties that teamed back in march 2019. On the other hand, we have nodes “76”, a right-wing political party, and “203”, a member of this political group. After applying the Polarization Louvain algorithm, those pairs are clustered into the same communities (see the original paper [1] for illustrations and more details).

ACKNOWLEDGMENT

This research has been partially supported by the Government of Spain, Grant Plan Nacional de I+D+i, MTM2015-70550-P, PGC2018096509-B-I00, TIN2015-66471-P and PID2019-106254RB-I00, and the CT17/17-CT18/17.

REFERENCES

- [1] Gutiérrez, I.; Guevara, J.A.; Gómez, D.; Castro, J.; Espínola, R. Community Detection Problem Based on Polarization Measures: An Application to Twitter: The COVID-19 Case in Spain. *Mathematics*. **2021**, *9*, 443.
- [2] Esteban, J.M.; Ray, D. On the measurement of polarization. *Econom. J. Econom. Soc.* **1994**, *62*, 819–851.
- [3] Guevara, J.A.; Gómez, D.; Robles, J.M.; Montero, J. Measuring Polarization: A Fuzzy Set Theoretical Approach. In Proceedings of the International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems, Lisbon, Portugal, 15–19 June 2020; Springer: Berlin, Germany, **2020**; pp. 510–522.
- [4] Sugeno, M. Fuzzy measures and fuzzy integrals: A survey. *Fuzzy Automata and Decision Processes*. **1977**, *78*, 89–102.
- [5] Gutiérrez, I.; Gómez, D.; Castro, J.; Espínola, R. A New Community Detection Algorithm Based on Fuzzy Measures. In *Advances in Intelligent Systems and Computing Series, Proceedings of the Intelligent and Fuzzy Techniques in Big Data Analytics and Decision Making INFUS 2019*; Springer: Cham, Switzerland, **2020**; 1029, pp. 133–140.
- [6] Beliakov, G. On random generation of supermodular capacities. *IEEE Trans. Fuzzy Syst.* **2020**.
- [7] Gutiérrez, I.; Gómez, D.; Castro, J.; Espínola, R. Fuzzy Measures: A solution to deal with community detection problems for networks with additional information. *J. Intell. Fuzzy Syst.* **2020**, *39*, 6217–6230.
- [8] Blondel, V.; Guillaume, J.; Lambiotte, R.; Lefevre, E. Fast unfolding of communities in large networks. *J. Stat.-Mech. Theory Exp.* **2008**, *10*.

¹<https://github.com/inmaggp/Community-Detection-Problem-Based-on-Polarization-Measures.-An-application-to-Twitter-the-COVID-19->

Rough-Set-Driven Approach for Attribute Reduction in Fuzzy Formal Concept Analysis

1st M. José Benítez-Caballero

Department of Mathematics
University of Cádiz
Cádiz, Spain
mariajose.benitez@uca.es

2nd Jesús Medina

Department of Mathematics
University of Cádiz
Cádiz, Spain
jesus.medina@uca.es

3rd Eloísa Ramírez-Poussa

Department of Mathematics
University of Cádiz
Cádiz, Spain
eloisa.medina@uca.es

4th Dominik Ślęzak

Institute of Informatics
University of Warsaw
Warsaw, Poland
slezak@mimuw.edu.pl

Abstract—This work concerns the research recently published in M. J. Benítez-Caballero, J. Medina, E. Ramírez-Poussa and D. Ślęzak, “Rough-Set-Driven Approach for Attribute Reduction in Fuzzy Formal Concept Analysis”, *Fuzzy Sets and Systems*, vol. 391, pp. 117–138, 2020.

Index Terms—fuzzy sets, attribute reduction, reduct, formal concept analysis, rough set theory

I. INTRODUCTION

The study of the knowledge stored in databases is one of the most important goals in several research fields. Formal Concept Analysis (FCA) [10] and Rough Set Theory (RST) [8] are two widely studied mathematical theories, devoted to obtain information from relational databases that contain uncertainty.

The reduction of size of database is a widely study issue in both theories, separately. In addition, in the literature, several papers can be found that establish the existing connections between these two mathematical tools, considering the classical framework [3], [5], [7], [9].

This paper introduces a novel mechanism to reduce the set of attributes in the fuzzy general framework of multi-adjoint concept lattices, considering the RST philosophy with tolerance relations. This study extends the one introduced in [3] in a fuzzy case, shown that the same properties are not satisfied. Besides that, the proposed mechanism has been enriched with other interesting properties, showing that the new procedure also keeps important features. One of these properties is that the reduction is directly applied to the context and the whole concept lattice is not needed to be computed. Moreover, the main structure, based on the join-irreducible elements, is preserved including no new join-irreducible element after the reduction procedure. The notions and results obtained considering this framework is deeply studied in [4].

II. REDUCTION IN MULTI-ADJOINT CONCEPT LATTICES

In this section, we will present the proposed reduction mechanism to multi-adjoint concept lattices, which can also be applied to any other fuzzy FCA framework. For that, the RST reduction philosophy and a family of tolerance relations

will be taken into account. In the following, we explain step by step how this reduction procedure is carried out.

Reduction in Multi-adjoint Concept Lattices from RST with tolerance relations

- S1.** The first step is to consider the associated information system (B, A) , with the mappings $\bar{a}: B \rightarrow V_a$, from a given multi-adjoint formal context (A, B, R, σ) and a frame $(L_1, L_2, P, \&_1, \dots, \&_n)$.
- S2.** From this context the sets V_a will be the poset P and the mappings $\bar{a}: B \rightarrow P$ will be defined as $\bar{a}(x) = R(a, x)$, for each $a \in A$ and $x \in B$.
- S3.** In this new environment, we consider the tolerance relations for each attribute, which are used to build the unidimensional $\mathcal{E}$ -discernibility function of the associated information system.
- S4.** We compute the $\mathcal{E}$ -information reducts $D_1, \dots, D_n$.
- S5.** Considering these reducts we reduce the original fuzzy context, obtaining the reduced ones $(D_1, B, R_{|D_1 \times B}), \dots, (D_n, B, R_{|D_n \times B})$.
- S6.** Finally, we build the concept lattices from the reduced contexts obtaining significant reductions in the context as well as in the size of the original lattice, preserving the discernibility among the objects.

From now on, a multi-adjoint frame $(L_1, L_2, P, \&_1, \dots, \&_n)$ and a multi-adjoint context (A, B, R, σ) will be fixed. First of all, different properties relating the concept-forming operators in the original context and in the reduced one will be introduced in the following proposition, which will be needed later on.

Proposition 1: Given a subset $D \subseteq A$, for any concept $\langle g, f \rangle$ obtained from the multi-adjoint context, the following statements hold:

- 1) $g^{\uparrow\downarrow} \preceq_2 g^{\uparrow D \downarrow D}$
- 2) $g^{\uparrow\downarrow\uparrow D} = g^{\uparrow D}$
- 3) $f^{\downarrow D \uparrow\downarrow} = f^{\downarrow D}$

The following proposition asserts that if a group of concepts in the original context are the same concept in the reduced one, then they have the structure of a join-semilattice.

Proposition 2: Let $D \subseteq A$ be a subset of attributes. The set $R_E = \{(\langle g_1, f_1 \rangle, \langle g_2, f_2 \rangle) \mid \langle g_1, f_1 \rangle, \langle g_2, f_2 \rangle \in$

$\mathcal{M}(A, B, R, \sigma), g_1^{\uparrow D \downarrow D} = g_2^{\uparrow D \downarrow D}$ is an equivalence relation and every class $[\langle g, f \rangle]_D$ of $\mathcal{M}(A, B, R, \sigma)/R_E$ is a join-semilattice with maximum element $\langle g^{\uparrow D \downarrow D}, g^{\uparrow D \downarrow D} \rangle$.

We can think that it is natural that properties presented in [3] will not be satisfied when fuzzy sets are considered. Concerning the structural properties, one of the most important feature is related to the join-irreducible elements. As it was proven in Theorem 3.5 of [3] for the classical framework of FCA, no new join-irreducible element appears after the reduction process. This main structural property is also preserved in this fuzzy framework, as the following result shows.

Theorem 3: Given an $\mathcal{E}$ -information reduct D in the corresponding context information system (B, A) , if the pair $\langle g, f \rangle$ is a join-irreducible concept in the reduced concept lattice built using the reduct D , then there exists an object $b \in B$ and a truth value $x \in L_2$ such that $\phi_{b,x}^{\uparrow D} = f$ and $\phi_{b,x}$ also generates a join-irreducible concept in the original concept lattice.

III. RELATED METHODOLOGIES

The attribute reduction mechanism based on tolerance relations proposed in this paper is different from the ones given in diverse papers. In this section, we will focus our attention on two general reduction procedures given in [1] and [6], since the rest of procedures are given in a more restrictive framework or are based on them.

The philosophy of the reduction presented in [1], considering similarities, is very different from the proposed one in this paper. Our reduction is directly applied to the context. This fact is very important since, for example, it has a direct impact on the construction of the attribute implications, such as, reducing the number of implications or creating equivalences among them. Meanwhile, the factorization presented by Bělohlávek reduces the concept lattice and it has no impact in the context. Moreover, the whole concept lattice must be computed before calculating the factorization, which is not necessary for the mechanism we propose. Furthermore, the reduction given in [1] provides a covering instead of a partition, which is the clustering we obtain from the $\frac{1}{2}$ -information reducts.

On the other hand, the reduction presented in [6] considers block relations, where the rows of this relation are intents and the columns are extents. The methodology proposed in this paper is different, since the tolerance relations are defined on the set of values of the attributes, and they are independent of the relation of the context. With our procedure, we try to group the similar attributes whose separated consideration does not provide relevant information and so, the main knowledge of the database is preserved after the reduction. Although both mechanisms are different, they are compatible. As a consequence, we can apply both procedures to the same database in order to obtain a bigger reduction, embedding both philosophies.

IV. CONCLUSIONS AND FUTURE WORK

In this paper, by means of different results and examples, we have introduced a mechanism to reduce attributes in fuzzy FCA, considering the reduction procedure with tolerance

relations introduced in RST [2]. This new method to reduce attributes provides a significant reduction of the original concept lattice. Some interesting properties of the new procedure have been presented throughout the paper. The most important one shows that the structure of the original concept lattice is partially preserved considering this new mechanism, that is, no new join-irreducible elements appear after carrying out the reduction procedure.

Moreover, the paper finishes exposing a comparison among this new mechanism and other interesting fuzzy reduction methods.

In the future, we are interested in establishing a comparison between the attribute reduction process given in multi-adjoint concept lattices and other multi-adjoint frameworks, as multi-adjoint property oriented concept lattice and multi-adjoint object oriented concept lattices. We also want to apply this study to real examples.

REFERENCES

- [1] R. Bělohlávek. Similarity relations in concept lattices. *Journal of Logic and Computation*, 10(6):823–845, 2000.
- [2] M. J. Benítez-Caballero, J. Medina, E. Ramírez-Poussa, and D. Ślęzak. Bireducts with tolerance relations. *Information Sciences*, 435:26 – 39, 2018.
- [3] M. J. Benítez-Caballero, J. Medina, E. Ramírez-Poussa, and D. Ślęzak. A computational procedure for variable selection preserving different initial conditions. *International Journal of Computer Mathematics*, 2019.
- [4] M. J. Benítez-Caballero, J. Medina, E. Ramírez-Poussa, and D. Ślęzak. Rough-set-driven approach for attribute reduction in fuzzy formal concept analysis. *Fuzzy Sets and Systems*, 2019.
- [5] J. Chen, J. Li, Y. Lin, G. Lin, and Z. Ma. Relations of reduction between covering generalized rough sets and concept lattices. *Information Sciences*, 304:16–27, 2015.
- [6] J. Konecny and M. Krupka. Complete relations on fuzzy complete lattices. *Fuzzy Sets and Systems*, 320:64 – 80, 2017. Theme Logic and Algebra.
- [7] J. Medina. Relating attribute reduction in formal, object-oriented and property-oriented concept lattices. *Computers & Mathematics with Applications*, 64(6):1992–2002, 2012.
- [8] Z. Pawlak. Rough sets. *International Journal of Computer and Information Science*, 11:341–356, 1982.
- [9] L. Wei and J.-J. Qi. Relation between concept lattice reduction and rough set reduction. *Knowledge-Based Systems*, 23(8):934–938, 2010.
- [10] R. Wille. Restructuring lattice theory: an approach based on hierarchies of concepts. In I. Rival, editor, *Ordered Sets*, pages 445–470. Reidel, 1982.

Measuring Consistency of Fuzzy Logic Theories

Nicolás Madrid

Departamento de Matemática Aplicada
Universidad de Málaga
Málaga, Spain
Email: nicolas.madrid@uma.es

Manuel Ojeda-Aciego

Departamento de Matemática Aplicada
Universidad de Málaga
Málaga, Spain
Email: aciego@uma.es

Abstract—Fuzzy logic has shown to be a suitable framework to handle contradictions in which, unsurprisingly, the notion of inconsistency can be defined in different ways. This paper starts with a short survey of different ways to define the notion of inconsistency in fuzzy logic systems. As a result, we provide a first notion of inconsistency by means of the absence of models. Subsequently, we define two measures of consistency that belong purely to the fuzzy paradigm; in the sense that both measures coincide with the crisp notion of consistency when the set of truth values is $\{0, 1\}$. Accordingly, we can state that the two provided measures of consistency are notions of consistency based on degrees, bringing back the spirit of fuzzy logic into the notion of consistency.

Based on the paper *A measure of consistency for fuzzy logic theories*, to appear in the journal *Mathematical Methods in the Applied Sciences*, 2021.

I. INTRODUCTION

Since its introduction, fuzzy sets and fuzzy logic have shown to be an interesting research topic. One can find lots of papers ranging from the development of algebraic theories of fuzzy structures or the underlying mathematics of fuzzy logic, to fuzzy modelling or automated control in terms of sets of fuzzy rules. From the theoretical standpoint; in [1] the notion of relational Galois connection is extended to be applied between transitive fuzzy directed graphs in a framework in which the components of the connection are crisp relations satisfying certain reasonable properties; from the practical standpoint, in [7] it is shown how a control application can leverage (even) from a set of inconsistent rules; and in [2] we can see a fuzzy logic-based mathematical model of a sequence of earthquakes using tools from fuzzy reasoning.

Being the fuzzy realm a matter of degrees, a number of papers have focused on measuring the degree of inconsistency of a set of fuzzy rules, and a number of different inconsistency indices have been introduced. For instance [3] introduces the so-called knowledge-based consistency index for deriving priorities from fuzzy pairwise comparison matrices in multiple-criteria decision-making problems; other approaches introduce means for both measuring and repairing inconsistency, for example [8] presents a family of measures aimed at determining the amount of inconsistency in knowledge bases with graded truth and considers minimal adjustments in the truth-degrees of the propositions necessary to make the knowledge-base to be consistent within a given frame (in that case the Łukasiewicz semantics); last but not least [4] deals with the

definition of measures of inconsistency in the residuated-logic-programming paradigm under the fuzzy answer set semantics and provides a soft mechanism to control the amount of information inferred, thus, controlling the inconsistencies by modifying slightly the truth values of some rules. The number of possible measures of inconsistency that can be found in the literature somehow suggests the existence of a problem with inconsistency in a fuzzy setting, namely, its definition: there is not a consensus on how to interpret inconsistency in a fuzzy system.

In this paper we briefly survey the main properties and equivalent characterisations of inconsistency in classical (crisp) logic and then, we focus on, under of point of view, the more natural way to define inconsistency in a logic theory, namely: the absence of models. This consideration as definition of inconsistency keeps some of the most important properties of inconsistency in the fuzzy paradigm, e.g., explosive reasoning. However, we also lose an important issue, we lose degrees; which is the soul of fuzzy logic. For such a reason, we propose a generalization of consistence by means of two measures of consistency. Specifically, we define two measures of consistency that belong purely to the fuzzy paradigm. In other words, both measures coincide with the crisp notion of consistency when the set of truth values is $\{0, 1\}$. Moreover, we provide a set of properties for both measures of consistency in order to motivate the use of them to represent the consistency of fuzzy logic theories.

II. MEASURES OF CONSISTENCY FOR FUZZY LOGIC THEORIES.

In this section we follow a common procedure in the definition of measures of (in-)consistency in crisp logic: given a logic theory Γ , we consider subsets of consistent formulas contained in Γ . At this point, in crisp logic we can measure the inconsistency by considering the ratio or the absolute number of removed formulas. Interestingly enough, in a fuzzy environment, we can proceed differently: for instance, we can measure the consistency Mc of the removed formulas with respect to the remaining ones.

Definition 1: Let Γ be a fuzzy logic theory defined on a residuated lattice $(L, \leq, *, \rightarrow)$ and consider $\alpha \in L$. A formula ψ is said to be α -feasible w.r.t. Γ if $\Gamma \models \psi \rightarrow \alpha$. If $\Gamma = \emptyset$

The degree of consistency of ψ with respect to Γ is:

$$Mc(\psi, \Gamma) = \min\{\alpha \mid \psi \text{ is } \alpha\text{-feasible w.r.t } \Gamma\}$$

In order to properly understand the rationale in the following definition, let us consider an arbitrary fuzzy logic theory of three formulas $\Gamma = \{\psi_1, \psi_2, \psi_3\}$. Assume that Γ is inconsistent and that ψ_1 is non-contradictory. The question is how consistent are formulas ψ_2 and ψ_3 in the logic theory $\Gamma^* = \{\psi_1\}$. One could think about measuring separately the consistence for both formulas and then aggregate them, but that is not possible because $\{\psi_1, \psi_2\}$ and $\{\psi_1, \psi_3\}$ could be consistent and, then, both measures would be 1. Therefore, the only reasonable option is to combine ψ_2 and ψ_3 into one formula. Note, that the consistency of $\{\psi_1, \psi_2, \psi_3\}$ is given by assuming on the one hand ψ_1 and, on the other hand, both ψ_2 and ψ_3 at the same time; the latter means that we are assuming $\psi_2 \wedge \psi_3$. Therefore, the consistency generated by ψ_2 and ψ_3 in the logic theory $\Gamma^* = \{\psi_1\}$ is $Mc(\psi_2 \wedge \psi_3, \Gamma^*)$.

Definition 2: Let Γ be a fuzzy logic theory defined on a residuated lattice $(L, \leq, *, \rightarrow)$, then we define the measure of consistency $Mc^*(\Gamma)$ as

$$\sup \left\{ Mc \left(\bigwedge_{\psi_i \in \Gamma \setminus \Gamma^*} \psi_i, \Gamma^* \right) \mid \Gamma^* \subseteq \Gamma \text{ is consistent} \right\}.$$

At first sight, the reader may think that considering all the set of combinations of consistent sub-theories of a fuzzy logic theory Γ may be impractical, however, the following result shows that only one consistent subtheory must be considered to compute the measure Mc^* , namely, the empty theory.

Theorem 1: Let Γ be a fuzzy logic theory defined on a residuated lattice $(L, \leq, *, \rightarrow)$, then:

$$Mc^*(\Gamma) = Mc \left(\bigwedge_{\psi_i \in \Gamma} \psi_i, \emptyset \right).$$

As a direct consequence of the previous theorem, we have the following corollary that shows that Mc^* satisfies those properties of a measure of consistency (i.e., the opposite properties of a measure of inconsistency).

Corollary 1: Let Γ and Γ' be fuzzy logic theories defined on a residuated lattice $(L, \leq, *, \rightarrow)$, then:

- a) $Mc^*(\Gamma) \geq Mc^*(\Gamma \cup \Gamma')$;
- b) If Γ is consistent then, $Mc^*(\Gamma) = 1$;
- c) If $Mc^*(\Gamma) \neq 1$ then, Γ is inconsistent;
- d) If L is finite and totally ordered, then $Mc^*(\Gamma) = 1$ implies Γ is consistent.

The measure of consistency Mc^* is related to the k -models which, in turn, are related to the so-called x -consistency [9] and α -cuts models [6]. The underlying idea in the k -models is to guarantee the satisfiability of formulas in at least truth-degree $k \in L$, and it is given in the following definition.

Definition 3: Let Γ be a fuzzy logic theory defined on a residuated lattice $(L, \leq, *, \rightarrow)$ and consider $k \in L$. We say that an interpretation M is a k -model of Γ if $M(\psi) \geq k$ for all $\psi \in \Gamma$.

The k -models were introduced in the context of Fuzzy Logic Programming aiming at providing “partial” models to a given inconsistent logic program (i.e., fuzzy logic theory). Later, it

was proved that the existence of models is guaranteed by very general requirements in Fuzzy Logic Programming [5] and, then, k -models faded away. However, in the general context we are working on in this approach, the existence of models cannot be guaranteed easily and k -models may be valuable here. The following result relates the measure Mc^* with k -models.

Theorem 2: Let Γ be a fuzzy logic theory defined on a residuated lattice $(L, \leq, *, \rightarrow)$. If $Mc^*(\Gamma) = \alpha$ then, there is not β -model of Γ with $\beta > \alpha$.

III. CONCLUSIONS AND FUTURE WORK

We have presented two different measures of consistency. The first one measures how much compatible a formula is with respect to a given theory in the sense *the closer to 0, the more inconsistent; and the closer to 1, the more consistent*. The second measure determines a degree of consistency of a logic theory by means of consistent subtheories. Both definitions coincide with the standard notion of consistency when we restrict to crisp logic, and both definitions satisfy convenient properties in order to be considered measures of consistency.

There are two main lines of future research. On the one hand it is convenient to keep digging up some measures of inconsistency in fuzzy paradigms. To have a notion of inconsistency based on degrees (as the ones proposed in this paper) may allow to incorporate a paraconsistent reasoning into inconsistent fuzzy logic theories without leaving the fuzzy paradigm aside. On the other hand, it is interesting to find out an application of the measures of consistence. For instance, we think they can be used to deal with contradictions in databases obtained from fails or system errors.

Acknowledgment

Partially supported by projects PGC2018-095869-B-I00 and UMA2018-FEDERJA-001.

REFERENCES

- [1] I. Cabrera, P. Cordero, E. Muñoz-Velasco, M. Ojeda-Aciego, and B. De Baets. Relational galois connections between transitive fuzzy digraphs. *Mathematical Methods in the Applied Sciences*, 43(9):5673–5680, 2020.
- [2] A. Konguetsof, N. Mylonas, and B. Papadopoulos. Fuzzy reasoning in the investigation of seismic behavior. *Mathematical Methods in the Applied Sciences*, 43(13):7747–7757, 2020.
- [3] S. Kubler, W. Derigent, A. Voisin, J. Robert, Y. L. Traon, and E. Herrera-Viedma. Measuring inconsistency and deriving priorities from fuzzy pairwise comparison matrices using the knowledge-based consistency index. *Knowl. Based Syst.*, 162:147–160, 2018.
- [4] N. Madrid and M. Ojeda-Aciego. Measuring inconsistency in fuzzy answer set semantics. *IEEE Trans. Fuzzy Syst.*, 19(4):605–622, 2011.
- [5] N. Madrid and M. Ojeda-Aciego. On the existence and unicity of stable models in normal residuated logic programs. *Intl J of Computer Mathematics*, 89(3):310–324, 2012.
- [6] J. Močkoř. α -cuts and models of fuzzy logic. *International Journal of General Systems*, 42(1):67–78, 2013.
- [7] N. R. Pal and K. Pal. Handling of inconsistent rules with an extended model of fuzzy reasoning. *J. Intell. Fuzzy Syst.*, 7(1):55–73, 1999.
- [8] D. Picado-Muñoz. Measuring and repairing inconsistency in knowledge bases with graded truth. *Fuzzy Sets Syst.*, 197:108–122, 2012.
- [9] D. Van Nieuwenborgh, M. De Cock, and D. Vermeir. An introduction to fuzzy answer set programming. *Annals of Mathematics and Artificial Intelligence*, 50(3):363–388, 2007.

Detecting equivalence classes being sublattices from an attribute reduction in FCA

Roberto G. Aragón*
Department of Mathematics
University of Cádiz
Cádiz, Spain
roberto.aragon@uca.es

Jesús Medina
Department of Mathematics
University of Cádiz
Cádiz, Spain
jesus.medina@uca.es

Eloísa Ramírez-Poussa
Department of Mathematics
University of Cádiz
Cádiz, Spain
eloisa.ramirez@uca.es

Abstract—This work concerns the research recently published in R. G. Aragón, J. Medina, and E. Ramírez-Poussa. *Identifying Non-Sublattice Equivalence Classes Induced by an Attribute Reduction in FCA. Mathematics*, 9(5), 2021.

Index Terms—Formal Concept Analysis, equivalence relations, attribute reduction

I. INTRODUCTION

An appealing goal in different frameworks is detecting redundant or irrelevant variables (attributes) in data sets, such as in Formal Concept Analysis (FCA) in which the removal of redundant data becomes essential. However, the elimination of such variables may have some impact in the concept lattice, which is closely related to the algebraic structure of the obtained quotient set and their classes [4].

In [2], [3], local congruences were introduced as equivalence relations defined on lattices whose equivalence classes are sublattices of the original lattice. Furthermore, local congruences were intended to complement the attribute reductions of formal contexts in order to ensure that the equivalence classes, $[C]_D$, be sublattices of the original concept lattice. If its infimum $C_m = \bigwedge_{C_i \in [C]_D} C_i$ belongs to the equivalence class, we can assert that the class is already a sublattice, since a join-semilattice with a minimum element is a lattice. As a consequence, in this case, the application of a local congruence, as a complementary mechanism to attribute reduction, does not provide any modification in this particular class. Therefore, it is significant to characterize the required conditions under which these cases arise.

II. CHARACTERIZING EQUIVALENCE CLASSES

This research line was initiated in [2] and was continued in [1]. Namely, we determined in this last paper a sufficient condition to ensure that the equivalence class of C_m is generated by an attribute-concept and presented the following enhanced version of the characterization of the infimum of elements belonging to a non-singleton classes:

* Corresponding author

Partially supported by the 2014–2020 ERDF Operational Programme in collaboration with the State Research Agency (AEI) in project PID2019-108991GB-I00, and with the Department of Economy, Knowledge, Business and University of the Regional Government of Andalusia in project FEDER-UCA18-108612, and by the European Cooperation in Science & Technology (COST) Action CA17124.

Theorem 1: Given a context (A, B, R) , a subset of attributes $D \subseteq A$, and a concept $C \in \mathcal{C}(A, B, R)$ such as its equivalence class $[C]_D$ of the induced equivalence relation is not a singleton. We have that $C_m \notin [C]_D$ if and only if one of the following statements is satisfied:

- There exists at least one attribute $a \in D$ such that $C_m = (a^\downarrow, a^\uparrow)$.
- There exists a concept $C^* \in \mathcal{C}(A, B, R)$, such as $C^* = (a^{*\downarrow}, a^{*\uparrow})$ with $a^* \in D$, $C^* \notin [C]_D$ and $C_M \not\leq C^*$. Moreover, $\overline{C^*}$ is in a meet-irreducible decomposition $\{\overline{C_j} \in M_F(D, B, R|_{D \times B}) \mid j \in J\}$ of $\overline{C_m}$.

In addition, we analyzed Theorem 1 when the considered subset of attributes does not contain unnecessary attributes. This fact is the usual case in FCA attribute reduction and simplifies the detection of equivalence classes which are not convex sublattices of the original concept lattice. Furthermore, under this consideration, we also proved that when the original concept lattice is isomorphic to a distributive lattice the induced equivalence classes by the reduction are always sublattices.

III. CONCLUSIONS

Consequently, all the results shown in this paper have a significant importance, for example, in the application of local congruences, due to characterizing the cases when classes are not sublattices, we will be useful to know which classes will be affected when a local congruence be applied after an attribute reduction mechanism.

REFERENCES

- [1] R. G. Aragón, J. Medina, and E. Ramírez-Poussa. Identifying non-sublattice equivalence classes induced by an attribute reduction in fca. *Mathematics*, 9(5), 2021.
- [2] R. G. Aragón, J. Medina, and E. Ramírez-Poussa. Impact of local congruences in attribute reduction. In M.-J. Lesot, S. Vieira, M. Z. Reformat, J. P. Carvalho, A. Wilbik, B. Bouchon-Meunier, and R. R. Yager, editors, *Information Processing and Management of Uncertainty in Knowledge-Based Systems*, volume 1239, pages 748–758, Cham, 2020. Springer International Publishing.
- [3] R. G. Aragón, J. Medina, and E. Ramírez-Poussa. Reducing concept lattices by means of a weaker notion of congruence. *Fuzzy Sets and Systems*, 2020. <https://doi.org/10.1016/j.fss.2020.09.013>.
- [4] M. J. Benítez-Caballero, J. Medina, E. Ramírez-Poussa, and D. Ślęzak. A computational procedure for variable selection preserving different initial conditions. *International Journal of Computer Mathematics*, 97(1-2):387–404, 2020.

Non-monotonic multi-adjoint logic programming with constraints

M. Eugenia Cornejo
Department of Mathematics
University of Cádiz
Cádiz, Spain
mariaeugenia.cornejo@uca.es

David Lobo*
Department of Mathematics
University of Cádiz
Cádiz, Spain
david.lobo@uca.es

Jesús Medina
Department of Mathematics
University of Cádiz
Cádiz, Spain
jesus.medina@uca.es

Abstract—This work concerns the research recently published in M. E. Cornejo, D. Lobo and J. Medina, “Extended multi-adjoint logic programming”, *Fuzzy Sets and Systems*, vol. 388, pp. 124–145, 2020. The reader is referred to that source for a full discussion and examples of the work.

Index Terms—multi-adjoint logic programming, non-monotonic logic programming, negation operator, stable models

I. INTRODUCTION

Medina et al. introduced multi-adjoint logic programming [8] in order to generalize different non-classical logic programming frameworks [4], [9]. The foundations of a multi-adjoint logic program is a complete lattice endowed with different adjoint pairs, which allow to obtain several generalizations of modus ponens. Recently, the inclusion of a negation operator in multi-adjoint logic programming was carried out in [1], giving rise to a first non-monotonic multi-adjoint logic approach, which generalizes other current frameworks [6], [7].

In [2], we defined a general non-monotonic logic programming language which shares the multi-adjoint philosophy. Besides permitting different adjoint pairs, the syntax of the so-called extended multi-adjoint logic programming is characterized by the use of constraints and of a special type of aggregator operator. Namely, such aggregation enables to consider, among others, different negation operators in the body of the same rule of the logic program. In addition to defining the syntax and the semantics of this logic programming paradigm, a mechanism for obtaining a multi-adjoint normal logic program from an extended multi-adjoint logic program has been shown, which entails a correlation between the semantics of both logic programming languages.

II. EXTENDED MULTI-ADJOINT LOGIC PROGRAMMING

The underlying idea of extended multi-adjoint logic programs can be stated as follows:

* Corresponding author

Partially supported by the 2014–2020 ERDF Operational Programme in collaboration with the State Research Agency (AEI) in project PID2019-108991GB-I00, and with the Department of Economy, Knowledge, Business and University of the Regional Government of Andalusia in project FEDER-UCA18-108612, and by the European Cooperation in Science & Technology (COST) Action CA17124.

- On the one hand, constraints simulate the limitation of a certain property, attribute or characteristic by an upper bound. Therefore, the inclusion of constraint rules of the form $\langle c \leftarrow_i B; \top \rangle$ in a logic programming language might be enormously useful in what regards applications, being c an element of a lattice representing a degree of limitation.
- On the other hand, as well as the multi-adjoint paradigm aims at relaxing limitations for modus ponens, it would be desirable to provide freedom for the non-monotonic behaviour of propositional symbols. For instance, it could be advantageous allowing different negations in the body of the same rule. This can be done by considering *extended aggregators*, that is, an n -ary mapping $@^e$ defined as:

$$@^e(x_1, \dots, x_n) = @^e[x_1, \dots, x_m; x_{m+1}, \dots, x_n]$$

such that it is order-preserving in the first m arguments and order-reversing in the last $n - m$ arguments.

According to the foregoing remarks, the syntax of extended multi-adjoint logic programming is defined as follows.

Definition 1: Let $(L, \preceq, \leftarrow_1, \&_1, \dots, \leftarrow_n, \&_n, @_1^e, \dots, @_k^e)$ be an extended multi-adjoint lattice with greatest element $\top$. An *extended multi-adjoint logic program* is a finite set of weighted rules of the form

$$\langle p \leftarrow_i @^e[p_1, \dots, p_m; p_{m+1}, \dots, p_n]; \vartheta \rangle$$

and constraint rules of the form

$$\langle c \leftarrow_i @^e[p_1, \dots, p_m; p_{m+1}, \dots, p_n]; \top \rangle$$

where $i \in \{1, \dots, n\}$, $@^e \in \{@_1^e, \dots, @_k^e\}$, $\vartheta, c \in L$ and $p_{s_1} \neq p_{s_2}$, for all $s_1, s_2 \in \{1, \dots, n\}$ with $s_1 \neq s_2$.

The semantics of extended multi-adjoint logic programs is defined similarly to the stable model semantics of multi-adjoint normal logic programs [1]. Namely, the stable model semantics lies on the principles of the Gelfond-Lifschitz reduct [5].

A detailed procedure to transform an extended multi-adjoint logic program into a semantically equivalent multi-adjoint normal logic program was introduced in [2]. Such a procedure is shown in two steps:

- 1) Firstly, given an extended multi-adjoint normal logic program, its constraint rules are converted into regular

rules, in such a way that the stable models of the original and the final logic programs coincide.

- 2) Then, the remaining rules are written in terms of a single non-monotonic unary mapping, which turns out to be an involutive negation. It needs to be stressed that such transformation is carried out by means of continuous mappings. An interesting outcome of this fact is the possibility to apply a sufficient condition for the existence of stable models in multi-adjoint normal logic programs [1] for constraint-free extended multi-adjoint logic programs.

III. CONCLUSIONS

The research carried out in [2] enables to make use of a flexible language like extended multi-adjoint logic programs in order to model real-world problems, and then translate them into multi-adjoint normal logic programs to handle compact simple programs with the same meaning. Furthermore, this procedure can be complemented with the methods shown in [3], where a multi-adjoint normal logic program is translated into a core fuzzy answer set program. As highlighted in [6], core fuzzy answer set programs are easier to implement and to reason about from a computational point of view.

The completed transformation considerably increases the potential of extended multi-adjoint logic programs to model real-life problems, since modelling the information contained in a text or in a database by decision rules and the interpretation of those rules will be easier through extended multi-adjoint logic programs, and its translation into a core fuzzy answer set program will facilitate the simulation and computation of the consequences/deductions from the program.

REFERENCES

- [1] M. E. Cornejo, D. Lobo and J. Medina, "Syntax and semantics of multi-adjoint normal logic programming", *Fuzzy Sets and Systems*, vol. 345, pp. 41–62, 2018.
- [2] M. E. Cornejo, D. Lobo and J. Medina, "Extended multi-adjoint logic programming", *Fuzzy Sets and Systems*, vol. 388, pp. 124–145, 2020.
- [3] M. E. Cornejo, D. Lobo and J. Medina, "Relating multi-adjoint normal logic programs to core fuzzy answer set programs from a semantical approach", *Mathematics*, vol. 8 (6), 881, 2020.
- [4] C.V. Damásio, L. M. Pereira, "Monotonic and residuated logic programs", in *Symbolic and Quantitative Approaches to Reasoning with Uncertainty (ECSQARU'01)*, in *Lecture Notes in Artificial Intelligence*, vol. 2143, pp. 748–759, 2001.
- [5] M. Gelfond and V. Lifschitz, "The stable model semantics for logic programming", in *Logic Programming: Proceedings of the Fifth International Conference and Symposium*, MIT Press, pp. 1070–1080, 1988.
- [6] J. Janssen, S. Schockaert, D. Vermeir and M. De Cock, "A core language for fuzzy answer set programming", *International Journal of Approximate Reasoning*, vol. 53 (4), pp. 660–692, 2012.
- [7] N. Madrid and M. Ojeda-Aciego, "On the existence and unicity of stable models in normal residuated logic programs", *International Journal of Computer Mathematics*, vol. 89 (3), pp. 310–324, 2012.
- [8] J. Medina, M. Ojeda-Aciego and P. Vojtáš, "Multi-adjoint logic programming with continuous semantics" in *Logic Programming and Non-Monotonic Reasoning (LPNMR'01)*, vol. 2173, pp. 351–364, 2001.
- [9] P. Vojtáš, "Fuzzy logic programming", *Fuzzy Sets and Systems*, vol. 124 (3), pp. 361–370, 2001.

Índice de autores

A

Abad-Navarro, Francisco, 806
Afsar, Sezin, 465
Aguado, Esther, 116
Aguilera-Martos, Ignacio, 685
Aguirre, Aitor, 108
Alba, Enrique, 23, 219, 941, 947, 969, 970, 1013
Alberto da Silva Abreu, Luiz, 168
Alcalde, Cristina, 262, 305
Aledo, Juan A., 128, 247, 566
Alegre, Fernando, 989
Alfaro, Juan C., 128
Almagro Cádiz, Mario, 1026
Alonso, Maria Teresa, 629
Alonso, Pedro, 317
Alonso Ayuso, Antonio, 593
Alonso Betanzos, Amparo, 613, 735
Alonso-Moral, Jose M., 274, 298
Alonso-Rial, Adrián, 955
Alonso-Weber, Juan M., 217
Alvarado Díaz, Jorge, 505, 537, 1009
Álvarez, Verónica, 215
Alvarez-Valdes, Ramon, 629
Amaya, Jhon Edgar, 499
Andrés Arroyo, Marta, 235
Andreu-Vilarroig, Carlos, 566
Antunes-Santos, Felipe, 83
Aragón Royón, Francisco, 775
Araujo, Lourdes, 824
Arbelaitz, Olatz, 189, 818
Arias, Joaquín, 190
Arias-Rodil, Manuel, 685, 846
Arnaiz-González, Álar, 691, 783, 785, 787
Arnaiz-Rodríguez, Adrián, 667
Arregui Pérez de Loza, Iñigo, 763
Arruti, Andoni, 201, 203
Arzamendia, Mario, 77
Asencio Martín, Lucía, 188
Atencia Payares, Luz Karime, 889
Aznar, Fidel, 102

B

Baena-García, Manuel, 648
Bahamonde, Antonio, 29, 735
Ban, Yue, 936
Barbudo Lunar, Rafael, 90
Barranco, Carlos D., 717
Barrenetxea Berasategi, Xabier, 180
Barri Khojasteh, Samad, 800
Bautista-Valhondo, Joaquín, 471
Baz, Juan, 331
Beltrán-Royo, César, 167
Benítez, José M., 775
Benítez-Caballero, María José, 393
Benito-Picazo, Jesús, 729
Bharatia, Harshal, 96
Bibiloni, Pedro, 211, 1031
Bidaurrezaga Barrueta, Arkaitz, 1062
Bielza, Concha, 840
Blanco-Mallo, Eva, 1041

Bobillo, Fernando, 351, 989
Bodi, Maria, 793
Bolón Canedo, Verónica, 735, 1041
Borrajo, Daniel, 186
Botana-López, Alberto, 955
Bueno, Sara, 1023
Bueno-Crespo, Andrés, 543
Bugarín-Diz, Alberto, 292, 298
Burusco, Ana, 262, 305
Bustince, Humberto, 223, 317, 358, 370, 376, 511, 613, 652

C

Cabrera, Inma P., 389
Cabrera, Marcelino, 434
Calleja, Pablo, 53
Calvo-Lluch, África, 723
Camacho, David, 523, 619
Camero, Andrés, 219
Campillo-Ageitos, Elena, 824
Campomanes-Álvarez, B. Rosario, 747
Campomanes-Álvarez, Carmen, 747
Campos, Manuel, 812
Campos Knupp, Diego, 168
Canabal-Juanatey, Mariña, 298
Cano Jiménez, Jesús, 889
Canovas-Segura, Bernardo, 812
Carabe, Luis, 209
Carmona, Cristobal J., 227, 231
Carmona, Enrique, 1023
Carmona, Javier, 446
Carmona del Jesús, Cristóbal José, 661
Carranza-García, Manuel, 658, 1091
Carrasco, Jacinto, 685
Carrasco-Villalón, Rocio, 717
Casado Ceballos, Alejandra, 472, 490
Casado-García, Ángela, 76, 134
Casado-Vara, Roberto, 712
Casanova, Jorge, 936
Casas-Martínez, Pedro, 490
Casas-Ortiz, Alberto, 997
Cascallar-Fuentes, Andrea, 292
Casillas, Jorge, 706
Castaño Rosa, Raúl, 919
Castejón, Federico, 1023
Castell Díaz, Javier, 806
Castellano-Quero, Manuel, 205
Castiblanco Ruiz, Fabian Alberto, 329
Castillo, Jhonny, 907
Castillo, Pedro, 505, 977
Castillo-López, Aitor, 376, 511
Castrillon-Santana, Modesto, 64
Castro, Jairo, 174
Castro, Javier, 241, 391
Catala, Alejandro, 298
Cavero Díaz, Sergio, 452, 1047
Ceberio Uribe, Josu, 453, 459
Cermeno, Eduardo, 209
Cernuda, Carlos, 108
Chanava, Walter, 757
Charte, David, 235

Charte, Francisco, 231
 Chávez, Francisco, 505, 537
 Chica, Manuel, 280
 Chicano, Francisco, 941, 947
 Cichos, Frank, 213
 Cikel, Kevin, 77
 Cintrano, Christian, 969, 970
 Cisneros, Eusebio, 901
 Cobo, Manuel Jesus, 337
 Colmenar, José Manuel, 587, 1116
 Colomine, Feijoo, 429, 517
 Conti, Dante, 763
 Corchado, Juan M., 712
 Cordero, Pablo, 389
 Cordon, Óscar, 180, 280, 601, 613, VII
 Cornejo, M. Eugenia, 398
 Corral de la Calle, Javier, 806
 Correia da Silva Jardim, Lucas, 168
 Cortés, Juan Carlos, 555, 566
 Costa Cortez, Nahuel, 1110
 Cotta, Carlos, 429, 499, 517
 Cristóbal, Álvaro, 351
 Cruz Corona, Carlos, 168
 Cuadrado, David, 793
 Cuesta-Infante, Alfredo, 114, 167
 Cuesta Llavona, Cristina, 985

D

DAHI, Zakaria Abdelmoiz, 941, 947
 de la Caba, Koro, 262
 de La Cal, Enrique, 800
 de la Morena Barrio, M. Eugenia, 806
 de La Paz López, Félix, 1026
 del Campo Ávila, José, 1079
 Del Jesús Díaz, María José, 41, 227
 Del Ser Lorente, Javier, 529
 Delgado, Alvaro, 572
 Delgado-Osuna, José Antonio, 753
 Diaz, Irene, 192, 1017
 Díaz, Josefa, 505
 Díaz-Díaz, Norberto, 717
 Diez, Francisco Javier, 221, 846
 Díez, Jorge, 29, 741, 985
 Díez-Oliván, Alberto, 180
 Diosdado, Daniel, 166
 Domínguez, César, 58, 89, 795
 Dominguez, Enrique, 1053
 Dominguez-Catena, Iris, 345
 Dorado Llamas, Ignacio, 459
 Duarte, Abraham, 452, 472, 601, 1047, 1097, 1103, 1116

E

E-Martín, Yolanda, 186
 Echanobe, Javier, 931, 936
 Elizondo, David, 231
 Elorza, Anne, 1074
 Escobar Mimbrenra, Ángela, 41
 Esnaola-Gonzalez, Iker, 221
 Espínola, Rosa, 241, 391
 Esteban Toscano, Aurora, 673
 Estévez, Sonia, 679

F

Fañez Kertelj, Mirko, 800
 Fdez-Olivares, Juan, 115

Fedorova, Stanislava, 961
 Fernández, Alberto, 235
 Fernández, Guillermo, 247
 Fernández, Javier, 223, 358, 376, 511, 613
 Fernández Ares, Antonio, 977
 Fernández-Breis, Jesualdo Tomás, 802, 806
 Fernández de Vega, Francisco, 411, 505, 537, 1009
 Fernandez-Izquierdo, Alba, 122
 Fernández Leiva, Antonio J., 429, 499, 517
 Fernández-Madrigal, Juan-Antonio, 205
 Fernández Serrano, Ricardo, 153
 Ferrer, Javier, 23, 1013
 Ferrero-Jaurrieta, Mikel, 358
 Fischer, Alexander, 213
 Flaminio, Tommaso, 256
 Flores Gallego, Maria Julia, 70
 Fontenla-Romero, Oscar, 141
 Franco De Los Ríos, Camilo, 329
 Fresno Fernández, Víctor, 1026
 Fuentes-Fernández, Rubén, 976
 Fuentetaja, Raquel, 186
 Fumanal Idocin, Javier, 370, 511

G

Gabilondo, Iñigo, 818
 Galar, Mikel, 345
 Galeano, Jesús, 446
 Gallego-Fernández, Javier, 292
 Gallego-Ledesma, Luis, 658
 Gámez, José A., 128, 247, 640
 Garanganga, Bradwell, 895
 Garcia-Andoin, Mikel, 931
 García-Aragón, Roberto, 397
 García Arenas, Maribel, 977
 García-Barzana, Marta, 685
 García-Castro, Raúl, 122
 García-Cerezo, Alfonso, 205
 García-Domínguez, Manuel, 76, 795
 García-Gil, Diego, 685
 García Gómez, Pablo, 147
 García-González, Jorge, 729
 García-Gutiérrez, Jorge, 1091
 García Herrero, Jesús, 881
 García-Magariño, Iván, 976
 García-Martínez, Carlos, 753
 García-Osorio, César, 783, 787
 García Pardo, Eduardo, 452, 490, 581, 593, 1047
 García-Pedrero, Angel, 836
 García-Retuerta, David, 712
 García-Rodríguez, Sara, 1017
 García Sánchez, Pablo, 499, 977
 Garcia-Vico, Angel M., 227, 231
 García-Zamora, Diego, 323
 Garcarena, Unai, 405
 Garitano, Iñaki, 108
 Garmendia, Luis, 364
 Garnica, Oscar, 153, 549
 Garrido, Roberto, 834
 Garrido-Labrador, José Luis, 667
 Garrido-Merchán, Eduardo C., 159, 188, 703
 Garrocho, Alejandro, 174
 Gatt, Albert, 274
 Gauchola, Jaume, 832
 Gayoso Heredia, Marta, 923

Gibaja, Eva, 697
 Gibert, Karina, 763
 Gil Borrás, Sergio, 593
 Gil-Merino, Rodrigo, 941, 947
 Giménez-Palacios, Iván, 629
 Giménez-Solano, Vicente Miguel, 802
 Giráldez-Cru, Jesús, 280
 Godo, Lluís, 256
 Gomez, Daniel, 241, 364
 Gómez, Daniel, 391
 Gomez, Fernando, 174
 Gomez, Josep, 793
 Gómez, Juan Antonio, 913
 Gómez, Marisol, 370
 Gómez-Sanz, Jorge J., 976
 Gómez-Valadés, Alba, 838
 González, Antonio, 268
 Gonzalez, Víctor, 800
 González-Briones, Alfonso, 976
 González-Carvajal, Santiago, 703
 González Doncel, Gaspar, 153
 González García, Pedro, 41, 227, 231, 661
 González-Hidalgo, Manuel, 211, 1031
 González-Moya, Francisco Javier, 691
 González Naharro, Luis, 70
 Gonzalez-Pardo, Antonio, 607
 González-Pérez, Beatriz, 560
 González-Reina, Adán, 723
 González Rodríguez, Inés, 147
 Gonzalez-Sanchez, Belen, 160, 478
 Gonzalo Martin, Consuelo, 834, 836
 Gordo-Martín, Daniel, 955
 Gragera, Alba, 186
 Gragera, Miguel Ángel, 543
 Gregor, Derlis, 77
 Guerrero, Juan José, 23, 1013
 Guerrero, Pedro, 262
 Guevara, Juan Antonio, 391
 Guijarro-Berdiñas, Bertha, 141
 Guijo-Rubio, David, 641, 660
 Gurrutxaga, Ibai, 818
 Gutiérrez, Alberto, 549
 Gutiérrez, Inmaculada, 241, 391
 Gutierrez, Juan Carlos, 174
 Gutierrez, Nuria, 834
 Gutiérrez, Pedro Antonio, 641, 660
 Gutiérrez-Avilés, David, 639
 Gutiérrez Reina, Daniel, 5, 11, 17, 77

H

Halodova, Patricie, 153
 Hemberg, Erik, 233
 Heras, Jónathan, 58, 76, 89, 134, 795
 Heredia Vizcaino, Diana, 907
 Hermoso Peralo, Roberto, 889
 Hernández-Díaz, Alfredo G., 627
 Hernández-García, Sergio, 114
 Hernando, Leticia, 1074
 Herrán, Alberto, 587
 Herrera, Francisco, 225, 235, 280, 529, 652, 685
 Hervás-Martínez, César, 641, 660
 Hidalgo, José Ignacio, 153, 549, 572
 Holubek, Viktor, 213
 Huertas-García, Álvaro, 619

Huertas-Tato, Javier, 523, 619
 Huidobro, Pedro, 317
 Huitzil, Ignacio, 351, 989

I

Idocin, Javier, 613
 Iglesias, Natalia, 812
 Iglesias Martínez, José Antonio, 923
 Iglesias-Rey, Sara, 83
 Inés Armas, Adrián, 58, 76
 Irazabal-Urrutia, Oier, 189
 Irigoyen, Eloy, 203
 Iturbe, Mikel, 108

J

Jabba, Daladier, 907
 Janis, Vladimír, 317
 Jara, Leonardo, 268
 Jariego Pérez, Luis C., 159
 Jelić, Marko, 221
 Jiménez, Daniel, 453
 Jiménez Agudelo, Yury Andrea, 757
 Jimenez Linares, Luis, 286
 Jiménez Merino, Ernesto, 593
 Jiménez-Navarro, M. J., 659
 Jin, Yaochu, VIII
 José da Silva Neto, Antônio, 168
 Juarez, Jose M., 812
 Juez-Gil, Mario, 783, 785, 787, 993

K

Kronberger, Gabriel, 153

L

Labella Romero, Álvaro, 323, 385, 1037
 Lafuente, Julio, 652
 Lanchares, Juan, 549
 Lara-Benítez, Pedro, 658
 Larrañaga, Pedro, 840
 Leceta, Itsaso, 262
 Ledezma Espino, Agapito, 869, 913, 919, 923
 Legaz-García, María del Carmen, 802
 Lillo-Saavedra, Mario, 836
 Llamazares Rodríguez, Bonifacio, 357
 Llerena Caña, Juan Pedro, 881
 Llopis-Ibor, Laura, 167
 Lobo, David, 398
 López, Beatriz, 832
 López, David, 685
 Lopez, Samuel, 174
 Lopez, Victoria, 679, 757
 Lopez-Molina, Carlos, 83, 376
 López-Navarro, Elena, 555
 López-Nozal, Carlos, 783, 785, 993
 López-Riobó Botana, Iñigo Luis, 735
 López-Rubio, Ezequiel, 729
 López-Sánchez, Ana Dolores, 627
 López Vargas, Ascension, 919
 Lorenzo-Navarro, Javier, 64
 Losada Casado, Adrián, 881
 Lozano, José Antonio, 215, 453, 1074, 1085, 1106
 Lozano-Orsorio, Isaac, 601, 1103
 Luaces, Oscar, 29
 Luengo, Julián, 685
 Luna, Francisco, 446

Luna, Jose Maria, 769
 Luna-Romera, José M., 658
 Luque, Gabriel, 941, 947
 Luque, Manuel, 846
 Luque-Baena, Rafael M., 1001, 1005
 Luzón, María V., 225

M

Madrid, Nicolas, 395
 Magan Lopez, Elena, 875
 Magdalena, Luis, 364
 Maldonado, José Alberto, 802
 Mandow, Lawrence, 135, 187
 Marcos, Mar, 802
 María-Armingol, José, 857, 863
 Marín-Plaza, Pablo, 857, 863
 Mariotti, Ettore, 274
 Martí, Rafael, 631, 633
 Martín, Alejandro, 523, 619
 Martin, Jose Ignacio, 201, 203
 Martín-Albo, Sergio, 135
 Martín Chozas, Patricia, 47, 53
 Martín Santamaría, Raúl, 587, 1116
 Martínez, Aritz, 529
 Martínez, Luis, 323, 385, 1037
 Martínez, María, 331
 Martínez-Álvarez, F., 659
 Martínez-Cámara, Eugenio, 225
 Martínez-Cendán, Juan-Pedro, 543
 Martínez Crespo, Jorge, 919
 Martínez-Eguiluz, Maitane, 189, 818
 Martínez España, Raquel, 543
 Martínez-Gavara, Anna, 627, 631, 633
 Martínez-Más, José, 543
 Martinez Rodriguez, Raquel, 201, 203
 Martinez-Romo, Juan, 824
 Martínez-Salvador, Begoña, 802
 Martínez-Sanchez, Antonio, 207
 Martínez Tomás, Rafael, 838
 Massana, Joaquim, 832
 Massanet, Sebastia, 223
 Mata, Eloy Javier, 58, 89, 795
 Mazuelas, Santiago, 215
 Medina, Jesús, 393, 397, 398
 Medina Burga, Melissa, 901
 Melgar-García, Laura, 639
 Menasalvas Ruiz, Ernestina, 804, 834
 Mendiburu, Alexander, 405
 Mercado Palomino, Elia, 775
 Merino, María, 1106
 Miguel, Diego, 993
 Millán, Laura, 153
 Miñana, Juan José, 304, 311
 Mir-Fuentes, Arnau, 83
 Molina-Cabello, Miguel A., 1001, 1005
 Molina Cabrera, Daniel, 529
 Molina López, José Manuel, 881
 Montañana, Ricardo, 640
 Montero, Javier, 329, 364
 Montes, Susana, 253, 317, 331, 1017
 Montiel-Ponsoda, Elena, 47, 53
 Mora, Antonio, 977
 Mora López, Llanos, 1079
 Morales, Antonio, 812

Morales-Bueno, Rafael, 648
 Morell Martínez, José Ángel, 970
 Moreno Garcia, Juan, 286
 Moreno-Rebato, Mar, 190
 Morita Hernández, Martín, 411, 1009
 Moscardó-García, Ana, 555
 Móstoles Rodríguez, Roberto, 881
 Muguerza, Javier, 201, 203, 818
 Muiños-Landín, Santiago, 213, 955
 Muñoz, Andres, 543
 Murombedzi, James, 895
 Murueta-Goyena, Ane, 818

N

Nesmachnow, Sergio, 191
 Nikolov-Vasilev, Vladislav, 115
 Novoa-Hernández, Pavel, 484
 Novoa-Paradela, David, 141
 Nuñez, Concepción, 560
 Núñez-Molina, Carlos, 115
 Núñez-Peiró, Miguel, 923
 Nyakutambwa, Trymore, 895

O

Ojeda-Aciego, Manuel, 395
 Ojeda-Hernandez, Manuel, 389, 1059
 Olivares-Gil, Alicia, 667
 Onaindia, Eva, 166
 O'Reilly, Una-May, 233
 Orellana, Aitor Manuel, 337
 Orengo, Juan C., 566
 Ortégón, Juan Carlos, 434
 Ortiz, Alberto, 304, 311
 Ortiz, Esaú, 304, 311
 Ortiz-González, Ana, 543
 Ortiz-Reina, Sebastián, 543
 Ortiz-Toro, César Antonio, 834, 836
 Ossowski, Sascha, 190
 Otero-Tranchero, Jacobo, 955
 Ottonello, Nicolás, 191

P

Pacheco, Santiago, 191
 Palacios, Juan Jose, 465
 Pantrigo, Juan José, 167
 Parra, Daniel, 549
 Parreño, Francisco, 629
 Pascual, Vico, 58, 89, 795
 Pascual Triana, José Daniel, 235
 Paternain, Daniel, 345
 Pavón, Juan, 976
 Pedraza, Tatiana, 387
 Pelta, David, 434, 484
 Peralta, Federico, 5, 11, 17
 Peregrin, Antonio, 174, 417
 Pereira Dimuro, Graçaliz, 358, 370
 Perez, Aritz, 1062
 Pérez, Carlos J., 423
 Perez, Ignacio Javier, 337
 Pérez, Raúl, 115, 268
 Pérez-Asurmendi, Patrizia, 357
 Pérez de la Cruz, José Luis, 135, 187
 Perez-Fernandez, Raul, 331
 Pérez-Martín, Jorge, 846
 Pérez Martos, Luis Alfonso, 661

Pérez-Núñez, Pablo, 29, 741, 985
 Pérez-Peló, Sergio, 472, 607, 633, 1097
 Perona, Iñigo, 818
 Pineda-Palencia, Ismael, 1001
 Pineda Torres, Ivo, 851
 Pinheiro Domingos, Roberto, 168
 Porras Castaño, Javier, 1068
 Poveda-Villalón, María, 122
 Poyatos Amador, Javier, 529
 Pozanco, Alberto, 186
 Prieto Santamaría, Lucía, 804
 Puerta Callejón, José Miguel, 70, 247, 640
 Pujić, Dea, 221
 Pulido, Inmaculada, 174

Q

Quiñones, Alejandro, 446
 Quirós, Pelayo, 747

R

R. Vela, Camino, 147, 192, 465
 Rais, Jasmina, 102
 Ramírez, Aurora, 35
 Ramírez Poussa, Eloisa, 393, 397
 Ramírez-Sanz, José Miguel, 667
 Ramos-Jiménez, Gonzalo, 729
 Ramos-Soto, Alejandro, 292
 Ranchal-Parrado, Daniel, 753
 Raya, Óscar, 832
 Reguera-Bakhache, Daniel, 108
 Remeseiro, Beatriz, 1041
 Remezal-Solano, Manuel, 543
 Riaño, David, 793
 Riaño Sánchez, Miguel A, 840
 Rico, Noelia, 192, 1017
 Rico Prieto, Iván, 889
 Riera, Juan Vicente, 223
 Rincón, Mariano, 838
 Riquelme, José C., 1091
 Rodríguez, Alejandro, 793
 Rodríguez, Álvaro, 17
 Rodríguez, J. Tinguaro, 329
 Rodríguez, Juan J., 691, 783, 785, 787
 Rodríguez, Miguel Angel, 417
 Rodríguez, Rosa M., 323, 385, 1037
 Rodríguez-Barroso, Nuria, 225
 Rodríguez-Benitez, Luis, 286
 Rodríguez Doncel, Victor, 53
 Rodríguez-García, José A., 190
 Rodríguez Gómez, Francisco, 1079
 Rodríguez González, Alejandro, 804, 834
 Rodríguez-López, Jesús, 387
 Rodríguez-Martínez, Iosu, 358, 652
 Rodríguez-Revuelta, Antonio, 741
 Roiz Pagador, Joaquín, 723
 Rojas Ramos, Fernando, 851
 Romero, Alberto, 166
 Romero, José Raúl, 35, 90
 Rossi, Claudio, 116
 Royo, Didac, 76
 Rubio-Escudero, Cristina, 639
 Rubio Perona, Fernando, 70
 Ruiz, Roberto, 723
 Ruiz-Casado, Jose Luis, 1005
 Ruiz-de-la-Cuadra, Alejandro, 857, 863

Ruiz-Rivas Hernando, Ulpiano, 919
 Rus Rus, Rafael, 869
 Rutabingwa, Frank, 895

S

Saborido, Rubén, 23, 229, 1013
 Said Hung, Elías, 907
 Salamanca, Juan Jesús, 253
 Salazar-Ramirez, Asier, 201, 203
 Salerno, Mauricio, 186
 Sánchez, Almudena, 537
 Sanchez, Arturo, 901
 Sánchez, José L., 560
 Sanchez-Gomez, Jesus M., 423
 Sánchez-Guevara Sánchez, Carmen, 923
 Sánchez-Nielsen, Elena, 64
 Sánchez-Oro, Jesús, 472, 490, 1097, 1103
 Sánchez Ramos, Luciano, 1110
 Sanchis, Araceli, 217, 875
 Sanmartin, Paul, 907
 Santamaría, Gonzalo, 89
 Santana, Roberto, 405, 642, 1062, 1085
 Santander-Jiménez, Sergio, 160, 440, 478
 Santos, Olga C., 997
 Sanz, Ricardo, 116
 Sanz-Fernández, Ana, 923
 Sarmiento-Ramírez, Jonay, 64
 Saunders, Anthony, 292
 Sedano, Javier, 800
 Sedova, Ekaterina, 187
 Seker, Huseyin, 227
 Sesma, Mikel, 376, 652
 Sesmero Lorente, María Paz, 217, 875
 Sierra, Basilio, 201
 Sierra, Ignacio, 717
 Simmonds, Jose, 913
 Slezak, Dominik, 393
 Sousa, Leonel, 440

T

Tabik, Siham, 529
 Taillefer, Eugenia, 648
 Talavera-Martínez, Lidia, 211, 1031
 Tapia, Alejandro, 17
 Tena, Félix, 572
 Teran, Diego, 1053
 Tomasevic, Nikola, 221
 Toral, Sergio, 5, 11, 17, 77
 Torres, J. F., 659
 Torrontegui, Erik, 936
 Toutouh, Jamal, 219, 233, 969, 970
 Trasierras, Antonio Manuel, 769
 Troncoso, Alicia, 639, 659

U

Unanue, Imanol, 1106
 Urdanpilleta, Marta, 262
 Ureña Espa, Aurelio, 775
 Urio, Asier, 370
 Ursúa, Pablo, 652

V

Vadillo, Jon, 642, 1085
 Valenzuela, Juan Francisco, 446
 Valero, Óscar, 304, 311, 387

Valverde, Gabriel, 560
Varela, Alejandro, 572
Vargas, Víctor Manuel, 641, 660
Vázquez-Flores, Karen Leticia, 47
Vega-Rodríguez, Miguel A., 160, 423, 440, 478
Velasco, José Manuel, 537, 560, 572
Velez-Estevez, Antonio, 337
Vellido, Alfredo, 830
Vellido, Ignacio, 115
Ventura, Sebastián, 35, 90, 673, 753, 769
Vidaurre, Carmen, 370
Villanueva, Rafael-J., 555, 566
Villar, David, 706
Villar, José Ramón, 800

Villar-Rodríguez, Guillermo, 523
Villegas Cortez, Juan, 411
Villota, María, 89

Y

Yagüe-Cuevas, David, 857, 863
Yanes, Samuel, 5, 11, 17
Yuste, Javier, 581

Z

Zafra Gómez, Amelia, 673, 697
Zapata, Miguel Ángel, 76
Zapata, Pablo, 446
Zurutuza, Urko, 108