

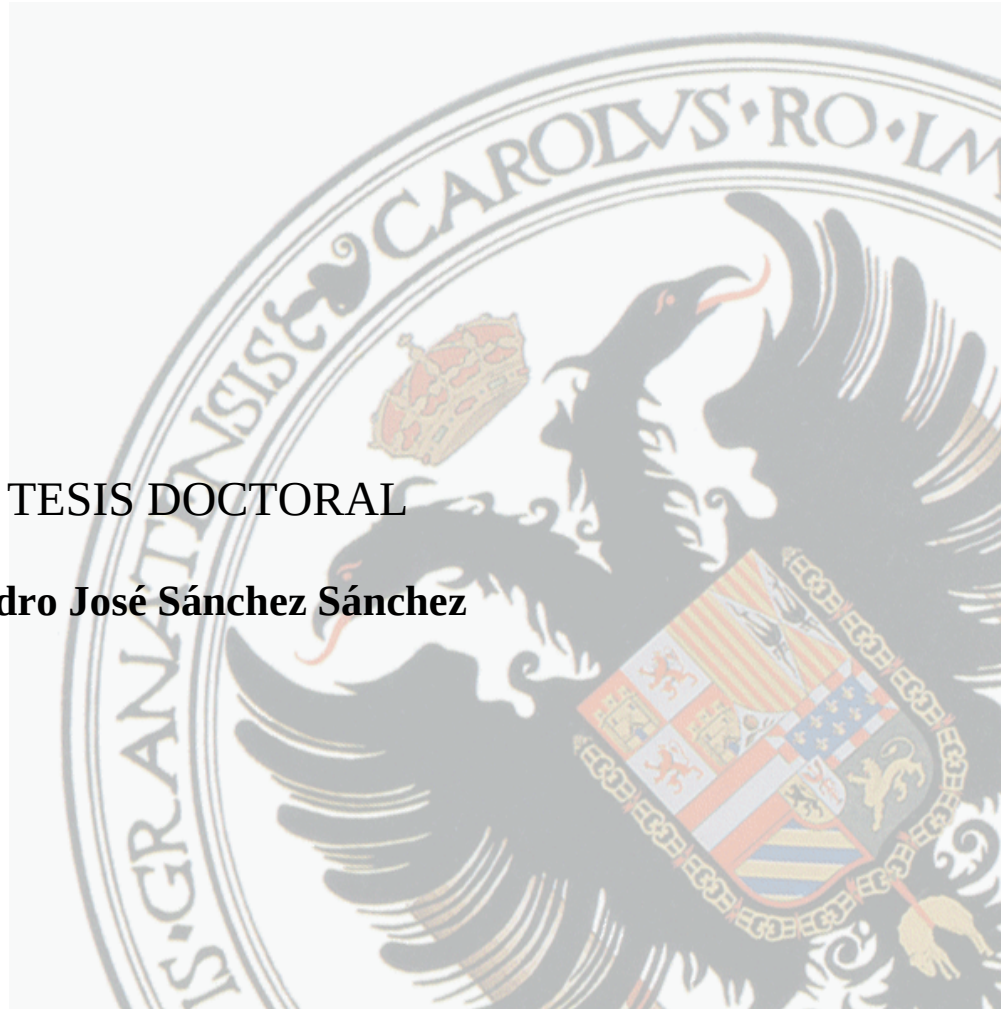


UNIVERSIDAD DE GRANADA
E.T.S. DE INGENIERÍA INFORMÁTICA
Departamento de Ciencias de la Computación e Inteligencia Artificial

MODELOS PARA LA COMBINACIÓN DE PREFERENCIAS EN TOMA DE DECISIONES: HERRAMIENTAS Y APLICACIONES

TESIS DOCTORAL

Pedro José Sánchez Sánchez





UNIVERSIDAD DE GRANADA

**MODELOS PARA LA COMBINACIÓN DE
PREFERENCIAS EN TOMA DE DECISIONES:
HERRAMIENTAS Y APLICACIONES**

MEMORIA QUE PRESENTA
PEDRO JOSE SÁNCHEZ SÁNCHEZ
PARA OPTAR AL GRADO DE DOCTOR

DIRECTORES

FRANCISCO HERRERA TRIGUERO
DEPARTAMENTO DE CIENCIAS DE LA COMPUTACIÓN E
INTELIGENCIA ARTIFICIAL
UNIVERSIDAD DE GRANADA

LUIS MARTÍNEZ LÓPEZ
DEPARTAMENTO DE INFORMÁTICA
UNIVERSIDAD DE JAÉN

E.T.S. DE INGENIERÍA INFORMÁTICA

UNIVERSIDAD DE GRANADA

La memoria de tesis titulada **Modelos para la Combinación de Preferencias en Toma de Decisiones: Herramientas y Aplicaciones**, que presenta **D. Pedro José Sánchez Sánchez** para optar al grado de Doctor en Informática, ha sido realizada en el **Departamento de Ciencias de la Computación e Inteligencia Artificial** de la Universidad de Granada bajo la dirección de los doctores **D. Francisco Herrera Triguero** y **D. Luís Martínez López**.

Fdo.: Francisco Herrera Triguero

Director

Fdo.: Luis Martínez López

Director

D. Pedro José Sánchez Sánchez

Doctorando

Agradecimientos

Son muchas las personas a las que quiero reconocer su ayuda durante la realización de esta memoria pero, entre ellas, hay algunas que me gustaría destacar.

En primer lugar y como no podría ser de otra forma, quiero expresar mi más sentido agradecimiento a las dos personas que han hecho posible la elaboración de esta memoria de investigación, como son mis directores de tesis, Dr. Luis Martínez López y Dr. Francisco Herrera Trigero. Gracias por su dedicación, esfuerzo y confianza depositada en mí a lo largo de todo este tiempo.

Gracias a todos mis compañeros del Departamento de Informática de la Universidad de Jaén, y en especial al apoyo y amistad del grupo con el que comparto temas de trabajo e investigación como son Luis G., Franciso, Macarena y Manolo.

Y por último agradecer el apoyo recibido desde mi familia.

Gracias a todos ...

Índice general

Introducción	1
1. Problemas de Toma de Decisión en Ambiente Difuso: Modelado de Preferencias y el Enfoque Lingüístico Difuso	11
1.1. El Problema de la Toma de Decisión	12
1.1.1. Según el Número de Criterios	13
1.1.2. Según el Ambiente de Decisión	15
1.1.3. Según el Número de Expertos	17
1.2. Modelado de Preferencias	18
1.2.1. Estructuras para la Representación de Preferencias	19
1.2.2. Dominios de Expresión de Preferencias	24
1.3. Nociones y Conceptos Básicos de la Teoría de Conjuntos Difusos	30
1.3.1. Conjuntos Difusos y Función de Pertenencia	32
1.3.2. Definición Básica de Conjunto Difuso	33
1.3.3. Tipos de Funciones de Pertenencia	35
1.3.4. Principio de Extensión	36
1.3.5. Número difuso	37
1.4. El Enfoque Lingüístico Difuso	38
1.4.1. Elección del Conjunto de Términos Lingüísticos	41

1.4.2. Semántica del Conjunto de Términos Lingüísticos	44
1.4.3. Modelo de Representación Lingüístico Basado en 2-tuplas	47
2. Tratamiento de Información Heterogénea en Problemas de Toma de Decisión	51
2.1. Problema de Toma de Decisiones Definido en un Contexto Heterogéneo	53
2.2. Herramientas para Manejar Información Heterogénea	59
2.2.1. Selección del dominio de expresión unificado: CBTL	60
2.2.2. Funciones de Transformación a $F(S_T)$	62
2.3. Conclusiones	82
3. Modelo de Decisión para Problemas de Toma de Decisión en Contextos Heterogéneos	83
3.1. El Modelo de Decisión	83
3.1.1. Adquisición de la Información	85
3.1.2. Proceso de Agregación	86
3.1.3. Proceso de Explotación	92
3.2. Estudio de TDME sobre la adecuación de la instalación de un ERP	93
3.2.1. Introducción: Enterprise Resource Planning	93
3.2.2. Adecuación de un Sistema ERP: Proceso de Decisión	95
3.2.3. TD sobre la Instalación de un ERP	96
3.3. Conclusiones	109
4. Especificación UML para el Tratamiento de Información Heterogénea y el Modelo de Decisión Heterogéneo	111
4.1. Lenguaje Unificado de Modelado, UML	114
4.2. Biblioteca para el tratamiento de la información heterogénea	116

4.3. Biblioteca del modelo de toma de decisión heterogénea	125
4.4. Biblioteca para la instalación de un sistema ERP	133
4.4.1. Conclusiones	138
5. Evaluación de la Calidad Docente en las Universidades Andaluzas	139
5.1. Arquitectura del Sistema de Evaluación	141
5.2. Biblioteca de clases para el Proceso de Evaluación Docente	145
5.3. Prototipo del Sistema de Evaluación Docente	151
5.3.1. Interfaz del Administrador	152
5.3.2. Módulo de resultados	160
5.3.3. Interfaz del Alumno	167
5.4. Conclusiones	176
Conclusiones	179

Índice de figuras

1.	Fases de un Proceso de Toma de Decisión	4
1.1.	Ejemplos de números difusos	39
1.2.	Número difuso trapezoidal	44
1.3.	Definición semántica de la variable lingüística altura	45
1.4.	Conjunto de 5 etiquetas lingüísticas uniformemente distribuidas	46
1.5.	Distribuciones diferentes del concepto “excelente”	47
2.1.	Fases de un Proceso de Toma de Decisión	56
2.2.	Modelo de decisión en contexto heterogéneo	57
2.3.	Proceso de agregación de información heterogénea	57
2.4.	Unificación de la información heterogénea	60
2.5.	Un CBTL con 15 términos simétricamente distribuidos.	62
2.6.	Transformación de información numérica en $F(S_T)$	67
2.7.	Transformación de un valor numérico en un conjunto difuso en S	68
2.8.	Transformación de información intervalar en $F(S_T)$	69
2.9.	Los 4 casos no triviales de solapamiento de A y B	71
2.10.	Función de pertenencia de un intervalo	75
2.11.	Ejemplo de transformación de un intervalo	76
2.12.	Transformación de información lingüística en $F(S_T)$	77

2.13. Transformación $l_1 \in S$ en un conjunto difuso en S_T	79
2.14. Transformación de conjuntos difusos sobre el CBTL en 2-tupla lingüísticas	80
2.15. CBTL	81
3.1. Modelo de decisión en contexto heterogéneo	84
3.2. Conjunto S utilizado	86
3.3. CBTL Utilizado	88
3.4. Semántica del conjunto de etiquetas A	98
3.5. Semántica del conjunto de etiquetas B	98
3.6. Semántica del conjunto de etiquetas C	98
3.7. Semántica del conjunto de etiquetas D	99
3.8. CBTL de 15 etiquetas distribuidas simétricamente.	102
4.1. Ejemplo de diagrama de clase UML	115
4.2. Diagrama de clases para el tratamiento de la información heterogénea	117
4.3. Diagrama de clases del modelo de decisión heterogéneo	126
4.4. Diagrama de clases para la instalación de un sistema ERP	134
5.1. Encuesta para la evaluación de la calidad docente universitaria . . .	140
5.2. Arquitectura Sistema de Evaluación de la Calidad Docente	143
5.3. Diagrama de clases para el Proceso de Evaluación Docente	147
5.4. Interfaz para la creación de una encuesta	153
5.5. Creación de una encuesta	154
5.6. Seleccionar encuesta	154
5.7. Encuesta ya creada con anterioridad	155
5.8. Añadir preguntas a la encuesta	156
5.9. Complementando una pregunta de la encuesta	157

5.10. Dominio numérico	157
5.11. Dominio lingüístico	158
5.12. Dominio enumerado	158
5.13. Modificar una pregunta	159
5.14. Encuesta terminada	160
5.15. Resultados de la encuesta	161
5.16. Selección de la encuesta a consultar	162
5.17. Encuesta seleccionada	163
5.18. Valoración calidad docente	164
5.19. Valoración encuesta actual	165
5.20. Valores globales por pregunta	166
5.21. Valores en cada pregunta	167
5.22. Entrada de opiniones	168
5.23. Selección del perfil de la encuesta	169
5.24. Preguntas de la encuesta a rellenar	170
5.25. Respuesta numérica	172
5.26. Respuesta intervalar	173
5.27. Respuesta lingüística	175
5.28. Aviso de encuesta no enviada	176

Introducción

Motivación

La toma de decisiones (TD) es un proceso complejo y una de las actividades fundamentales de los humanos. Algunos autores argumentan que la TD en situaciones complejas es una característica fundamental que diferencia al género humano de los animales [19]. Constantemente nos enfrentamos a situaciones en las que existen varias alternativas y, al menos en algunas ocasiones, tenemos que decidir cuál es mejor, o cuál llevar a cabo.

Un problema clásico de decisión tiene los siguientes elementos básicos:

1. Un conjunto de alternativas o decisiones posibles.
2. Un conjunto de estados de la naturaleza que definen el contexto de definición del problema.
3. Un conjunto de valores de utilidad, cada uno de los cuales está asociado a un par formado por una alternativa y un estado de la naturaleza.
4. Una función que establece las preferencias del experto o decisor sobre los posibles resultados.

En la TD nos podemos encontrar distintas situaciones de decisión dependiendo del contexto del problema [49]:

1. *Ambiente de certidumbre*: La utilidad de cada alternativa se conoce con exactitud y precisión.
2. *Ambiente de riesgo*. El conocimiento sobre las alternativas consiste en sus distribuciones de probabilidad.
3. *Ambiente de incertidumbre*. En esta situación no conocemos la probabilidad de las alternativas. La utilidad de cada una de ellas, se caracteriza de forma aproximada.

La TD se aplica en distintas disciplinas, tales como, las Ciencias Sociales, la Economía, la Ingeniería, la Psicología, etc. Esta amplia gama de campos de aplicación tiene como consecuencia la existencia de diferentes modelos de toma de decisión [22, 25, 35] que han dado lugar a la Teoría de Decisión. La Teoría Clásica de Decisión proporciona gran cantidad de modelos sobre las distintas situaciones enumeradas anteriormente. Sin embargo, los métodos clásicos no son adecuados en situaciones de incertidumbre, es decir, en problemas que presentan información vaga e imprecisa. En estas situaciones hablamos de problemas de decisión en contexto difuso o de toma de decisiones Difusa [27, 34, 100, 136].

Hemos de considerar que existen diferentes modelos de problemas de decisión atendiendo al número de decisores que toman parte en él, pudiendo ir de uno en los problemas de decisión unipersonal como a múltiples decisores en los problemas multi-experto o de decisión en grupo. Igualmente atendiendo a los aspectos o atributos que se valoren sobre cada alternativa podemos hablar de problemas de decisión multi-atributo cuando se valoran diferentes atributos para cada alternativa. Por tanto, en los casos de múltiples expertos y/o múltiples atributos dependiendo

del conocimiento que los expertos tengan sobre las alternativas del problema y de la naturaleza de los atributos a valorar sobre las mismas (cuantitativa o cualitativa), el contexto de definición y el modelo de decisión puede variar. Lo que nos lleva a que en estos tipos problemas de decisión, el contexto de definición del problema puede ser heterogéneo ya que según la naturaleza de los atributos y/o alternativas que se valoran y según el conocimiento de cada uno de los expertos que toman parte en el problema, las preferencias emitidas por los mismos pueden estar valoradas en diferentes dominios de expresión y con información de distinta naturaleza.

Observamos pues, que el Modelado de Preferencias [124] juega un papel clave en la de TD ya que definirá la naturaleza y organización de la información que los expertos utilizan para expresar su conocimiento, gustos, preferencias, etc. En contextos de certidumbre y riesgo en los que suelen valorarse aspectos cuantitativos el uso de información numérica e intervalar es adecuada [84, 109, 153, 157], sin embargo en contextos con incertidumbre, en los que se valoran aspectos cualitativos, el uso del Enfoque Lingüístico Difuso [161] basado en la Teoría de Conjuntos Difusos [46, 160] se ha mostrado útil a la hora de modelar este tipo de preferencias [8, 39, 77, 146]. El uso del enfoque lingüístico implica la necesidad de realizar procesos de operar con palabras, denominados en inglés *Computing with Words* (CW) [143, 145, 164]. Estos procesos se han llevado a cabo en la TD difusa utilizando distintos modelos computacionales [39, 41, 68].

Un esquema de resolución clásico para cualquier problema de *toma de decisiones*, donde intervienen varios expertos y/o criterios, tiene las siguientes fases [123]:

1. *Un proceso de agregación.* En el que se transforma un conjunto de valores de preferencias marginales asociadas a diferentes expertos y/o criterios en un conjunto de valores de preferencia colectiva aplicando un operador de

agregación.

2. *Un proceso de explotación.* A partir de los valores de preferencia colectiva y aplicando un criterio de selección se obtiene un conjunto solución.

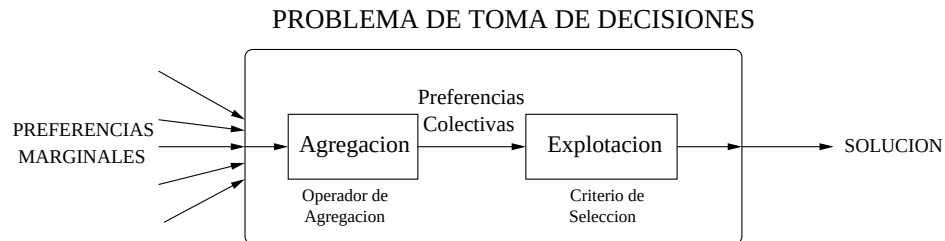


Figura 1: Fases de un Proceso de Toma de Decisión

En la mayoría de los modelos de decisión se fuerza a todos los expertos a expresar sus preferencias en un único dominio de información, sin embargo tal y como hemos señalado antes, en los problemas de TD con múltiples expertos y/o múltiples atributos es común encontrar que las preferencias suministradas por los distintos expertos, definen un contexto heterogéneo de información. Porque los expertos pueden tener un conocimiento preciso sobre algunos atributos, y vago e impreciso sobre otros o la naturaleza de los mismos puede ser diferente (cuantitativa o cualitativa). Esto hace que nos podamos encontrar que la información de entrada al problema esté definida en contextos no homogéneos sobre la que hay que operar (procesos de agregación, ordenación, etc.).

La posibilidad de ofrecer una mayor flexibilidad a la hora de suministrar sus preferencias a los expertos, permitiéndoles utilizar distintos dominios de expresión, para información de distinta naturaleza nos permitirá obtener información más ajustada al conocimiento de los expertos y mejores resultados.

En la literatura se han presentado distintos modelos, operadores y procesos de

decisión que han permitido abordar problemas de TD definidos en algunos de estos contextos heterogéneos [28, 40, 52, 63, 71, 104, 166]. Sin embargo siguen existiendo algunas limitaciones para abordar problemas de decisión definidos en contextos heterogéneos compuestos por información numérica, lingüística e intervalar:

- *Herramientas, modelos y operadores:* en la literatura existen algunos operadores, modelos y herramientas como, modelado lingüístico difuso junto con la aritmética difusa [46, 161] y modelado lingüístico basado en 2-tuplas [68] entre otros, que nos facilitan el operar y manejar información definida en algunos contextos no homogéneos:

1. *Numérico y Lingüístico* [40, 69, 104].
2. *Lingüístico Multi-granular* [28, 71, 139].
3. *Numérico e Intervalar* [37, 93, 135].

Sin embargo, la capacidad de manejar contextos de información con información modelada en todos los dominios anteriores no ha sido abordada de una manera formal.

- *Complejidad de los resultados:* El tratamiento de información no homogénea que además involucra incertidumbre, ha implicado el uso de la aritmética difusa [46] para operar con dicha información. Dando lugar a resultados que son expresados mediante números difusos que son difíciles de manejar y entender por los participantes en dichos procesos de decisión.

Debido a estas razones y a la necesidad de enfrentarse a problemas de TD definidos en contextos heterogéneos, parece lógico y adecuado desarrollar herramientas, modelos y operadores que sean capaces de resolver problemas de TD, tales que, la información pueda estar modelada en dominios numéricos, intervalares y

lingüísticos, a la vez que los resultados que produzcan dichos operadores, modelos y herramientas sean fácilmente inteligibles para los participantes involucrados en dichos procesos de decisión.

También hemos de destacar que durante el proceso de estudio e investigación de problemas y procesos de decisión, vimos cómo en la literatura se han utilizado con éxito la resolución de procesos de evaluación basados en modelos de decisión [31, 32, 103, 104, 105], por tanto, vimos la necesidad, utilidad y posibilidad de utilizar las herramientas, modelos y operadores que nos habíamos planteado inicialmente para problemas de TD en contextos heterogéneos a procesos de evaluación en los que también pueden encontrarse fácilmente contextos con información de distinta naturaleza. Habiéndose obtenido resultados de interés que son presentados durante el desarrollo de esta memoria de investigación.

Objetivos

Esta memoria plantea las siguientes propuestas como objetivos a desarrollar a lo largo de la misma:

- *Desarrollo de herramientas, modelos y operadores que nos permitan ofrecer una mayor flexibilidad a los expertos a la hora de expresar sus preferencias mediante información numérica, intervalar y lingüística en procesos de decisión.* Además de subsanar las limitaciones para el tratamiento de este tipo de contextos y que produzcan resultados que sean fácilmente comprensibles por los decisores que toman parte en el proceso de decisión y/o evaluación.
- *Análisis, Diseño e Implementación de un sistema de apoyo a la decisión en lenguaje JAVA,* que nos permita definir distintos tipos de modelos de decisión, ya sea desde un punto de vista para problemas de evaluación o de toma

de decisiones. A la vez que, los contextos de definición de los mismos puedan modelar la información en cualquier dominio de expresión y que ofrezca resultados fáciles de entender.

Resumen

La presente memoria se encuentra estructurada en 5 capítulos. A continuación presentamos un breve resumen de cada uno de los mismos:

- En el Capítulo 1, se presenta una breve revisión de la toma de decisiones y sus distintos modelos. A continuación se hará una revisión del modelado de preferencias atendiendo tanto a las estructuras de representación como a los modelos de representación de la información utilizados habitualmente en los problemas de decisión. En la revisión del modelado de representación de la información se hará un especial hincapié en el modelo de representación lingüístico basado en 2-tuplas [68], debido a la importancia del mismo en las propuestas posteriores que se presentan en esta memoria.
- En el Capítulo 2, se presentan una serie de operadores y herramientas difusas que son necesarias para el tratamiento de información expresada en contextos heterogéneos con información numérica, intervalar y lingüística necesarias para la propuesta que hacemos de un modelo de decisión que se presenta en el siguiente capítulo.
- En el Capítulo 3, se propone un modelo de decisión para problemas de TD definidos en contextos heterogéneos con información numérica, intervalar y lingüística y que proporciona los resultados en un contexto cualitativo y simbólico, fácil de comprender para los expertos que toman parte en el problema de TD.

- En el Capítulo 4, se realiza el diseño, por medio de diagramas UML, para un sistema de soporte a la decisión para problemas de TD en general o sistemas de evaluación. Para ello se diseñará una biblioteca para el manejo de la información heterogénea. Además, se diseñará también una biblioteca para el modelo de decisión para problemas de TD en contextos heterogéneos del capítulo anterior.
- En el Capítulo 5, se presenta el diseño y la implementación de un prototipo inicial para un sistema de evaluación de la calidad docente en las Universidades Andaluzas, que ha sido desarrollado para la Unidad de Calidad de las Universidades Andaluzas (UCUA), actualmente Agencia Andaluza de Evaluación (AGAE), bajo la cobertura de un proyecto de la Convocatoria de Grupos Específicos y Análisis de la Calidad de las Universidades Andaluzas. En dicho sistema se han utilizado todos los resultados presentados en los capítulos anteriores para implementar un sistema de evaluación basado en un modelo de decisión para información heterogénea que *facilite* la realización de encuestas que se realizan en las Universidades Andaluzas para medir la calidad de sus docentes. El concepto de facilitar la encuesta anterior se interpreta desde distintos puntos de vista:
 1. *Tecnológico*: el sistema permitirá automatizar la realización de las encuestas.
 2. *Cognitivo*: al utilizarse los modelos presentados en los capítulos iniciales de la memoria, el sistema permitirá modelar la información de forma natural y no forzar a los encuestados a expresar sus opiniones en dominios de información que están lejanos a lo que es su conocimiento sobre la encuesta.
 3. *Gestión*: los resultados proporcionados por este sistema son más fáciles

de comprender e interpretar en el entorno propio de la encuesta que los que se venían utilizando habitualmente.

- Finalmente, se presentan las conclusiones más relevantes obtenidas a lo largo de la memoria y una breve descripción de las líneas futuras de desarrollo a partir de la misma. La memoria finaliza con una recopilación bibliográfica que recoge las contribuciones más destacadas en la materia estudiada.

Capítulo 1

Problemas de Toma de Decisión en Ambiente Difuso: Modelado de Preferencias y el Enfoque Lingüístico Difuso

En este capítulo se hace una revisión de los problemas de toma de decisión definidos en un contexto difuso, es decir, problemas de decisión en los que participan varios individuos y en los que se trabaja con información vaga e imprecisa.

En primer lugar describiremos las características de los problemas de toma de decisión. A continuación se abordará el concepto de modelado de preferencias, haciendo una distinción entre las estructuras que utilizan para expresar y representar las preferencias y el dominio de la información en el que los expertos las expresan. En tercer lugar se hará una breve revisión de nociones y conceptos básicos de la Teoría de Conjuntos Difusos que nos va a servir para introducir los conceptos más

importantes relacionados con el Enfoque Lingüístico Difuso.

1.1. El Problema de la Toma de Decisión

En un sentido amplio, tomar una decisión consiste en elegir la mejor opción o alternativa de entre un conjunto de alternativas posibles. Muchas de las actividades humanas precisan en algún momento tomar decisiones. Diariamente nos enfrentamos a situaciones en las que debemos decidir qué hacer o qué alternativa tomar en función del entorno en el que nos encontramos. Por citar un ejemplo cercano a todos nosotros, la elección de qué carrera universitaria estudiar supuso en su día un problema de toma de decisión que seguro nos obligó a sopesar, analizar y comparar las distintas alternativas con el propósito de elegir la más adecuada.

Cada vez que se plantea la necesidad de tomar una decisión, ésta va acompañada de un conjunto de posibles alternativas que a su vez tienen una serie de consecuencias que pueden hacernos dudar sobre la idoneidad de cada una de ellas. La incertidumbre suele ser una compañera presente en los procesos de toma de decisión que produce malestar e inseguridad a los individuos que deben tomar las decisiones.

La toma de decisiones, como apuntan Keeney y Raiffa [90], intenta ayudar a los individuos a tomar decisiones difíciles y complejas de una forma racional. Esta racionalidad implica el desarrollo de métodos y/o modelos que permitan representar fielmente cada problema y analizar las distintas alternativas con criterios objetivos. Partiendo de disciplinas clásicas como la Estadística, la Economía y la Matemática, a las que se les han unido otras más recientes como la Inteligencia Artificial, se han desarrollado teorías y modelos en el campo de la toma de decisiones que han permitido estructurar de una forma lógica y racional el proceso de toma de decisión y facilitar esta tarea a los individuos encargados de llevarla a cabo.

Los problemas clásicos de decisión presentan los siguientes elementos básicos [33]:

1. Uno o varios objetivos por resolver.
2. Un conjunto de alternativas o decisiones posibles para alcanzar dichos objetivos.
3. Un conjunto de factores o estados de la naturaleza que definen el contexto en el que se plantea el problema de decisión.
4. Un conjunto de valores de utilidad o consecuencias asociados a los pares formados por cada alternativa y estado de la naturaleza.

Ante la gran variedad de situaciones o problemas de decisión que se pueden presentar en la vida real, la Teoría de la Decisión ha establecido una serie de criterios que permiten clasificar los problemas atendiendo a diferentes puntos de vista:

1. Según el número de criterios o atributos que se han de valorar en la toma de decisión.
2. Según el ambiente de decisión en el que se han de tomar las decisiones.
3. Según el número de expertos que participan en el proceso de decisión.

En los siguientes apartados se hace una breve revisión de las características que definen cada uno de estos puntos de vista.

1.1.1. Según el Número de Criterios

El número de criterios (también denominados atributos) que se tienen en cuenta en los procesos de decisión para obtener la solución también permite clasificar a los problemas de decisión en dos tipos [25, 34, 56, 81, 100, 121, 141]:

1. Problemas con un sólo criterio o atributo. Problemas de decisión en los que para evaluar las alternativas se tiene en cuenta un único valor que representa la valoración dada a esa alternativa. La solución se obtiene como la alternativa que mejor resuelve el problema teniendo en cuenta la valoración de dicha alternativa como único criterio de decisión.
2. Problemas multicriterio o multiatributo. Problemas de decisión en los que para evaluar las alternativas se tienen en cuenta los valores de dos o más criterios o atributos que definen las características de cada alternativa. La alternativa solución será aquella que mejor resuelva el problema considerando todos estos criterios o atributos.

Ambos tipos se pueden diferenciar perfectamente con el siguiente ejemplo. Supongamos un problema de decisión en el que nos planteamos cambiar de trabajo y nos ofrecen tres posibles alternativas, cada una de ellas caracterizada por tres atributos como son el sueldo, la ubicación geográfica y tipo de trabajo a desarrollar. Este problema puede ser muy simple si para tomar la decisión consideramos como único criterio de decisión elegir la alternativa con mejor sueldo. Sin embargo, este mismo problema se complicaría y el proceso para resolverlo sería diferente si además de considerar el sueldo también tuviésemos en cuenta el tipo de trabajo y/o la ubicación geográfica del mismo. En este segundo caso estaríamos ante un problema en el que hemos de considerar varios atributos o criterios antes de tomar una decisión y por lo tanto, estaríamos hablando de un problema de decisión multicriterio o multiatributo.

Los problemas de toma de decisión multicriterio son más complejos de resolver que los problemas en los que sólo hay que tener en cuenta un criterio para obtener la solución. Cada criterio puede establecer un orden de preferencia particular y diferente sobre el conjunto de alternativas. A partir del conjunto de órdenes de

preferencia particulares será necesario establecer algún mecanismo que permita construir un orden global de preferencia.

El número de criterios en problemas de decisión multicriterio se asume que es finito. Sean $X = \{x_1, x_2, \dots, x_n\}$ y $C = \{c_1, c_2, \dots, c_m\}$ el conjunto de alternativas y el conjunto de criterios que caracterizan una situación de decisión determinada. Entonces, una forma de representación de la información del problema puede expresarse mediante la siguiente tabla:

Alternativas	<i>Criterios</i>			
(x_i)	c_1	c_2	$\dots$	c_m
x_1	y_{11}	y_{12}	$\dots$	y_{1m}
$\dots$	$\dots$	$\dots$	$\dots$	$\dots$
x_n	y_{n1}	y_{n2}	$\dots$	y_{nm}

Cuadro 1.1: Esquema general de un problema de toma de decisiones multicriterio

Cada entrada de la tabla y_{ij} indica la preferencia de la alternativa x_i respecto del criterio c_j . Según el contexto de definición del problema, cada y_{ij} podrá estar valorado en un dominio de expresión de preferencias determinado (numérico, lingüístico,...).

1.1.2. Según el Ambiente de Decisión

El ambiente de decisión viene definido por las características y contexto en el que se va a llevar a cabo la toma de decisiones. La Teoría Clásica de la Decisión distingue tres situaciones o ambientes de decisión [49, 119]:

1. Ambiente de certidumbre.

Un problema de decisión está definido en un ambiente de certidumbre cuan-

do son conocidos con exactitud todos los elementos y/o factores que intervienen en el problema. Esta situación permite asignar valores precisos de utilidad a cada una de las alternativas presentes en el problema.

Como ejemplo, supongamos que disponemos de una determinada cantidad de dinero que queremos invertir en alguno de los diferentes productos financieros del mercado que nos garantice la inversión realizada (ej., imposición a plazo fijo). Asumiendo que conocemos con exactitud la rentabilidad de cada producto, los gastos de gestión, la duración del mismo, deberemos decidir en que producto invertir para maximizar la inversión realizada. En este caso conocemos todos los factores que se han de tener en cuenta para la toma de decisión y el problema consistirá en estructurar correctamente esta información y establecer las preferencias entre las alternativas de forma que nos permita elegir aquella que maximice el beneficio esperado.

2. Ambiente de riesgo.

Un problema de decisión está definido en un ambiente de riesgo cuando alguno de los elementos o factores que intervienen están sujetos a las leyes del azar. En estos casos estos problemas son resueltos utilizando la Teoría de la Probabilidad.

Por ejemplo, si queremos realizar una inversión en un deposito ligado a resultados deportivos, inmediatamente surgen dudas sobre los resultados de equipos de los que depende la inversión. En este caso el enfoque del problema ha de ser diferente y se podrá utilizar una distribución de probabilidad para reflejar la posible subida o bajada dependiendo de los resultados deportivos que influirá en la utilidad de cada una de las posibles alternativas en las que invertir el dinero.

3. Ambiente de Incertidumbre.

Un problema de decisión está definido en un ambiente de incertidumbre cuando la información disponible sobre las distintas alternativas puede ser incompleta, vaga o imprecisa, lo que implica que la utilidad asignada a cada alternativa tenga que ser valorada de forma aproximada. Esta incertidumbre surge a raíz del intento de modelar la imprecisión propia del comportamiento humano o la inherente a ciertos fenómenos que por su naturaleza son inciertos.

Los métodos clásicos no son adecuados para tratar situaciones en las que la incertidumbre se debe a la aparición de información vaga e imprecisa, como por ejemplo si queremos realizar una inversión sobre un nuevo valor en bolsa, donde los expertos intentan valorar fenómenos relacionados con apreciaciones subjetivas sobre una posible subida o bajada en la cotización de las acciones en las que se invierta el dinero. Esto ha generado la necesidad de recurrir a la definición de nuevos modelos basados en la Teoría de los Conjuntos Difusos [160] para modelar la incertidumbre como pueden ser los Rough Sets [48, 58, 78], Conjuntos Difusos Intuicionistas [7, 23, 24], etc.

1.1.3. Según el Número de Expertos

Finalmente, otro punto de vista a la hora de clasificar los problemas de decisión hace referencia al número de expertos o fuentes de información que toman parte en el proceso. Un proceso de toma de decisión en el que participan varios expertos es más complejo que otro en el que la toma de decisión se realiza de forma individual. Sin embargo, el hecho de que intervengan varios expertos con puntos de vista diferentes puede ofrecer una solución más satisfactoria al problema.

Atendiendo al número de expertos o fuentes de información que toman parte

en la toma de decisión [8, 87, 144, 149], los problemas de decisión se pueden clasificar en dos tipos:

1. Unipersonales o individuales. Las decisiones son tomadas por un sólo experto.
2. En Grupo o Multiexperto. Las decisiones son tomadas en conjunto por un grupo de expertos que intentan alcanzar una solución en común al problema.

Los problemas abordados en esta memoria pertenecen al segundo tipo.

1.2. Modelado de Preferencias

El modelado de preferencias es una de las actividades inevitables en los problemas de toma de decisión [54, 114, 117, 124], independientemente del área en el que se esté trabajando (Economía [5, 38], Psicología [26, 36, 88], etc.). Los expertos en base a su conocimiento, experiencias y creencias han de emitir sus valoraciones sobre el conjunto de alternativas y establecer un orden de preferencia sobre la idoneidad de cada una ellas como solución al problema.

En los problemas de decisión los expertos utilizan modelos de representación de preferencias que les resulten cercanos a sus disciplinas o campos de trabajo. Por ejemplo, expertos que pertenecen a áreas técnicas se pueden sentir muy cómodos representando sus preferencias mediante valores numéricos. Sin embargo, expertos que pertenecen a otro tipo de disciplinas menos técnicas como pueden ser las pertenecientes a áreas sociales (Psicología, Sociología, ...), pueden preferir expresar sus preferencias utilizando expresiones más cercanas al lenguaje humano tales como palabras o términos lingüísticos. Para manejar este tipo de valoraciones se han definido diferentes mecanismos que permiten transformar las preferencias de

los expertos en representaciones formales que admiten un tratamiento matemático, racional y consistente de dicha información.

El modelado de preferencias es un área de trabajo dentro de la toma de decisión dedicada a la representación de las preferencias de los expertos [45, 114, 124]. Para hacer una breve revisión sobre la misma, vamos a considerar dos puntos de vista claramente diferenciados pero igualmente trascendentales, como son:

- a) *La estructura de información* utilizada por los expertos para la representación de sus preferencias.
- b) *El dominio de la información* en el que se expresan las preferencias sobre el conjunto de alternativas al problema.

Ambos puntos de vista presentan diferentes variantes que dependerán del problema que se esté tratando.

1.2.1. Estructuras para la Representación de Preferencias

En esta sección haremos un repaso de las estructuras de información más utilizadas en la literatura para la representación de las preferencias de los expertos [74, 111, 134]:

- Vectores de Utilidad
- Órdenes de Preferencia
- Relaciones de Preferencia

Cada una de ellas se representa y se interpreta de forma diferente tal y como se recoge en los siguientes apartados.

1.2.1.1. Vectores de Utilidad

Los vectores de utilidad han sido una estructura de representación de información ampliamente utilizada en la literatura clásica para representar las preferencias de los expertos [44, 99, 103, 134]. Es una estructura muy simple basada en un vector donde cada elemento se interpreta como la preferencia o utilidad de una de las alternativas del problema.

Ejemplo

Sea $E = \{e_1, \dots, e_m\}$ ($m \geq 2$) un conjunto finito de expertos que han de expresar sus preferencias sobre un conjunto finito de alternativas $X = \{x_1, x_2, \dots, x_n\}$ ($n \geq 2$). Las preferencias dadas por los expertos sobre el conjunto de alternativas X utilizando vectores de utilidades U^i serían la siguientes:

$$U^i = \{u_1^i, \dots, u_n^i\},$$

donde u_j^i representa la utilidad o valoración dada por el experto i a la alternativa j . Se asume que cuanto mayor sea el valor de u_j^i , más satisface la alternativa j el objetivo del problema según la opinión del experto i .

En la Sección 1.2.2 donde se presentarán los diferentes dominios para el modelado de preferencias, aparecen más ejemplos de representación de preferencias utilizando vectores de utilidad.

1.2.1.2. Órdenes de Preferencia

En esta estructura de representación de preferencias se establece un ranking u orden de alternativas que representa la idoneidad de cada alternativa como solución al problema según el punto de vista de cada experto [111, 130, 133].

Un orden de preferencia O^i representa un orden dado por el experto i sobre el conjunto de alternativas X atendiendo a sus preferencias. Se representa mediante

un vector ordenado decreciente del conjunto de alternativas,

$$O^i = \{o^i(1), \dots, o^i(n)\}$$

Para todo orden de preferencia O^i , suponemos sin pérdida de generalidad que cuanto menor es la posición de una alternativa en dicho orden, esta alternativa es más preferida que el resto para resolver el problema según la opinión del experto i .

Ejemplo

Sea $E = \{e_1, \dots, e_m\}$ ($m \geq 2$) un conjunto finito de expertos que han de expresar sus preferencias sobre un conjunto finito de alternativas $X = \{x_1, x_2, \dots, x_n\}$ ($n = 4$). Las preferencias dadas por los expertos 1 y 2 sobre el conjunto de alternativas X utilizando órdenes de preferencia podrían ser las siguientes:

$$O^1 = \{x_3, x_2, x_1, x_4\}$$

$$O^2 = \{x_2, x_3, x_1, x_4\}$$

En este ejemplo, el experto 1 considera que la mejor alternativa para resolver el problema es x_3 y la peor x_4 . Sin embargo, para el experto 2 la mejor alternativa es x_2 y la peor también es x_4 .

1.2.1.3. Relaciones de Preferencia

En el Modelado de Preferencias [124], las preferencias sobre un conjunto de alternativas $X = \{x_1, \dots, x_n\}$ se pueden modelar como relaciones binarias entre pares de alternativas $x_l R x_k$ ($x_l, x_k \in X$), que se interpretan como la intensidad o el grado de preferencia de la alternativa x_l sobre la alternativa x_k .

Cuando se trabaja con conjuntos de alternativas finitos, una estructura de información capaz de soportar este tipo de relaciones binarias entre alternativas son las relaciones de preferencia.

Una relación de preferencia individual se representa como una matriz $P_{e_i} \subset X \times X$, donde el valor $\mu_{P_{e_i}}(x_l, x_k) = p_i^{lk}$ representa el grado de preferencia de la alternativa x_l sobre la alternativa x_k [96, 134, 159],

$$P_{e_i} = \begin{pmatrix} p_i^{11} & \dots & p_i^{1n} \\ \vdots & \ddots & \vdots \\ p_i^{n1} & \dots & p_i^{nn} \end{pmatrix}$$

Clásicamente, en los problemas de TDG, los expertos expresan sus preferencias sobre el conjunto de alternativas X utilizando relaciones de preferencia valoradas numéricamente en $[0, 1]$ [30, 53, 86, 150, 151],

$$P_{e_i}: X \times X \rightarrow [0, 1]$$

Partiendo de esta representación, de los valores p_i^{lk} se debe destacar que:

- $p_i^{lk} = 1/2$, significa que hay indiferencia sobre la preferencia entre ambas alternativas.
- $p_i^{lk} \geq 1/2$, significa que la alternativa x_l es preferida sobre la x_k .
- $p_i^{lk} = 1$, significa que la alternativa x_l es totalmente preferida sobre la x_k .

En los problemas de decisión es importante que las opiniones de los expertos sean consistentes. Para garantizar esta consistencia a las relaciones de preferencia se les puede requerir que satisfagan algunas de las siguientes propiedades [73, 128]:

- Reciprocidad: $p_i^{lk} + p_i^{kl} = 1, \forall l, k = 1, \dots, n$
- Completitud: $p_i^{lk} + p_i^{kl} \geq 1, \forall l, k = 1, \dots, n$

- Transitividad max-min: $p_i^{lk} \geq \min(p_i^{lj}, p_i^{jk}), \forall l, j, k = 1, \dots, n$
- Transitividad max-max: $p_i^{lk} \geq \max(p_i^{lj}, p_i^{jk}), \forall l, j, k = 1, \dots, n$
- Transitividad max-min restrictiva: $p_i^{lj} \geq 0.5, p_i^{jk} \geq 0.5 \Rightarrow p_i^{jk} \geq \min(p_i^{lj}, p_i^{jk}), \forall l, j, k = 1, \dots, n$
- Transitividad max-max restrictiva: $p_i^{lj} \geq 0.5, p_i^{jk} \geq 0.5 \Rightarrow p_i^{jk} \geq \max(p_i^{lj}, p_i^{jk}), \forall l, j, k = 1, \dots, n$
- Transitividad aditiva: $p_i^{lj} + p_i^{jk} - 0.5 = p_i^{lk}, \forall l, j, k = 1, \dots, n$

En las relaciones de preferencia es habitual no definir los elementos de la diagonal principal o en el caso de hacerlo asignarles el valor $p_i^{ll} = 1/2$. Esto se debe a que no es útil comparar una alternativa consigo misma.

Ejemplo

Sea $E = \{e_1, \dots, e_m\}$ ($m \geq 2$) un conjunto finito de expertos que han de expresar sus preferencias sobre un conjunto finito de alternativas $X = \{x_1, x_2, \dots, x_n\}$ ($n = 4$). Las preferencias dadas por el experto 1 sobre el conjunto de alternativas X definido en un dominio numérico en $[0, 1]$ utilizando una relación de preferencia difusa P_{e_1} tendría el siguiente aspecto:

$$P_{e_1} = \begin{pmatrix} - & 0.3 & 0.7 & 0 \\ 0.7 & - & 0.6 & 0.6 \\ 0.3 & 0.4 & - & 0.2 \\ 1 & 0.4 & 0.8 & - \end{pmatrix}$$

Finalmente destacar que las relaciones de preferencia han sido utilizadas satisfactoriamente por muchos autores para resolver problemas decisión en grupo

[51, 67, 73, 83, 85, 146, 148], siendo también utilizadas en alguna de las propuestas de esta memoria para representar las preferencias de los expertos.

1.2.2. Dominios de Expresión de Preferencias

En problemas de decisión entendemos por dominio de expresión de preferencias el dominio de información utilizado por los expertos para expresar sus preferencias.

En la literatura encontramos que en la mayoría de los problemas de toma de decisión los expertos expresan sus preferencias en el mismo dominio de información, hablándose de problemas definidos en contextos homogéneos [4, 8, 14, 27, 41, 50, 59, 64, 95, 98, 102, 118, 147], y algunos problemas en los que los expertos utilizan dominios de información diferentes, conocidos como problemas definidos en contextos heterogéneos [40, 52, 71, 73, 104, 105, 166].

En problemas de TDG, la elección de un dominio de información para expresar las preferencias puede deberse a varios motivos:

- a) Expertos con diferente grado de conocimiento sobre el problema. La experiencia de los expertos en la resolución de problemas similares puede implicar que unos expertos opten por elegir dominios de expresión de preferencias precisos como valores numéricos exactos (0, 1, 100, 2500, ...) frente a otros expertos con menos experiencia y que se sientan más cómodos utilizando dominios más flexibles como los intervalos.
- b) Pertenencia de los expertos a diferentes áreas de conocimiento. Siempre que sea posible, cada experto tenderá a utilizar un dominio de información que le resulte cercano al tipo de información con el que esté acostumbrado a trabajar en su respectiva área de trabajo. Así, expertos pertenecientes a áreas técnicas se sentirán cómodos utilizando valoraciones numéricas mientras que

aquellos pertenecientes a áreas sociales pueden preferir utilizar otro tipo de valoraciones no numéricas como las lingüísticas.

- c) Naturaleza cuantitativa o cualitativa de la información con la que se esté trabajando. La naturaleza del atributo que se esté evaluando puede condicionar el dominio utilizado para su valoración. Atributos de naturaleza cuantitativa admiten mucho mejor valoraciones de tipo numérico que aquellos otros de naturaleza cualitativa en los que al tratarse por ejemplo sensaciones o percepciones de los expertos, el uso de otro tipo de valoraciones como palabras o términos lingüísticos (“bueno”, “malo”, “mejor”, ...) suele ser mucho más apropiado [29, 59, 70, 82, 97, 103, 152].

Adaptar el modelado de preferencias al contexto en el que se desarrolla el problema de decisión consigue que los expertos se sientan más seguros a la hora de valorar sus preferencias y por lo tanto que la solución final tenga mayor garantía de éxito [45].

En la literatura [3, 40, 50, 73, 92, 166] encontramos que los expertos utilizan principalmente tres tipos de dominios de información para expresar sus preferencias:

- Dominio Numérico
- Dominio Intervalar
- Dominio Lingüístico

Ejemplos, características y una breve justificación de las circunstancias en las que es adecuado utilizar un dominio u otro se presentan en los siguientes apartados.

1.2.2.1. Dominio Numérico

El uso del dominio numérico para modelar las preferencias implica que los expertos expresen sus preferencias mediante valores numéricos. Dentro del área de investigación en la que estamos trabajando, sobre este dominio podemos encontrar dos variantes:

a) Numérico Binario. Como su propio nombre indica, el modelado numérico binario se caracteriza por utilizar exclusivamente dos valores para cuantificar la utilidad de cada alternativa. Normalmente se utilizan los valores $\{0, 1\}$, donde el 0 representa una valoración negativa de la alternativa y el 1 representa una valoración positiva.

Ejemplo

Sea $E = \{e_1, \dots, e_m\}$ ($m \geq 2$) un conjunto finito de expertos que han de expresar sus preferencias sobre un conjunto finito de alternativas $X = \{x_1, \dots, x_n\}$ ($n = 4$). Cada experto i utiliza un vector de utilidades $U^i = \{u_1^i, \dots, u_n^i\}$ para expresar sus preferencias, donde a cada alternativa de X se le asigna un valor de utilidad del dominio binario, $u_j^i \in \{0, 1\}$. Los valores dados por e_1 y e_2 pueden ser los siguientes:

$$U^1 = \{1, 0, 0, 1\}$$

$$U^2 = \{0, 0, 1, 0\}$$

donde las alternativas x_1 y x_4 son valoradas positivamente por el experto 1 y las alternativas x_1, x_2 y x_4 reciben una valoración negativa por parte del experto 2.

El modelado de preferencias utilizando dominios binarios ha formado parte de la visión clásica y “crisp” de la toma de decisión en la que los expertos sólo podían indicar si una alternativa era considerada como buena o mala para resolver el problema, no teniendo la posibilidad de introducir cierta incertidumbre sobre la utilidad o bondad de cada alternativa como solución al problema. Esta forma

tan nítida o “crisp” de expresión de preferencias ha ido cambiando con el tiempo y ha tendido hacia el uso de dominios menos restrictivos que permiten reflejar la incertidumbre presente en los problemas de decisión.

b) Numérico normalizado en el intervalo $[0, 1]$. Los expertos utilizan un valor numérico dentro del intervalo $[0, 1]$ para modelar la preferencia sobre cada alternativa [53, 96]. A diferencia del dominio anterior donde sólo se admiten dos posibles valores, ahora se pueden utilizar valores reales dentro de este intervalo que permiten establecer un orden de preferencia entre las distintas alternativas en función de la utilidad asignada a cada una de ellas.

Ejemplo

Sea $E = \{e_1, \dots, e_m\}$ ($m \geq 2$) un conjunto finito de expertos que han de expresar sus preferencias sobre un conjunto finito de alternativas $X = \{x_1, \dots, x_n\}$ ($n = 4$). Un ejemplo de preferencias dadas por los expertos 1 y 2 sobre el conjunto de alternativas X utilizando un dominio de expresión dentro del intervalo $[0, 1]$ y vectores de utilidad podría ser el siguiente:

$$U^1 = \{1, 0.2, 0, 0.6\}$$

$$U^2 = \{0, 0.4, 0.7, 0.9\}$$

Podemos comprobar como el experto 1 considera que la alternativa x_1 es la mejor y le asigna una utilidad máxima. Por el contrario, considera la alternativa x_3 peor que la x_2 asignándole una utilidad de 0 y 0.2 respectivamente. Para el experto 2, la mejor alternativa sería x_4 y la peor la x_1 .

1.2.2.2. Dominio Intervalar

El hecho de considerar la incertidumbre en los problemas de decisión ha originado la necesidad de definir modelados de preferencias más flexibles capaces de

recoger dicha incertidumbre, siendo el modelado intervalar uno de ellos. La valoración de alternativas por medio de intervalos $[\underline{a}, \bar{a}]$ ($\underline{a} \leq \bar{a}$) se ha mostrado como una técnica eficaz para tratar la incertidumbre en ciertos problemas de decisión [2, 92, 135]. Los expertos han de valorar alternativas sobre las que no tienen un conocimiento lo suficientemente preciso para asignarles valores exactos mediante un valor numérico. En estos casos la utilización del modelado de preferencias mediante valores intervalares hace que los expertos se sientan más seguros en sus valoraciones y que los resultados de problema aunque no sean exactos estén delimitados.

Un ejemplo de problema de decisión donde puede apreciarse como el uso de intervalos resuelve el problema de la incertidumbre podría ser el siguiente. Supongamos que los expertos han de opinar sobre el consumo de combustible de diferentes vehículos con el propósito de elegir el vehículo de menor consumo. Sería arriesgado e inapropiado que estimasen el consumo de combustible mediante un valor numérico exacto debido a que si se llevase a cabo una prueba de consumo real, el resultado también dependería de elementos externos no controlados por los expertos (viento, pericia del conductor, etc) que añadirían incertidumbre y condicionarían el resultado de la misma. En este caso el resultado del problema sería más fiable si los expertos utilizasen intervalos $[\underline{a}, \bar{a}]$ (ej., entre [5.0,6.0] litros/100km) para expresar sus preferencias en lugar de utilizar valores numéricos precisos que no reflejarían la incertidumbre del problema.

En los problemas consultados en la literatura donde se utilizan intervalos [73, 93], los expertos expresan sus preferencias mediante intervalos del tipo $[\underline{a}, \bar{a}] \in [0, 1]$. En el caso de que no estuviesen definidos dentro de este rango tan sólo habría que aplicar un proceso de normalización en $[0, 1]$.

Ejemplo

Sea $E = \{e_1, \dots, e_m\}$ ($m \geq 2$) un conjunto finito de expertos que han de ex-

presar sus preferencias sobre un conjunto finito de alternativas $X = \{x_1, \dots, x_n\}$ ($n = 4$). Los expertos 1 y 2 pueden expresar sus preferencias sobre el conjunto de alternativas X utilizando un dominio de expresión intervalar en $[0, 1]$ y vectores de utilidad como:

$$U^1 = \{[0.5, 0.7], [0.2, 0.5], [0, 0.2], [0.7, 1]\}$$

$$U^2 = \{[0, 0.3], [0.3, 0.7], [0.7, 0.8], [0.8, 1]\}$$

En este ejemplo, la alternativa mejor valorada por el experto 1 es la numero 4, pero debido a la existencia de incertidumbre, el experto ha preferido utilizar un intervalo de utilidad $[0.7, 1]$ en lugar de un valor numérico preciso.

1.2.2.3. Dominio Lingüístico

Los expertos pueden utilizar un modelado de preferencias lingüístico [55, 59, 132, 142, 161, 163] en aquellas situaciones de decisión en las que la información disponible es demasiado imprecisa o se valoran aspectos cuya naturaleza recomienda el uso de valoraciones cualitativas. En estas situaciones, el experto puede considerar más adecuado utilizar una palabra o término lingüístico para expresar sus preferencias que un valor numérico más o menos preciso.

Son muchas las ocasiones en las que los expertos se sienten más cómodos utilizando este tipo de dominios lingüísticos, sobre todo si han de valorar aspectos relacionados con percepciones humanas muchas veces expresadas de forma imprecisa y donde es habitual utilizar palabras del lenguaje natural en lugar de números. Como ejemplo podemos citar el propuesto en [97] para valorar el nivel de confort de un vehículo. En este caso concreto, los expertos pueden preferir utilizar palabras como “malo”, “bueno”, “aceptable” para expresar su opinión sobre el grado de confort de un vehículo en lugar de utilizar valores numéricos.

Ejemplo

Sea $E = \{e_1, \dots, e_m\}$ ($m \geq 2$) un conjunto finito de expertos que han de expresar sus preferencias sobre un conjunto finito de alternativas $X = \{x_1, \dots, x_n\}$ ($n = 4$). Sea $S = \{muy_malo, malo, normal, bueno, muy_bueno\}$ el conjunto de términos o etiquetas lingüísticas utilizadas por los expertos para expresar sus preferencias sobre el conjunto de alternativas X . Las preferencias dadas por los expertos 1 y 2 utilizando vectores de utilidad podrían ser las siguientes:

$$U^1 = \{muy_malo, bueno, malo, muy_bueno\}$$

$$U^2 = \{muy_bueno, malo, muy_malo, normal\}$$

En este ejemplo, según la opinión del experto 1, la alternativa mejor valorada es x_4 y la peor valorada x_1 . Por el contrario, el experto 2 considera que la mejor alternativa es x_1 y la peor x_3 .

Dentro de la toma de decisión Difusa, el Enfoque Lingüístico Difuso [161] ha sido la disciplina encargada de modelar las preferencias de los expertos que utilizan valoraciones lingüísticas para expresar sus preferencias [1, 3, 4, 8, 15, 41, 100, 146, 151].

Debido a la importancia del Enfoque Lingüístico Difuso en nuestras investigaciones se abordará en profundidad dentro de este capítulo en la Sección 1.4.

1.3. Nociones y Conceptos Básicos de la Teoría de Conjuntos Difusos

La Teoría de los Conjuntos Difusos propuesta por L. Zadeh en la década de los 60 [160] tuvo por objeto modelar aquellos problemas donde los enfoques clásicos

resultaban insuficientes o no operativos. Dicha teoría generaliza la noción clásica de conjunto e introdujo el concepto de “conjunto difuso” como aquel conjunto cuya frontera no es precisa. Los conjuntos difusos surgen como una nueva forma de representar la imprecisión y la incertidumbre [91, 167] diferente al tratamiento tradicional llevado a cabo por la Teoría Clásica de Conjuntos y Teoría de la Probabilidad. A lo largo de las cinco décadas de existencia de la Teoría de Conjuntos Difusos, gran cantidad de investigadores le han prestado atención en sus investigaciones y la han aplicado en dos vertientes principales [116]:

1. Como una teoría matemática formal [76, 108], ampliando conceptos e ideas de otras áreas de la matemática como el Álgebra, la Teoría de Grafos, la Topología, etc., al aplicar conceptos de la Teoría de Conjuntos Difusos a dichas áreas.
2. Como una potente herramienta para tratar situaciones del mundo real en las que aparece incertidumbre (imprecisión, vaguedad, inconsistencia, etc.). Debido a la generalidad de esta teoría, ésta se adapta con facilidad a diferentes contextos y problemas: Teoría de Sistemas [21, 115], Teoría de la Decisión [53, 56], Bases de Datos [16, 158, 165], etc. En muchas ocasiones esto implicará adaptar los conceptos originales de la Teoría de los Conjuntos Difusos a los diferentes contextos en los que se esté trabajando.

En este apartado haremos una pequeña revisión de los conceptos básicos de la Teoría de Conjuntos Difusos que utilizaremos en esta memoria. Esta introducción no pretender ser exhaustiva sino una breve presentación de dichos conceptos. Para mayor detalle, véase [91].

1.3.1. Conjuntos Difusos y Función de Pertenencia

La noción de conjunto refleja la idea de agrupar colecciones de objetos que cumple una o varias propiedades que caracterizan a dicho conjunto. Una propiedad puede ser considerada como una función que a cada elemento del universo de discurso X le asigna un valor en el conjunto $\{0, 1\}$, de forma que si el elemento pertenece al conjunto, es decir, cumple la propiedad se le asigna el valor 1 o en caso contrario el valor 0.

Definición 1.1. Sea A un conjunto en el universo X , la función característica asociada a A , $A(x)$, $x \in X$, se define como:

$$A(x) = \begin{cases} 1, & \text{si } x \in A \\ 0, & \text{si } x \notin A. \end{cases}$$

La función $A : X \longrightarrow \{0, 1\}$ induce una restricción, con un límite bien definido, sobre los objetos del universo X que pueden ser asignados al conjunto A .

El concepto de conjunto difuso lo que hace es *relajar* esta restricción y admite valores intermedios en la función característica, que pasa a denominarse *función de pertenencia*.

Esta relajación permite una interpretación más realista de ciertos contextos de trabajo. La mayoría de las categorías que describen los objetos del mundo real no tienen unos límites claros y bien definidos, por ejemplo, ordenador *potente*, buen sabor, coche *veloz*, etc. (las palabras en itálica identifican fuentes de imprecisión). Si un objeto pertenece a una categoría con un grado de pertenencia que puede ser expresado por un número real en el intervalo $[0, 1]$, cuanto más cercano a 1 sea el grado, indicará mayor pertenencia a esa categoría determinada y cuanto más cercano a 0 indicará menor pertenencia a dicha categoría.

1.3.2. Definición Básica de Conjunto Difuso

Un conjunto difuso puede definirse como una colección de objetos con valores de pertenencia entre 0 (exclusión completa) y 1 (pertenencia completa). Los valores de pertenencia expresan los grados con los que cada objeto es compatible con las propiedades o características distintivas de la colección. Formalmente podemos definir los conjuntos difusos como sigue [160]:

Definición 1.2. *Un conjunto difuso $\tilde{A}$ sobre X está caracterizado por una función de pertenencia que transforma los elementos de un dominio, espacio, o universo del discurso X en el intervalo $[0, 1]$.*

$$\mu_{\tilde{A}} : X \longrightarrow [0, 1].$$

Así, un conjunto difuso $\tilde{A}$ en X puede representarse como un conjunto de pares ordenados de un elemento genérico x , $x \in X$ y su grado de pertenencia $\mu_{\tilde{A}}(x)$:

$$\tilde{A} = \{(x, \mu_{\tilde{A}}(x)) / x \in X, \mu_{\tilde{A}}(x) \in [0, 1]\}$$

Claramente, un conjunto difuso es una generalización del concepto de conjunto cuya función de pertenencia toma sólo dos valores $\{0, 1\}$.

Ejemplo

Consideremos el concepto “persona joven”, en un contexto donde se clasifica a las personas atendiendo exclusivamente a la edad que oscila en el intervalo $E = [1, 110]$ años. Una persona cuya edad sea menor o igual a 30 años se puede considerar como joven y por lo tanto se le asignará un valor 1 a su grado de pertenencia al conjunto difuso de personas jóvenes. Una persona con una edad igual o superior a 65 años no puede considerarse como una persona joven y de ahí que se le asigne el valor 0 al grado de pertenencia al conjunto difuso de persona joven. La cuantificación del resto de valores puede llevarse a cabo mediante una función

de pertenencia $\mu_{\tilde{J}} : E \longrightarrow [0, 1]$ que caracteriza el conjunto difuso $\tilde{J}$ de personas jóvenes en el universo $E = [1, 110]$.

$$\mu_{\tilde{J}}(x) = \begin{cases} 1 & x \in [1, 30] \\ 1 - \frac{x-30}{35} & x \in (30, 65) \\ 0 & x \in [65, 110]. \end{cases}$$

Los conjuntos difusos pueden ser definidos sobre universos finitos o infinitos usando distintas notaciones. Si un universo X es discreto y finito, con cardinalidad n , el conjunto difuso puede expresarse con un vector n -dimensional cuyos valores son los grados de pertenencia de los correspondientes elementos de X . Por ejemplo, si $X = \{x_1, \dots, x_n\}$, entonces un conjunto difuso $\tilde{A} = \{(a_i/x_i) | x_i \in X\}$, donde $a_i = \mu_{\tilde{A}}(x_i)$, $i = 1, \dots, n$, puede notarse por [91]:

$$\tilde{A} = a_1/x_1 + a_2/x_2 + \dots + a_n/x_n = \sum_{i=1}^n a_i/x_i.$$

Cuando el universo X es continuo, para representar un conjunto difuso usamos la siguiente expresión:

$$\tilde{A} = \int_x a/x,$$

donde $a = \mu_{\tilde{A}}(x)$ y la integral debería ser interpretada de la misma forma que el sumatorio en el universo finito.

A continuación, introducimos otros conceptos básicos a la hora de trabajar con conjuntos difusos, como son el **soporte**, el **núcleo** y el α -**corte** de un conjunto difuso:

Definición 1.3. *El soporte de un conjunto difuso $\tilde{A}$, $\text{Soporte}(\tilde{A})$, es el conjunto de todos los elementos de $x \in X$, tales que, el grado de pertenencia sea mayor que cero.*

$$\text{Soporte}(\tilde{A}) = \{x \in X \mid \mu_{\tilde{A}}(x) > 0\}.$$

Si esta definición la restringimos a aquellos elementos del universo X con grado de pertenencia igual a 1, tendríamos el núcleo del conjunto difuso.

Definición 1.4. *El núcleo de un conjunto difuso $\tilde{A}$, $Núcleo(\tilde{A})$, es el conjunto de todos los elementos de $x \in X$, tales que el grado de pertenencia es igual a 1.*

$$Núcleo(\tilde{A}) = \{x \in X \mid \mu_{\tilde{A}}(x) = 1\}.$$

En muchas ocasiones puede ser interesante conocer no sólo los elementos que pertenecen en algún grado al conjunto difuso, sino también conocer el conjunto de aquellos elementos que lo hacen con un valor al menos igual o mayor que un umbral determinado α . Estos conjuntos se denominan α -cortes.

Definición 1.5. *Sea $\tilde{A}$ un conjunto difuso sobre el universo X , dado un número $\alpha \in [0, 1]$. Se define el α -corte sobre $\tilde{A}$, ${}^{\alpha}A$, como un conjunto que contiene todos los valores del universo X cuya función de pertenencia en $\tilde{A}$ sea mayor o igual al valor α :*

$${}^{\alpha}A = \{x \in X \mid \mu_{\tilde{A}}(x) \geq \alpha\}.$$

1.3.3. Tipos de Funciones de Pertenencia

En principio cualquier función $\mu_{\tilde{A}} : X \longrightarrow [0, 1]$, describe una función de pertenencia asociada a un conjunto difuso $\tilde{A}$ que depende no sólo del concepto que representa, sino también del contexto en el que se usa. Las gráficas de las funciones pueden tener diferentes representaciones o formas y pueden tener algunas propiedades específicas (ej., continuidad).

Los conjuntos difusos suelen representarse con familias de funciones paramétricas. Las más comunes son las siguientes:

1. Función Triangular:

$$\mu_{\tilde{A}}(x) = \begin{cases} 0 & \text{si } x \leq a \\ \frac{x-a}{b-a} & \text{si } x \in [a, b] \\ \frac{c-x}{c-b} & \text{si } x \in [b, c] \\ 0 & \text{si } x \geq c, \end{cases}$$

donde b es el punto modal de la función triangular y a y c los límites inferior y superior, respectivamente, para los valores no nulos de $\mu_{\tilde{A}}(x)$.

2. Función Trapezoidal:

$$\mu_{\tilde{A}}(x) = \begin{cases} 0, & \text{si } x \leq a \\ \frac{x-a}{b-a} & \text{si } x \in [a, b] \\ 1 & \text{si } x \in [b, d] \\ \frac{c-x}{c-d} & \text{si } x \in [d, c] \\ 0 & \text{si } x \geq c, \end{cases}$$

donde b y d indican el intervalo dónde la función de pertenencia vale 1.

3. Función Gaussiana:

$$A(x) = e^{-k(x-m)^2},$$

donde $k > 0$.

1.3.4. Principio de Extensión

El Principio de Extensión es un concepto básico de la Teoría de Conjuntos Difusos utilizado para generalizar conceptos matemáticos no difusos a conjuntos difusos. A lo largo del tiempo han aparecido diferentes formulaciones de este concepto [46, 91] que se puede definir como:

Definición 1.6. Sea X el producto cartesiano de los universos $X_1, \dots, X_r$ y sean $\tilde{A}_1, \dots, \tilde{A}_r$, r conjuntos difusos en $X_1, \dots, X_r$ respectivamente. Sea f una función definida desde el universo X , ($X = X_1 \times \dots \times X_r$), al universo Y , $y = f(x_1, \dots, x_r)$. El Principio de Extensión nos permite definir un conjunto difuso $\tilde{B}$ en Y , a partir de los conjuntos difusos $\tilde{A}_1, \dots, \tilde{A}_r$ representando su imagen a partir de la función f , de acuerdo a la siguiente expresión,

$$\tilde{B} = \{(y, \mu_{\tilde{B}}(y)) / y = f(x_1, \dots, x_r), (x_1, \dots, x_r) \in X\}$$

donde

$$\mu_{\tilde{B}}(y) = \begin{cases} \sup_{(x_1, \dots, x_r) \in f^{-1}(y)} \min\{\mu_{\tilde{A}_1}(x_1), \dots, \mu_{\tilde{A}_r}(x_r)\}, & \text{si } f^{-1}(y) \neq \emptyset \\ 0, & \text{en otro caso.} \end{cases}$$

Para $r = 1$, el Principio de Extensión se reduce a:

$$\tilde{B} = f(A) = \{(y, \mu_{\tilde{B}}(y)) / y = f(x), x \in X\}$$

donde

$$\mu_{\tilde{B}}(y) = \begin{cases} \sup_{x \in f^{-1}(y)} \mu_{\tilde{A}}(x), & \text{si } f^{-1}(y) \neq \emptyset \\ 0, & \text{en otro caso.} \end{cases}$$

1.3.5. Número difuso

Entre los distintos tipos de conjuntos difusos, tienen una especial significación aquellos que están definidos sobre el conjunto de los números reales, $\mathcal{R}$.

$$\tilde{A} : \mathcal{R} \longrightarrow [0, 1]$$

Bajo ciertas condiciones estos conjuntos difusos pueden ser vistos como “números difusos” o “intervalos difusos”, definiéndose el concepto de número difuso como [46, 161]:

Definición 1.7. Un número difuso $\tilde{A}$ es un subconjunto de $\mathcal{R}$ que verifica las siguientes propiedades:

1. *La función de pertenencia es convexa,*

$$\forall x, y \in \mathcal{R}, \forall z \in [0, 1], \mu_{\tilde{A}}(z) \geq \min\{\mu_{\tilde{A}}(x), \mu_{\tilde{A}}(y)\}$$

2. *Para cualquier $\alpha \in (0, 1]$, ${}^\alpha A$ debe ser un intervalo cerrado.*

3. *El soporte de $\tilde{A}$ debe ser finito.*

4. *$\tilde{A}$ está normalizado,*

$$\sup_x \mu_{\tilde{A}}(x) = 1$$

Casos particulares de números difusos [91]:

- Los números reales (Figura 1.1.a).
- Intervalos de números reales (Figura 1.1.b).
- Valores aproximados (Figura 1.1.c).
- Intervalos aproximados o difusos (Figura 1.1.d).

Para terminar esta sección dedicada a los conjuntos difusos tan sólo añadir que las operaciones aritméticas habituales sobre números reales se extienden a los números difusos mediante el Principio de Extensión presentado en el apartado 1.3.4

1.4. El Enfoque Lingüístico Difuso

Los problemas presentes en el mundo real presentan atributos que pueden ser de distinta naturaleza. Cuando dichos atributos son de naturaleza cuantitativa, éstos se valoran fácilmente utilizando valores numéricos más o menos precisos. Sin

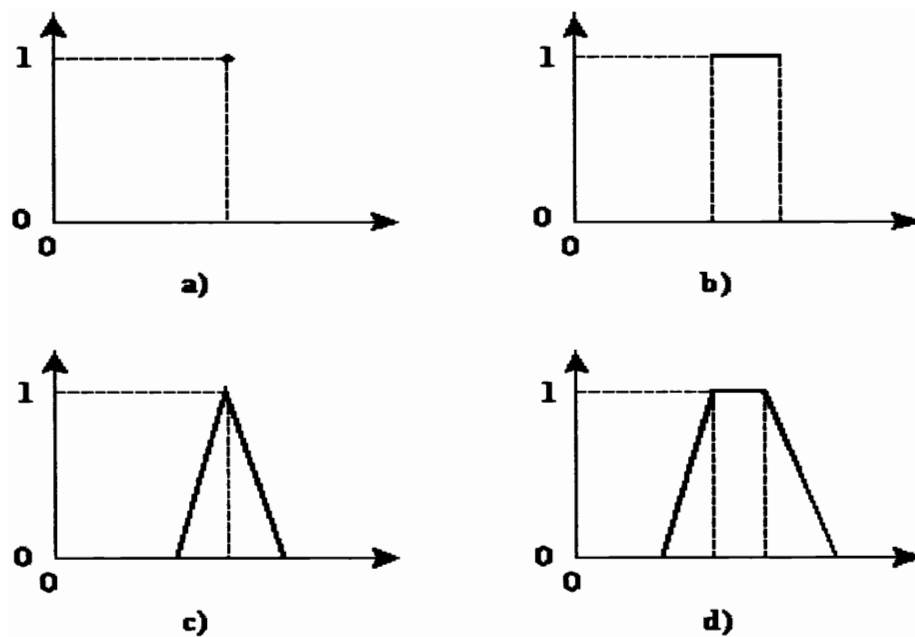


Figura 1.1: Ejemplos de números difusos

embargo, cuando se trabaja con información vaga e imprecisa o cuando la naturaleza de tales atributos no es cuantitativa sino cualitativa, no es sencillo ni adecuado utilizar un modelado de preferencias numérico, aconsejándose utilizar otro tipo de modelado como por ejemplo el lingüístico. Este tipo de atributos suelen aparecer frecuentemente en problemas en los que se pretende evaluar fenómenos relacionados con percepciones y relaciones de los seres humanos (diseño, gusto, ...). En estos casos se suele utilizar palabras del lenguaje natural (bonito, feo, dulce, salado, simpático, ...) en lugar de valores numéricos para emitir tales valoraciones. Tal y como se indica en [30, 163], el uso de un modelado lingüístico de preferencias puede deberse a varias razones:

1. La información disponible con la que trabajan los expertos es demasiado vaga o imprecisa para ser valorada utilizando valores numéricos precisos.
2. Situaciones en las que la información no puede ser cuantificada debido a su naturaleza, y así, sólo puede “medirse” utilizando términos lingüísticos (ej., como se vio anteriormente al evaluar el “confort” o el “diseño” de un coche [97], el uso de términos como “bueno”, “medio”, “malo” suelen ser habituales).
3. Información cuantitativa que no puede medirse porque no están disponibles los elementos necesarios para llevar a cabo una medición exacta o porque el coste de su medición es muy elevado. En este caso el uso de un “valor aproximado” que permita reflejar los distintos valores del problema puede ser adecuado (ej., imaginemos una situación en la que se pretende evaluar la velocidad de un coche y no disponemos de velocímetro, sirviéndonos tan sólo de nuestras percepciones, entonces se puede utilizar términos lingüísticos como “rápido”, “muy rápido”, “lento” para medir la velocidad en lugar de valores numéricos).

El Enfoque Lingüístico Difuso, que tiene como base teórica la Teoría de los Conjuntos Difusos, se ha mostrado como una técnica eficaz para valorar atributos de naturaleza cualitativa [3, 4, 8, 9, 39, 60, 100, 140, 151, 156]. Utiliza *variables lingüísticas* cuyo dominio de expresión son conjuntos de palabras o términos lingüísticos [161].

Una variable lingüística se caracteriza por un *valor sintáctico* o *etiqueta* y por un *valor semántico* o *significado*. La etiqueta es una palabra o frase perteneciente a un conjunto de términos lingüísticos y el significado de dicha etiqueta viene dado por un subconjunto difuso en un universo del discurso. Al ser las palabras

menos precisas que los números, el concepto de variable lingüística es una buena propuesta para caracterizar aquellos fenómenos que no son adecuados para poder ser evaluados mediante valores numéricos. Formalmente una variable lingüística se define de forma precisa como [161]:

Definición 1.8. *Una variable lingüística está caracterizada por una quintupla $(H, T(H), \mathcal{U}, G, M)$, en la que:*

- *H es el nombre de la variable.*
- *$T(H)$ es el conjunto de valores lingüísticos o etiquetas lingüísticas.*
- *$\mathcal{U}$ es el universo de discurso de la variable.*
- *G es una regla sintáctica (que normalmente toma forma de gramática) para generar los valores de $T(H)$.*
- *M es una regla semántica que asocia a cada elemento de $T(H)$ su significado. Para cada valor $L \in T(H)$, $M(L)$ será un subconjunto difuso de $\mathcal{U}$.*

Para resolver un problema desde el punto del vista del Enfoque Lingüístico Difuso es necesario llevar a cabo dos operaciones fundamentales:

1. Elección de un adecuado conjunto de términos lingüísticos, $T(H)$.
2. Definición de la semántica asociada a cada término lingüístico.

1.4.1. Elección del Conjunto de Términos Lingüísticos

Para que una fuente de información (experto, consultor, ...) pueda expresar con facilidad su información y/o conocimiento es necesario que disponga de un conjunto apropiado de descriptores lingüísticos. Un aspecto muy importante de

este conjunto es el número de etiquetas lingüísticas disponible para expresar la información, denominado *la granularidad de la incertidumbre* [11].

Se dice que un conjunto de términos lingüísticos tiene:

- Una granularidad baja o un tamaño de grano grueso cuando la cardinalidad del conjunto de etiquetas lingüísticas es pequeña. Esto significa que el dominio está poco particionado y que existen pocos niveles de distinción de la incertidumbre, produciéndose una pérdida de expresividad.
- Una granularidad alta o un tamaño de grano fino cuando la cardinalidad del conjunto de etiquetas lingüísticas es alta. Esta situación puede provocar un aumento de la complejidad en la descripción del dominio.

La cardinalidad de un conjunto de términos lingüísticos no debe ser demasiado pequeña como para imponer una restricción de precisión a la información que quiere expresar cada fuente de información, y debe ser lo suficientemente grande para permitir hacer una discriminación de las valoraciones en un número limitado de grados. Valores típicos de cardinalidad usados en modelos lingüísticos son valores impares, tales como 7 ó 9, donde el término medio representa una valoración de “aproximadamente 0.5”. El resto de los términos se distribuyen alrededor de éste [11]. Estos valores clásicos de cardinalidad parecen estar dentro de la línea de observación de Miller [107] sobre la capacidad humana, en la que se indica que se pueden manejar razonablemente y recordar alrededor de siete o nueve términos diferentes.

Una vez establecida la cardinalidad es necesario un mecanismo para generar los términos lingüísticos. Existen dos enfoques para esto, uno los define a partir de una gramática libre de contexto y el otro mediante un orden total definido sobre el conjunto de términos. A continuación analizamos ambos mecanismos.

1. Enfoque Basado en una Gramática Libre de Contexto

Una posibilidad para generar el conjunto de términos lingüísticos consiste en utilizar una gramática libre de contexto G , donde el conjunto de términos pertenece al lenguaje generado por G [9, 15, 162]. Una gramática generadora, G , es una 4-tupla (V_N, V_T, I, P) , siendo V_N el conjunto de símbolos no terminales, V_T el conjunto de símbolos terminales, I el símbolo inicial y P el conjunto de reglas de producción. La elección de estos cuatro elementos determinará la cardinalidad y forma del conjunto de términos lingüísticos. El lenguaje generado debería ser lo suficientemente grande para que pueda describir cualquier posible situación del problema.

De acuerdo con las observaciones de Miller [107], el lenguaje generado no tiene que ser infinito, sino mas bien fácilmente comprensible.

Por ejemplo, entre los símbolos terminales y no terminales de G podemos encontrar términos primarios (ej.: *alto*, *medio*, *bajo*), modificadores (ej.: *no*, *mucho*, *muy*, *más o menos*), relaciones (ej.: *mayor que*, *menor que*) y conectivos (ej.: *y*, *o*, *pero*). Construyendo I como cualquier término primario, el conjunto de términos lingüísticos $T(H) = \{muy\ alto, alto, alto\ o\ medio, \dots\}$ se genera usando P .

2. Enfoque basado en Términos Primarios con una Estructura Ordenada

Una alternativa para reducir la complejidad de definir una gramática consiste en dar directamente un conjunto de términos distribuidos sobre una escala con un orden total definido [14, 66, 156]. Por ejemplo, consideremos el siguiente conjunto de siete etiquetas $T(H) = \{N, MB, B, M, A, MA, P\}$:

$$\begin{aligned} s_0 &= N = Nada & s_1 &= MB = Muy_Bajo & s_2 &= B = Bajo \\ s_3 &= M = Medio & s_4 &= A = Alto & s_5 &= MA = Muy_Alto \\ s_6 &= P = Perfecto \end{aligned}$$

donde $s_i < s_j$ si y sólo si $i < j$.

Normalmente en estos casos es necesario que los términos lingüísticos satisfagan las siguientes condiciones adicionales:

1. Existe un operador de negación. Por ejemplo, $Neg(s_i) = s_j, j = g - i$ ($g+1$ es la cardinalidad de $T(H)$).
2. Tiene un operador de maximización: $máx(s_i, s_j) = s_i$ si $s_i \geq s_j$.
3. Tiene un operador de minimización: $mín(s_i, s_j) = s_i$ si $s_i \leq s_j$.

Nosotros utilizamos este último modo en nuestra memoria.

1.4.2. Semántica del Conjunto de Términos Lingüísticos

En la literatura existen varios enfoques para definir la semántica del conjunto de etiquetas lingüísticas [14, 138, 139, 154], siendo uno de los más utilizados el enfoque basado en funciones de pertenencia [11, 15, 41, 94, 100, 103].

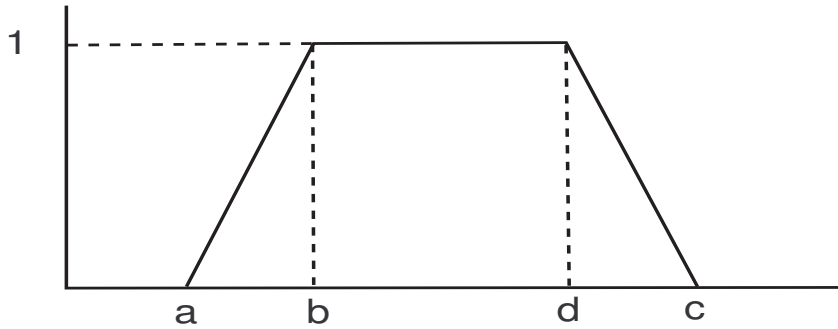


Figura 1.2: Número difuso trapezoidal

Este enfoque define la semántica del conjunto de términos lingüísticos utilizando números difusos en el intervalo $[0, 1]$, donde cada número difuso es descrito por una función de pertenencia.

Un método eficiente desde un punto de vista computacional para caracterizar un número difuso es usar una representación basada en parámetros de su función de pertenencia [9, 43]. Debido a que las valoraciones lingüísticas dadas por las fuentes de información son aproximaciones, algunos autores consideran que las funciones de pertenencia trapezoidales son lo suficientemente buenas para representar la vaguedad de dichas valoraciones lingüísticas [9, 11, 41, 136, 137]. Esta representación paramétrica se expresa usando una 4-tupla (a, b, d, c) se muestra en la Figura 1.2. Los parámetros “ b ” y “ d ” indican el intervalo en el que la función de pertenencia vale 1; mientras que “ a ” y “ c ” indican los extremos izquierdo y derecho de la función de pertenencia [9].

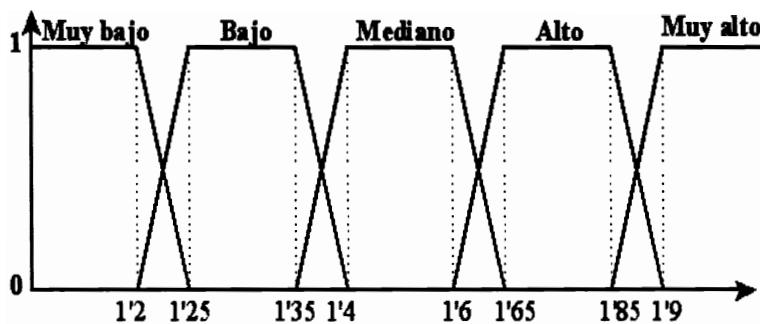


Figura 1.3: Definición semántica de la variable lingüística altura

En la Figura 1.3 se muestra la semántica de una variable lingüística que evalúa la altura de una persona utilizando números difusos definidos por funciones de pertenencia trapezoidales:

$$T(Altura) = \{\text{Muy bajo, Bajo, Mediano, Alto, Muy Alto}\}$$

$$Muy\ Bajo = (0, 0, 1.2, 1.25)$$

$$Bajo = (1.2, 1.25, 1.35, 1.4)$$

$$Mediano = (1.35, 1.4, 1.6, 1.65)$$

$$Alto = (1.6, 1.65, 1.85, 1.9)$$

$$Muy\ Alto = (1.85, 1.9, 2, 2)$$

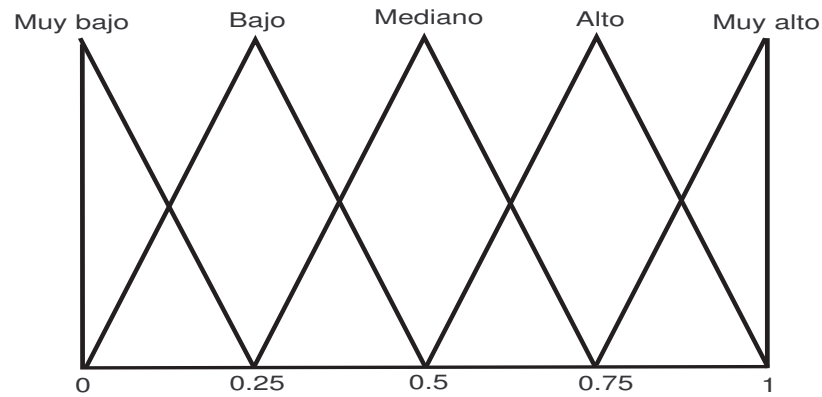


Figura 1.4: Conjunto de 5 etiquetas lingüísticas uniformemente distribuidas

Un caso particular de este tipo de representación son las funciones de pertenencia triangulares, en las que $b = d$. Se representan mediante una terna (a, b, c) , donde “ b ” es el valor donde la función de pertenencia vale 1, mientras que “ a ” y “ c ” indican los extremos izquierdo y derecho de la función.

La Figura 1.4 muestra el mismo conjunto visto anteriormente pero representado ahora con funciones de pertenencia triangulares.

Otros autores usan otro tipo de funciones como por ejemplo funciones Gaussianas [15].

Este enfoque implica establecer las funciones de pertenencia asociadas a cada etiqueta. Esta tarea presenta el problema de determinar los parámetros según los

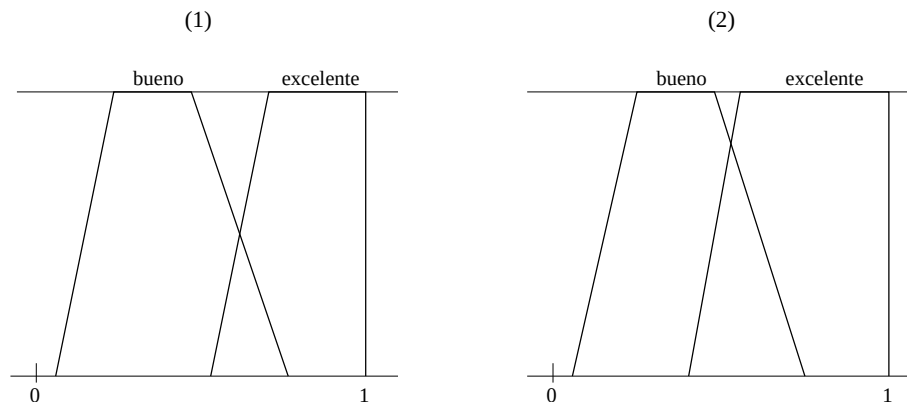


Figura 1.5: Distribuciones diferentes del concepto “excelente”

puntos de vista de todas las fuentes de información. En la realidad, es difícil que todas las fuentes de información propongan exactamente las mismas funciones de pertenencia asociadas a los términos lingüísticos primarios, debido a que cada una de ellas puede interpretar de forma parecida pero a la vez diferente el mismo concepto. Por ejemplo, en la Figura 1.5, vemos dos percepciones muy cercanas pero diferentes de la evaluación del concepto “excelente”.

Por lo tanto, pueden darse el caso de términos lingüísticos con una sintaxis similar pero con diferente semántica [60]. Este tipo de situaciones pueden presentarse en problemas con múltiples fuentes de información en contextos lingüísticos donde algunos expertos utilizan conjuntos de términos lingüísticos iguales en sintaxis y cardinalidad pero con diferente semántica.

1.4.3. Modelo de Representación Lingüístico Basado en 2-tuplas

El modelo de representación lingüístico basado en 2-tuplas fue presentado en [68] para mejorar los problemas de pérdida de información en los procesos de computación con palabras de otros modelos: (i) Modelo basado en el Principio de Extensión [39], (ii) Modelo simbólico [42]. Además se ha demostrado como un

modelo útil en el tratamiento de contextos de información no homogéneos [62, 71] puesto que nos permite mejorar su comprensión y su manipulación matemática con el objetivo de facilitar el proceso de ordenación para la fase de explotación de nuestro proceso de decisión propuesto en esta memoria.

Este modelo se basa en el concepto de traslación simbólica. Sea $S = \{s_0, \dots, s_g\}$ un conjunto de términos lingüísticos, y $\beta \in [0, g]$ un valor en el intervalo de granularidad de S .

Definición 1.9. *La Traslación Simbólica de un término lingüístico s_i es un número valorado en el intervalo $[-.5, .5)$ que expresa la “diferencia de información” entre una cantidad de información expresada por el valor $\beta \in [0, g]$ obtenido en una operación simbólica y el valor entero más próximo, $i \in \{0, \dots, g\}$, que indica el índice de la etiqueta lingüística (s_i) más cercana en S .*

A partir de este concepto desarrollaremos un nuevo modelo de representación para la información lingüística, el cuál usa como base de representación un par de valores o 2-tupla, (r_i, α_i) , donde $r_i \in S$ y $\alpha_i \in [-.5, .5)$

Este modelo de representación define un conjunto de funciones que facilitan las operaciones sobre 2-tuplas.

Definición 1.10. *Sea $S = \{s_0, \dots, s_g\}$ un conjunto de términos lingüísticos y $\beta \in [0, g]$ un valor que representa el resultado de una operación simbólica, entonces la 2-tupla lingüística que expresa la información equivalente a β se obtiene usando la siguiente función:*

$$\Delta : [0, g] \longrightarrow S \times [-.5, .5)$$

$$\Delta(\beta) = (s_i, \alpha), \text{ con } \begin{cases} s_i, & i = \text{round}(\beta) \\ \alpha = \beta - i, & \alpha \in [-.5, .5), \end{cases}$$

donde round es el operador usual de redondeo, s_i es la etiqueta con índice más cercano a β y α es el valor de la traslación simbólica.

Proposición 1.1. Sea $S = \{s_0, \dots, s_g\}$ un conjunto de términos lingüísticos y (s_i, α) una 2-tupla lingüística. Existe la función Δ^{-1} , tal que, dada una 2-tupla (s_i, α) esta función devuelve su valor numérico equivalente $\beta \in [0, g]$.

Demostración.

Es trivial si consideramos la siguiente función.

$$\begin{aligned}\Delta^{-1} : Sx[-.5, .5) &\longrightarrow [0, g] \\ \Delta^{-1}(s_i, \alpha) &= i + \alpha = \beta.\end{aligned}$$

Comentario: A partir de las definiciones 1.9 y 1.10 y de la proposición 1.1, la conversión de un término lingüístico en una 2-tupla consiste en añadir el valor cero como traslación simbólica:

$$s_i \in S \longrightarrow (s_i, 0)$$

Este modelo de representación tiene un modelo computacional asociado presentado en [68]:

1. **Agregación de 2-tuplas:** La agregación de 2-tuplas lingüísticas consiste en la obtención de un valor que resuma un conjunto de valores, por lo tanto, el resultado de la agregación del conjunto de 2-tuplas debe ser una 2-tupla lingüística. En [68] podemos encontrar varios operadores de agregación basados en los operadores de agregación clásicos.
2. **Comparación de 2-tuplas:** La información de comparación representada por las 2-tuplas la realizamos de acuerdo a un orden lexicográfico.
Sea (s_k, α_1) y (s_l, α_2) dos 2-tuplas que representan dos valoraciones:
 - Si $k < l$ entonces (s_k, α_1) es más pequeño que (s_l, α_2)

- Si $k = l$ entonces:
 - a) Si $\alpha_1 = \alpha_2$ entonces (s_k, α_1) y (s_l, α_2) representan el mismo valor.
 - b) Si $\alpha_1 < \alpha_2$ entonces (s_k, α_1) es más pequeño que (s_l, α_2) .
 - c) Si $\alpha_1 > \alpha_2$ entonces (s_k, α_1) es mayor que (s_l, α_2) .

3. **Operador de negación de una 2-tupla:** El operador de negación de una 2-tupla se define como:

$$Neg(s_i, \alpha) = \Delta(g - \Delta^{-1}(s_i, \alpha))$$

donde $g + 1$ es la cardinalidad de S , $s_i \in S = \{s_0, \dots, s_g\}$.

Capítulo 2

Tratamiento de Información Heterogénea en Problemas de Toma de Decisión

Como hemos visto en la teoría de la decisión, los problemas con múltiples fuentes de información han recibido una especial atención por parte de investigadores pertenecientes a un amplio rango de disciplinas tales como el consenso [14, 75], toma de decisión en grupo (TDG) [8, 149], toma de decisión con múltiples expertos (TDME) [77, 155] o evaluación [31, 103], etc.

Se observa que en los trabajos anteriores la información utilizada es modelada en un único dominio ya sea numérico [127], lingüístico [59] o intervalar [135]. El tipo de modelado de la información suele ser seleccionado atendiendo:

1. A la naturaleza de los atributos que forman parte del problema.
2. A la incertidumbre que existe en el conocimiento que tienen los expertos que participan en el problema.

3. A la necesidad de herramientas que pueden operar con la información del problema.

Nuestro objetivo es poder ofertar una mayor flexibilidad de expresión a los expertos que participan en el problema de decisión de forma que puedan expresar su conocimiento en un contexto heterogéneo con información numérica, intervalar y lingüística. El hecho de forzar a los distintos expertos que toman parte en un problema de decisión con múltiples fuentes de información a expresar su conocimiento mediante un único dominio y/o escala de información se ha mostrado poco adecuado en este tipo de problemas, ya que, en situaciones de decisión en las que participan varios individuos, cada uno de ellos tendrá su propio conocimiento sobre las alternativas que estemos estudiando en ese problema y cada uno de ellos, debido a su área de trabajo o al propio conocimiento sobre el problema o a la naturaleza de los atributos valorados, puede querer expresar sus preferencias en un dominio de expresión y/o escala diferente. En la literatura podemos encontrar algunas propuestas para manejar información no homogénea [40, 60, 72, 105, 166]. De esta forma podemos ofrecer a los expertos que tomen parte en un problema de TD que expresen sus preferencias en el dominio de información más adecuado según su conocimiento y la naturaleza de la alternativa o experto que esté valorando. Pero para hacer esto y según acabamos de ver, hace falta tener elementos como son las herramientas que permitan operar con información heterogénea que define el contexto del problema de TD.

Por tanto, en este capítulo en primer lugar mostraremos el esquema formal de definición de los problemas de TD que pretendemos resolver junto con un modelo de decisión para su resolución, que nos mostrará la necesidad de operar con información heterogénea. Y a continuación presentaremos una serie de herramientas que nos van a permitir realizar las operaciones necesarias del modelo de decisión

sobre información heterogénea con la que tratar dicha información.

2.1. Problema de Toma de Decisiones Definido en un Contexto Heterogéneo

En el Capítulo 1 hemos visto distintos tipos de problemas de TD y distintas posibilidades de modelado de preferencias que pueden utilizar los expertos que toman parte en ellos. Nuestra atención en esta memoria se centra en problemas de TD en los que interfieren varias fuentes de información, ya sea con un esquema de TDME o de TDG. Nuestro objetivo es ofrecer una mayor flexibilidad a la hora de expresar las preferencias por parte de los expertos, por lo que, en primer lugar presentaremos un esquema de cómo serán los problemas de TD definidos en un contexto heterogéneo.

En un problema de TDME o TDG tenemos un conjunto finito de alternativas $X = \{x_1, \dots, x_n\}$ con $n \geq 2$, así como un conjunto finito de expertos $E = \{e_1, \dots, e_m\}$ con $m \geq 2$ que tienen como objetivo alcanzar una solución al problema teniendo en cuenta las opiniones de todos los expertos. Los expertos suministran, de acuerdo a su conocimiento, las preferencias para cada una de las alternativas utilizando algún modelo de representación disponible como vectores de utilidad (TDME) [71, 77] o relaciones de preferencia (TDG) [65, 149]. De forma que si los atributos a valorar tienen una naturaleza cuantitativa expresar las preferencias mediante valoraciones numéricas suele ser adecuado. Sin embargo, a veces ocurre que estas valoraciones no pueden ser expresadas de forma exacta, o bien por su propia naturaleza o bien por depender de factores desconocidos o no predecibles. Como hemos visto anteriormente en estos casos, un modelado de preferencias utilizado en estas situaciones son los intervalos numéricos, pues éstos representan la incertidumbre asociada a la preferencia mediante un rango de va-

lores. Si trabajamos con atributos cualitativos los modelados anteriores no son los más idóneos debido a la propia naturaleza de los atributos a valorar, siendo el enfoque lingüístico difuso [161] más adecuado para estos casos debido a que se adecua mejor para expresar el conocimiento sobre este tipo de atributos y a los resultados satisfactorios que ha obtenido [1, 3, 4, 9, 68, 100, 150, 146].

El esquema de definición de nuestro problema es el siguiente:

Partimos de un conjunto finito de alternativas:

$$X = \{x_1, \dots, x_n\} \quad n \geq 2,$$

y un conjunto finito de expertos:

$$E = \{e_1, \dots, e_m\} \quad m \geq 2.$$

Los expertos expresarán sus preferencias o bien mediante relaciones de preferencia en un problema de TDG:

$$P_{e_k} = \begin{pmatrix} p_{11}^k & \cdots & p_{1n}^k \\ \vdots & \cdots & \vdots \\ p_{n1}^k & \cdots & p_{nn}^k \end{pmatrix}$$

en donde p_{ij}^k es el grado de preferencia de la alternativa i sobre la alternativa j dado por el experto k .

O bien los expertos expresarán sus preferencias mediante vectores de utilidad en problemas de TDME:

$$P_{e_k} = (p_1^k, \dots, p_n^k)$$

en donde p_i^k es la valoración de la alternativa i dado por el experto k sobre la alternativa i .

De esta forma, en un problema de TDG o un problema de TDME, algunas de las relaciones de preferencia o vectores de utilidad pueden ser valorados de forma numérica, y por lo tanto, tendrán la siguiente forma:

$$P_{e_k} = \begin{pmatrix} p_{11}^k & \cdots & p_{1n}^k \\ \vdots & \cdots & \vdots \\ p_{n1}^k & \cdots & p_{nn}^k \end{pmatrix}, p_{ij}^k \in [0, 1]$$

$$P_{e_k} = (p_1^k, \dots, p_n^k), p_i^k \in [0, 1]$$

Otras pueden ser evaluadas con información intervalar:

$$P_{e_k} = \begin{pmatrix} i_{11}^k & \cdots & i_{1n}^k \\ \vdots & \cdots & \vdots \\ i_{n1}^k & \cdots & i_{nn}^k \end{pmatrix}, i_{ij}^k \in \wp([0, 1])$$

$$P_{e_k} = (i_1^k, \dots, i_n^k), i_i^k \in \wp([0, 1])$$

O mediante etiquetas lingüísticas:

$$P_{e_k} = \begin{pmatrix} s_{11}^k & \cdots & s_{1n}^k \\ \vdots & \cdots & \vdots \\ s_{n1}^k & \cdots & s_{nn}^k \end{pmatrix}, s_{ij}^k \in S$$

$$P_{e_k} = (s_1^k, \dots, s_n^k), s_i^k \in S$$

Donde S es el conjunto de etiquetas que utilizaremos para evaluar las alternativas. No es de extrañar que en este caso, los expertos, deseen utilizar un conjunto de etiquetas propio adecuado a su grado de conocimiento. En este tipo de casos

hablamos de información lingüística multigranular en el que cada experto e_i expresará su conocimiento en un conjunto de etiquetas propio S_i que puede tener una semántica, sintáxis o granularidad distinta a los demás expertos.

Como hemos indicado en la introducción vamos a definir un modelo de para resolver problemas de toma de decisión, con múltiples fuentes de información, definidos en un contexto heterogéneo en donde las alternativas pueden estar modeladas en distintos dominios de expresión (numérico, lingüístico e intervalar). La propuesta para dicho modelo de decisión se basará en el modelo lingüístico de decisión que se presenta en la figura 2.1 y los cambios necesarios para tratar con múltiples fuentes de información se presentan en la figura 2.2 que a continuación pasamos a detallar:

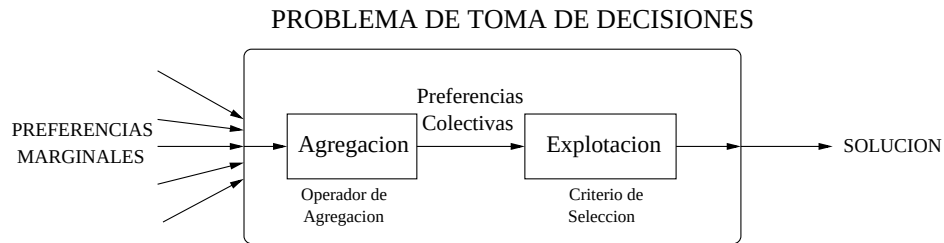


Figura 2.1: Fases de un Proceso de Toma de Decisión

1. *Adquisición de la información:* Los distintos expertos suministrarán sus preferencias en el dominio de información más adecuado a su grado de conocimiento sobre las alternativas del problema.
2. *Proceso de agregación:* esta fase tiene como objeto obtener un valor colectivo para cada una de las alternativas valoradas por los distintos expertos en los diferentes dominios (numérico, intervalar y lingüístico). Para ello, es necesario efectuar los siguientes pasos debido a que manejamos información

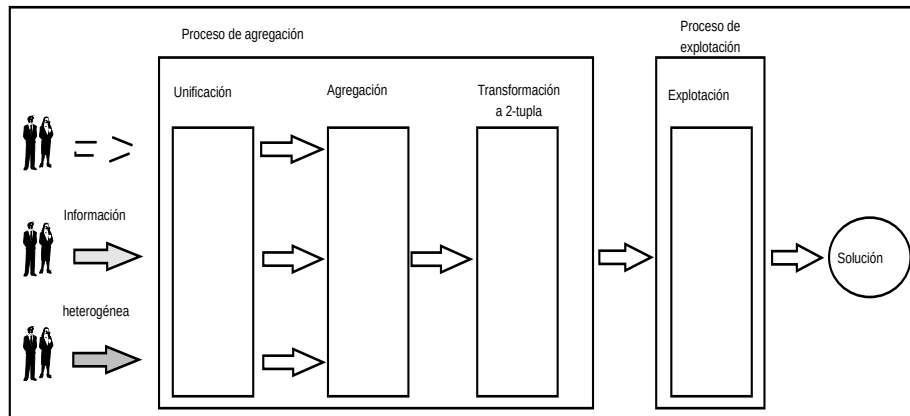


Figura 2.2: Modelo de decisión en contexto heterogéneo

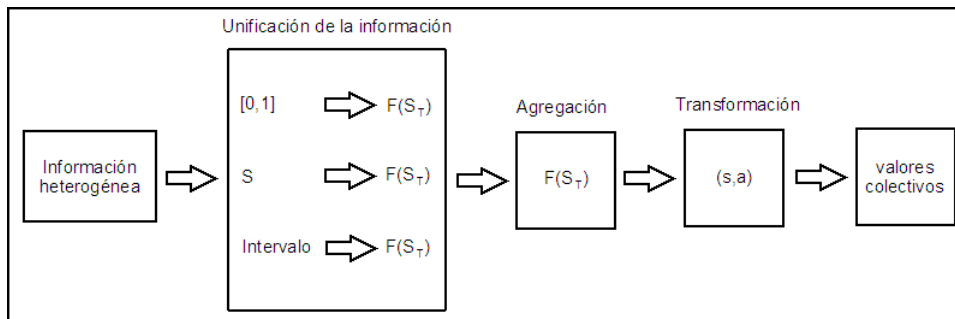


Figura 2.3: Proceso de agregación de información heterogénea

no homogénea:

- a) *Unificación de la información:* Antes de agregar la información heterogénea, necesitamos expresarla en un marco común para poder operar sobre la misma. Este proceso, consiste en unificar la información de entrada (heterogénea) en un único dominio de expresión. Entre los distintos dominios que podemos seleccionar nosotros hemos decidido unificar sobre el dominio lingüístico. Para realizar la unificación de la información se seleccionará un conjunto de etiquetas que denominare-

mos Conjunto Básico de Términos Lingüísticos(CBTL) y simbolizado por S_T . En este proceso de unificación cada valor de preferencia numérico, intervalar y lingüístico suministrado por un experto debe ser expresado en el CBTL, para realizar este proceso se propone convertir cada valor de entrada en un conjunto difuso sobre el CBTL, $F(S_T)$.

- b) *Agregación*: Una vez unificada la información queremos obtener para cada alternativa un valor colectivo. Para ello agregaremos los conjuntos difusos valorados en el CBTL que representan la valoración individual que cada experto ha asignado a esa alternativa, por medio de un operador de agregación. Por lo tanto, cada valor de preferencia colectiva será un conjunto difuso en el dominio lingüístico CBTL.
- c) *Transformación a 2-tupla*: El valor colectivo obtenido para cada una de las alternativas, expresado por medio de un conjunto difuso valorado en el CBTL. Se transformará en una 2-tupla lingüística [68] valorada en el CBTL para facilitar procesos de ordenación de las alternativas en función de dicha valoración y construir la solución del problema de forma fácil y directa [73].

3. *Proceso de Explotación*: Una vez que tenemos los valores de preferencias colectivas de cada una de las alternativas podemos obtener la mejor alternativa o el mejor conjunto de alternativas utilizando una función de selección, tales como, las presentadas en [113, 123]. Esta función de selección ordenará las alternativas de acuerdo a diferentes criterios a partir de esta ordenación obtendremos la mejor o el mejor conjunto de estas.

2.2. Herramientas para Manejar Información Heterogénea

Como hemos visto en la sección anterior para operar con información heterogénea en el proceso de decisión necesitamos una serie de herramientas que nos permitan manejar este tipo de información y su incertidumbre.

Según el modelo de resolución presentado en la figura 2.2 necesitamos definir un conjunto de funciones y operadores matemáticos basados en la lógica y aritmética difusa [46, 160] para el tratamiento de información heterogénea en problemas de decisión en su proceso de agregación. Que está dividido en los siguientes pasos:

- *Adquisición de la información*
- *Unificación de la información*
 - Selección del dominio de expresión unificado: CBTL.
 - Transformar los valores numéricos del intervalo $[0, 1]$ al $F(S_T)$.
 - Transformar los intervalos comprendidos en $[0, 1]$ al $F(S_T)$.
 - Transformar los términos lingüísticos al $F(S_T)$.
- *Agregación*
- *Transformación a 2-tupla*

A continuación presentamos en detalle cada una de las funciones de transformación necesarias para completar los pasos en los que está dividido el proceso de agregación de nuestro modelo de decisión que puede verse en la figura 2.3.

2.2.1. Selección del dominio de expresión unificado: CBTL

Como primer paso en el proceso de agregación deberá unificarse la información de entrada en un dominio de expresión común (CBTL) que gráficamente podemos verlo en la figura 2.4.

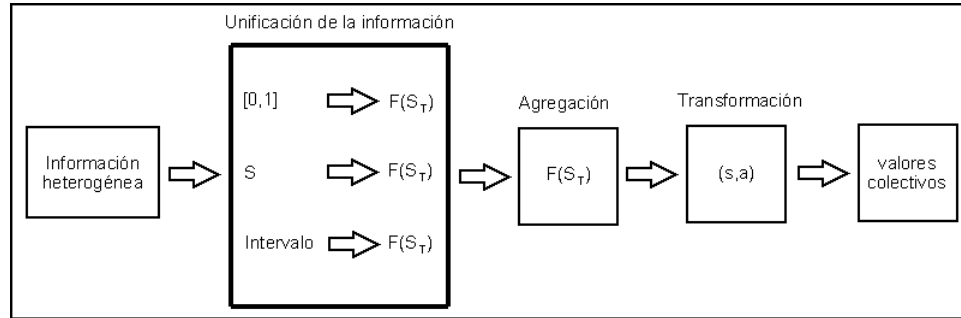


Figura 2.4: Unificación de la información heterogénea

En un problema de TD definido en un contexto heterogéneo la información suministrada por los expertos puede estar expresada en cualquiera de los siguientes dominios:

1. Numérico.
2. Intervalar.
3. Lingüístico o lingüístico multigranular.

No puede operar directamente sobre información heterogénea, en primer lugar necesitamos unificarla en un marco común de expresión para poder realizar procesos de cómputo sobre ella. Este proceso consiste en expresar la información de entrada en un único dominio de expresión. Tal y como ha sido indicado anteriormente la unificación se realizará en un dominio lingüístico denominado (Conjunto Básico

de Términos Lingüísticos), y simbolizado por S_T . En este dominio representaremos toda la información de entrada suministrada por los expertos transformándola en información homogénea mediante conjuntos difusos en el CBTL.

La selección del CBTL no es un proceso aleatorio, sino que sigue un proceso definido por las siguientes reglas:

1. En primer lugar se estudian los conjuntos de términos lingüísticos S_i que pertenecen al contexto de definición de un problema de decisión, es decir, los conjuntos de etiquetas usados por los distintos expertos. Buscaremos el S_i de mayor cardinalidad que llamaremos $S_l = \{s_o, \dots, s_g\}$ donde $g = \max\{|S_i|\}$ (si nos encontramos en un contexto lingüístico no multigranular realizaremos el estudio sobre el único conjunto de etiquetas S) y

SI (S_l es una partición difusa [125]) **Y** (Las funciones de pertenencia de estos términos son triangulares, $s_i = (a_i, b_i, c_i)$)

ENTONCES seleccionaremos S_l como el CBTL, debido al hecho de que, estas condiciones son necesarias y suficientes para que la transformación entre valores en $[0, 1]$ y 2-tuplas, pueda realizarse sin pérdida de información [69].

2. Si no se cumplen esas condiciones o si el contexto de definición del problema no contienen ningún conjunto de etiquetas lingüísticas, seleccionaremos como CBTL un conjunto de términos con un número de términos más grande de lo que una persona es capaz de discriminar (normalmente 11 ó 13, ver [107]) y que satisfaga las condiciones anteriores. Podemos escoger un CBTL con 15 términos simétricamente distribuidos, con la intención de mantener la máxima cantidad de información [60, 69], cuya semántica es (gráficamente, figura 2.5):

s_0	(0, 0, 0.07)	s_1	(0, 0.07, 0.14)	s_2	(0.07, 0.14, 0.21)
s_3	(0.14, 0.21, 0.28)	s_4	(0.21, 0.28, 0.35)	s_5	(0.28, 0.35, 0.42)
s_6	(0.35, 0.42, 0.5)	s_7	(0.42, 0.5, 0.58)	s_8	(0.5, 0.58, 0.65)
s_9	(0.58, 0.65, 0.72)	s_{10}	(0.65, 0.72, 0.79)	s_{11}	(0.72, 0.79, 0.86)
s_{12}	(0.79, 0.86, 0.93)	s_{13}	(0.86, 0.93, 1)	s_{14}	(0.93, 1, 1)

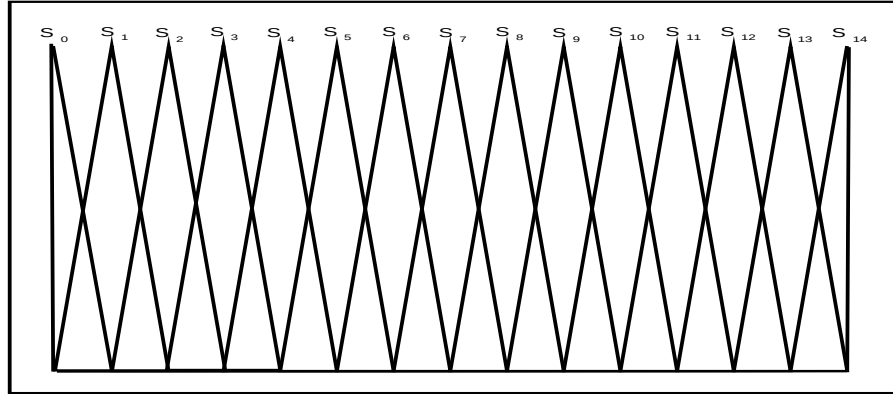


Figura 2.5: Un CBTL con 15 términos simétricamente distribuidos.

Una vez seleccionado el CBTL, es decir el marco común de expresión, se procederá a unificar la información de entrada en conjuntos difusos sobre el CBTL, $F(S_T)$. Para ello son necesarios funciones de transformación que definimos a continuación.

2.2.2. Funciones de Transformación a $F(S_T)$

Como hemos visto en los problemas de TD definidos en un contexto heterogéneo en los que centraremos nuestro estudio, cada experto puede expresar sus conocimientos o preferencias en un dominio diferente debido al área de conocimiento del experto, al propio conocimiento sobre el problema o a la naturaleza de los atributos valorados. Para poder agregar y explotar dicha información debemos

representarla sobre un dominio de expresión unificado que hemos seleccionado. Aquí definiremos las funciones de transformación para unificar la información heterogénea de entrada en conjuntos difusos sobre el CBTL.

Las funciones de transformación que proponemos en esta memoria, transforman la información de entrada en conjuntos difusos, por ello estas funciones realizarán procesos de comparación entre conjuntos difusos. A continuación haremos una breve revisión de las medidas de comparación y seleccionaremos la que utilizaremos para alcanzar nuestro objetivo.

2.2.2.1. Medidas de Comparación

En la literatura podemos encontrar diversos estudios sobre medidas de comparación de descripciones de objetos [17, 18]. Estas comparaciones son habituales en muchos dominios: psicología, analogía, ciencias físicas, procesado de imágenes, clustering, razonamiento deductivo, razonamiento basado en casos, etc., y frecuentemente se basan en medidas que intentan determinar que puntos tienen en común ambos objetos y en cuáles difieren. Las medidas de comparación tienen varias formas dependiendo del propósito de su utilización. En [18] se consideran cuatro tipos de medidas de comparación:

1. *Medidas de satisfacción*: se corresponde con una situación en la cual consideramos un objeto referencia y una clase y decidiremos si el nuevo objeto es compatible o satisface la referencia. Esta situación es típica en el razonamiento basado en prototipos y el nuevo objeto debe ser asociado con uno de ellos.
2. *Medidas de semejanza*: se utilizan para realizar una comparación entre las descripciones de dos objetos, del mismo nivel de generalidad, para decidir si ellos tienen muchas características comunes. Nos encontramos con esta

situación cuando trabajamos en sistemas de razonamiento basado en casos. También se emplea como base de la lógica de la similaridad [126].

3. *Medidas de inclusión*: la inclusión concierne también a situaciones con objetos referencia y medidas, vemos si los puntos en común de A y B son importantes con respecto a A . Puede ser utilizado en sistemas de gestión de base de datos para decidir si una clase está incluida en otra [112, 122].
4. *Medidas de disimilaridad*: la disimilaridad entre objetos evalúa hasta qué punto son diferentes. Este valor puede ser útil cuando, en la recuperación de información de los pasos en un sistema de razonamiento basado en casos, ningún caso es lo suficientemente similar al nuevo caso. Entonces es interesante ser capaz de establecer comparaciones con respecto a las diferencias entre las descripciones, y escoger el caso menos diferente con respecto al nuevo.

Las tres primeras son medidas de similitud.

Todas estas medidas trabajan en entornos difusos pues operan con valores imprecisos e inexactos. Para transformar la información de entrada en conjuntos difusos en el CBTL utilizaremos medidas de semejanza (ver lo parecido que son dos conjuntos difusos) pues son las más apropiadas para el problema que estamos considerando.

En [18] podemos encontrar la definición de una medida (*M-medida*) de semejanza, la cual es definida de la siguiente forma:

Definición 2.1. Una *M-medida* de semejanza es una *M-medida* de similitud S_C que satisface las propiedades:

- *Reflexiva*: $S_C(A, A) = 1$.
- *Simétrica*: $S_C(A, B) = S_C(B, A)$

Estas propiedades ya fueron estudiadas por Kaufmann cuando describió la semejanza [89].

Definimos una *M-medida* [18] como:

Definición 2.2. Una *M-medida* de comparación en Ω es una correspondencia $S_Q : F(\Omega) \times F(\Omega) \rightarrow [0, 1]$ tales que $S_Q(A, B) = F_{S_Q}(M(A \cap B), M(B - A), M(A - B))$, para una correspondencia dada $F_S : \mathbb{R}^+ \times \mathbb{R}^+ \times \mathbb{R}^+ \rightarrow [0, 1]$ y una medida difusa M en Ω .

Donde Ω es el conjunto de elementos que queremos comparar, $F(\Omega)$ es el conjunto de subconjuntos difusos de Ω y M es una medida de conjuntos difusos definida de la siguiente forma:

Definición 2.3. Una medida de conjuntos difusos M es una correspondencia:

$F(\Omega) \rightarrow \mathbb{R}^+$ tal que, para cada A y B en $F(\Omega)$:

- $M(\emptyset) = 0$
- Si $B \subseteq A$, entonces $M(B) \leq M(A)$.

Si los valores de M están restringidos a $[0, 1]$, M es una medida de conjuntos difusos.

Además, empleamos la siguiente definición de operador de diferencia de dos conjuntos:

Definición 2.4. Una operación en $F(\Omega)$ se denomina diferencia y se denotada por $-$, si satisface para cada A y B en $F(\Omega)$:

- Si $A \subseteq B$, entonces $A - B = \emptyset$.
- $B - A$ es monótona con respecto a B : $B \subseteq B'$ entonces $B - A \subseteq B' - A$.

Esta definición es ligeramente diferente de la dada por Kaufmann ([89]), o la de Roberts ([120]), o también la de Dubois y Prade ([47]), pero es compatible con

la definición de una diferencia en el caso de conjuntos crisp en Ω , donde $A - B$ puede ser definido como el complemento de $A \cap B$ en A o, equivalentemente a la intersección de A con el complementario de B en Ω .

Para el problema que estamos tratando, vamos a utilizar una medida de semejanza. En concreto, hemos utilizado una medida basada en una función de posibilidad:

$$S(A, B) = \max_x \min(\mu_A(x), \mu_B(x)), \quad (2.1)$$

donde μ_A y μ_B son funciones de pertenencia de los conjuntos difusos A y B respectivamente y que cumple con las condiciones enunciadas en [18] para ser una medida de comparación.

Hemos elegido esta función de semejanza porque utiliza operaciones fácilmente calculables, es simple, y se ha utilizado con éxito en otros problemas similares [62, 63, 72].

Una vez seleccionada la medida de semejanza que utilizaremos para comparar objetos de distintos dominios presentamos la definición y funcionamiento de las funciones de transformación que hemos definido para:

- La información numérica en $F(S_T)$.
- La información intervalar en $F(S_T)$.
- La información lingüística en $F(S_T)$.

A partir de la función de semejanza presentada en la ecuación (2.1).

2.2.2.2. Transformación de Información Numérica en $F(S_T)$

A continuación presentamos la definición para la función de transformación de información numérica. En la figura 2.6 podemos ver gráficamente en qué paso del

proceso de decisión nos encontramos en este momento.

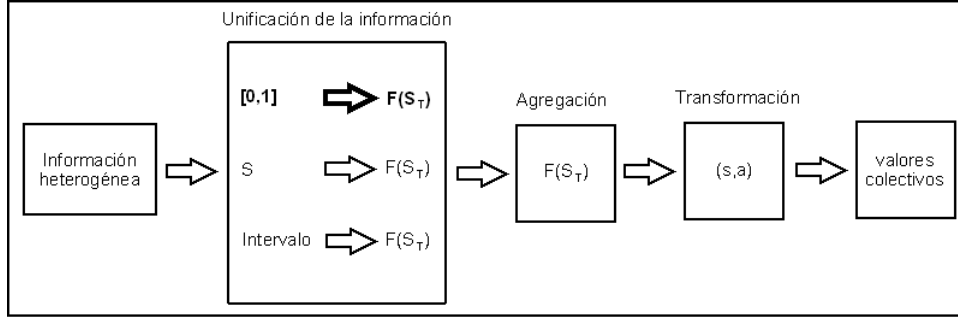


Figura 2.6: Transformación de información numérica en $F(S_T)$

Sea $F(S_T)$ el conjunto de conjuntos difusos definidos sobre el CBTL, $S_T = \{s_0, \dots, s_g\}$, dado un valor numérico $\vartheta \in [0, 1]$, lo transformaremos en un conjunto difuso en $F(S_T)$ calculando el valor de pertenencia de ϑ en las funciones de pertenencia asociadas con los términos lingüísticos de S_T según la siguiente función de transformación $\tau(\vartheta)$:

Definición 2.5. La función τ transforma un valor numérico en un conjunto difuso en S_T :

$$\begin{aligned} \tau : [0, 1] &\rightarrow F(S_T) \\ \tau(\vartheta) &= \{(s_0, \gamma_0), \dots, (s_g, \gamma_g)\}, s_i \in S_T \text{ y } \gamma_i \in [0, 1] \end{aligned} \quad (2.2)$$

$$\gamma_i = \mu_{s_i}(\vartheta) = \begin{cases} 0, & \text{si } \vartheta \notin \text{Soporte}(\mu_{s_i}(x)) \\ \frac{\vartheta - a_i}{b_i - a_i}, & \text{si } a_i \leq \vartheta \leq b_i \\ 1, & \text{si } b_i \leq \vartheta \leq d_i \\ \frac{c_i - \vartheta}{c_i - d_i}, & \text{si } d_i \leq \vartheta \leq c_i \end{cases}$$

Comentario: Consideramos que las funciones de pertenencia, $\mu_{s_i}(\cdot)$, para etiquetas lingüísticas, definida por una función paramétrica (a_i, b_i, d_i, c_i) . Un caso parti-

cular son las valoraciones lingüísticas cuyas funciones de pertenencia son triangulares ($b_i = d_i$).

Ejemplo

Sea $\vartheta = 0.78$ un valor numérico que queremos que sea transformado en un conjunto difuso en $S = \{s_0, \dots, s_4\}$. La semántica de este conjunto de términos puede verse en la figura 2.7:

$$\begin{aligned} s_0 &= (0, 0, 0.25) & s_1 &= (0, 0.25, 0.5) & s_2 &= (0.25, 0.5, 0.75) \\ s_3 &= (0.5, 0.75, 1) & s_4 &= (0.75, 1, 1) \end{aligned}$$

Por lo tanto, el conjunto difuso obtenido de transformar 0.78 sobre $F(S_T)$ será:

$$\tau(0.78) = \{(s_0, 0), (s_1, 0), (s_3, 0.88), (s_4, 0.12)\}$$

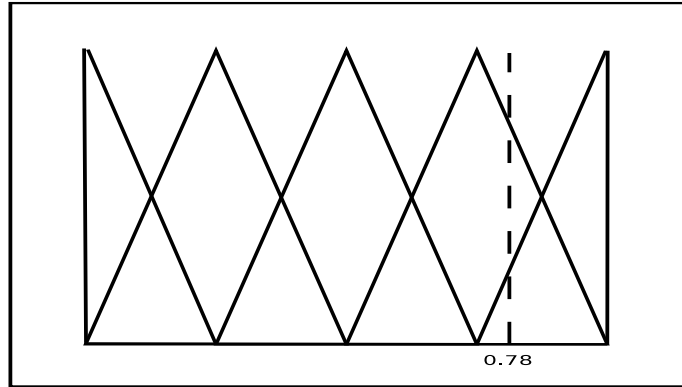


Figura 2.7: Transformación de un valor numérico en un conjunto difuso en S

2.2.2.3. Transformación de Información Intervalar en $F(S_T)$

A continuación pasamos a presentar la definición para la función de transformación intervalar. En la figura 2.8 podemos ver gráficamente en que paso del proceso de decisión nos encontramos en este momento.

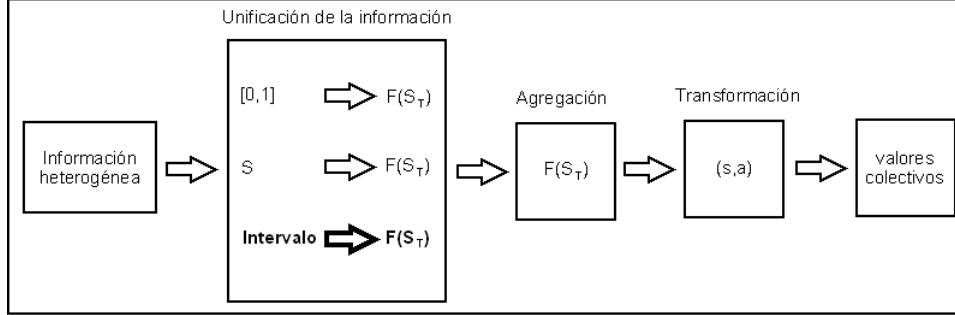


Figura 2.8: Transformación de información intervalar en $F(S_T)$

Para el modelado de preferencias intervalares hemos estudiado dos posibles modelos para su transformación a conjuntos difusos en $F(S_T)$:

1. Una primera función de transformación que utiliza el operador Left [93] y que transforma las preferencias intervalares en preferencias numéricas para ser transformadas en $F(S_T)$ según la transformación que acabamos de presentar en la sección anterior. Este modelo de transformación se ha demostrado útil sólo en problemas de TDME en que los expertos utilizan **vectores de utilidad** para expresar sus preferencias, pero no resulta apropiado para problemas de TDG.
2. Una función de transformación basada en un modelo de representación difusa de la información intervalar presentada en [73] y que resulta más general que la anterior ya que se puede utilizar fácilmente tanto con **relaciones de preferencia** como con **vectores de utilidad** (TDME, TDG).

Transformación de Intervalos en $F(S_T)$ Basada en el Operador Left.

En este caso las preferencias están expresadas mediante un vector de preferencias intervalares $([\underline{a}_1, \overline{a}_1], \dots, [\underline{a}_n, \overline{a}_n])$ sobre un conjunto de alternativas $\{x_1, \dots, x_n\}$ y queremos transformar las preferencias en conjuntos difusos sobre el CBTL.

En este caso el proceso de unificación se llevará a cabo en dos fases:

1. Transformamos los valores de preferencia intervalares en preferencias numéricas.
2. Transformamos los valores de preferencia numéricos en $F(S_T)$ tal y como hemos visto en la sección anterior.

Para la primera fase, la transformación de valores de preferencia intervalares del vector de utilidad en numéricos, los pasos a seguir son los siguientes:

1. **Construir una relación de preferencia:** a partir del vector de preferencia intervalar. Este proceso de construcción se realiza utilizando el operador $Left'(\cdot)$:

$$Left'(A, B) = P_{AB}(x < y)$$

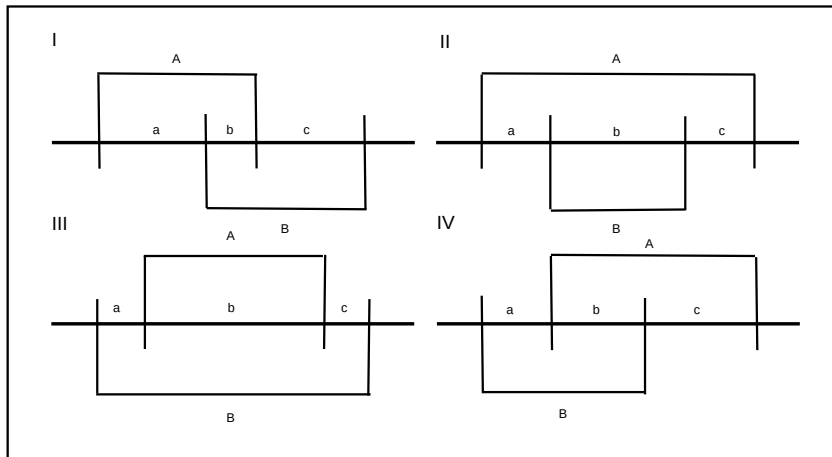
donde $x \in A$ y $y \in B$.

Los valores de $Left'(A, B)$ son:

- a) Sea $A = [\underline{a}, \overline{a}]$ y $B = [\underline{b}, \overline{b}]$ dos intervalos, si $\overline{a} \leq \underline{b}$ entonces

$$Left'(A, B) = 1, \text{ también, si } \underline{a} \geq \overline{b} \text{ entonces } Left'(A, B) = 0$$

- b) La figura 2.9 muestra los 4 casos no triviales de solapamiento de A y B , y la tabla 2.1 muestra los valores de $Left'(A, B)$ en cada uno de estos casos.

Figura 2.9: Los 4 casos no triviales de solapamiento de A y B

Caso		$Left'(A, B)$
I	$A = [0, a + b]$ $B = [a, a + b + c]$	$1 - \frac{b^2}{2(a+b)(b+c)}$
II	$A = [0, a + b + c]$ $B = [a, a + b]$	$\frac{2a+b}{2(a+b+c)}$
III	$A = [a, a + b]$ $B = [0, a + b + c]$	$\frac{b+2c}{2(a+b+c)}$
IV	$A = [a, a + b + c]$ $B = [0, a + b]$	$\frac{b^2}{2(a+b)(b+c)}$

Cuadro 2.1: Los valores de $Left'(A, B)$.

Comentario: consideramos que los valores de preferencia son independientes entre sí.

2. **Explotar la relación de preferencia:** se aplica un proceso de explotación, $\Lambda(\cdot)$, a la relación de preferencia anterior para obtener un valor numérico en

$[0, 1]$ para cada alternativa x_i que exprese la dominancia de x_i sobre el resto de alternativas:

$$\Lambda(x_i) = \frac{1}{n-1} \sum_{j=0|j \neq i}^n Left'(I_i, I_j)$$

donde n es el número de alternativas.

3. **Obtener una preferencia numérica en $[0, 1]$.** Combinaremos el grado de dominancia, $\Lambda(x_i)$, con el centro del intervalo I_i para alcanzar una preferencia numérica para x_i utilizando un operador de agregación. En este modelo hemos usado la media aritmética, aunque se podría haber utilizado otro operador de agregación [57].

Llegados a este punto y habiendo obtenido una representación numérica en $[0, 1]$ del intervalo, pasamos a la segunda fase del proceso, la transformación del valor numérico al $F(S_T)$, utilizando la función de transformación numérica que hemos visto en el apartado anterior.

Ejemplo

Partimos de las siguientes valoraciones intervalares proporcionadas por un experto sobre las alternativas $\{x_1, x_2, x_3, x_4\}$:

	x_1	x_2	x_3	x_4
Valoración	[.75, .95]	[.4, .6]	[.45, .7]	[.65, .8]

Queremos transformar la valoración intervalar al CBTL. Esta transformación se realizará en las dos fases indicadas tal que:

1. *Construir una relación de preferencia.*

A partir del vector de preferencia intervalar construimos una relación de preferencia aplicando el operador $Left'(\cdot)$ y obtenemos:

$$\begin{pmatrix} - & 1 & 1 & 0.96 \\ 0 & - & 0.22 & 0 \\ 0 & 0.77 & - & 0.03 \\ 0.04 & 1 & 0.97 & - \end{pmatrix}$$

2. *Transformamos los valores de preferencia intervalares del vector de utilidad en numéricos.*

Lo primero que haremos será calcular el grado de dominancia de cada una de las alternativas sobre la otra mediante el operador *Left*. Así obtendremos los siguientes valores:

$$\Lambda(x_1) = \frac{1}{3} \sum_{j=0|j \neq 1}^4 Left'(I_1, I_j) = 0.99$$

$$\Lambda(x_2) = \frac{1}{3} \sum_{j=0|j \neq 2}^4 Left'(I_2, I_j) = 0.08$$

$$\Lambda(x_3) = \frac{1}{3} \sum_{j=0|j \neq 3}^4 Left'(I_3, I_j) = 0.265$$

$$\Lambda(x_4) = \frac{1}{3} \sum_{j=0|j \neq 4}^n Left'(I_4, I_j) = 0.675$$

Donde I_i es el valoración intervalar de la alternativa alt_i .

El siguiente paso es obtener el valor numérico de las preferencias en $[0, 1]$. Para ello, combinamos el grado de dominancia, $\Lambda(x_i)$, con el centro del intervalo I_i utilizando un operador de agregación, en este caso hemos utilizado la media aritmética. Al final obtendremos los siguientes valores de preferencia numéricos para el valor de utilidad inicial:

	x_1	x_2	x_3	x_4
Valoración	0.92	0.29	0.42	0.70

3. Transformamos los numéricos al $F(S_T)$ tal y como hemos visto en la sección anterior.

Y obtendremos los siguientes resultados:

	x_1	x_2
Valoración	$\{0, 0, 0, 0, 0, 0.48, 0.52\}$	$\{0, 0.26, 0.73, 0, 0, 0, 0\}$
	x_3	x_4
Valoración	$\{0, 0, 0.46, 0.54, 0, 0, 0\}$	$\{0, 0, 0, 0, 0.83, 0.17, 0\}$

Comentario: El uso de esta transformación para información intervalar en $F(S_T)$ se ha mostrado útil en aquellos problemas de decisión en los que los expertos expresan su información mediante vectores de utilidad, pero no resulta apropiado su uso en problemas donde la información es expresada mediante relaciones de preferencia. Por lo que, a continuación mostraremos un modelo de transformación de información intervalar a $F(S_T)$ genérico para cualquier tipo de estructura de representación de información. Esta función de transformación será la que utilizaremos en el resto de la memoria por mostrarse más general y fácil de implementar y modelar.

Transformación de Intervalos en $F(S_T)$ Basada en una Representación Difusa

Nuestro objetivo es el mismo que en el apartado anterior, transformar un intervalo en $F(S_T)$. En este caso, dado un intervalo $I = [\underline{i}, \bar{i}]$ valorado en $[0, 1]$, queremos realizar una transformación en la que asumiremos que el intervalo tiene

una representación inspirada en una función de pertenencia de conjuntos difusos como la siguiente [73]:

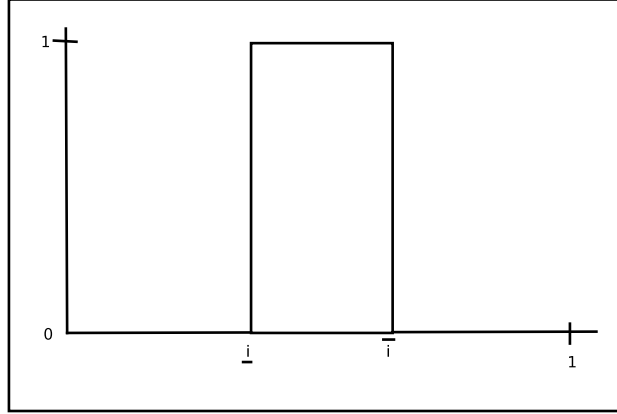


Figura 2.10: Función de pertenencia de un intervalo

$$\mu_I(\vartheta) = \begin{cases} 0, & \text{si } \vartheta < i \\ 1, & \text{si } i \leq \vartheta \leq \bar{i} \\ 0, & \text{si } \bar{i} < \vartheta \end{cases}$$

Donde ϑ es un valor en $[0, 1]$. En la figura 2.10 podemos observar la representación gráfica de un intervalo.

Para realizar la transformación utilizaremos la siguiente función de transformación:

Definición 2.6. Sea $S_T = \{s_0, \dots, s_g\}$ el CBTL. Entonces, la función τ_{IS_T} transforma un intervalo I en $[0, 1]$ en un conjunto difuso en S_T .

$$\begin{aligned} \tau_{IS_T} : I &\rightarrow F(S_T) \\ \tau_{IS_T}(I) &= \{(c_k, \gamma_k^i) / k \in \{0, \dots, g\}\} \\ \gamma_k^i &= \max_y \min \{\mu_I(y), \mu_{c_k}(y)\} \end{aligned} \tag{2.3}$$

donde $F(S_T)$ es el conjunto de conjuntos difusos definidos en S_T , y $\mu_I(\cdot)$ y $\mu_{c_k}(\cdot)$ son las funciones de pertenencia asociadas con el intervalo I y los términos c_k , respectivamente.

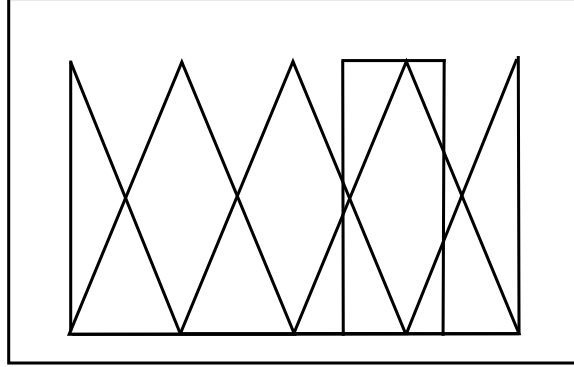


Figura 2.11: Ejemplo de transformación de un intervalo

Ejemplo

Sea $I = [0.6, 0.78]$ un intervalo que debe ser transformado en un conjunto difuso en S_T con cinco términos simétricamente distribuidos. El conjunto difuso obtenido después de aplicar τ_{IS_T} es (figura 2.11):

$$\tau_{IS_T}([0.6, 0.78]) = \{(s_0, 0), (s_1, 0), (s_2, 0.6), (s_3, 1), (s_4, 0.2)\}$$

2.2.2.4. Transformación de Información Lingüística en $F(S_T)$

A continuación pasamos a presentar la definición para la función de transformación lingüística. En la figura 2.12 podemos ver gráficamente en que paso del proceso de decisión nos encontramos en este momento.

Cuando trabajamos con información lingüística en el contexto heterogéneo ocurre que aunque el dominio común de expresión para la información es el lingüístico, la representación es mediante conjuntos difusos, por lo que tenemos que

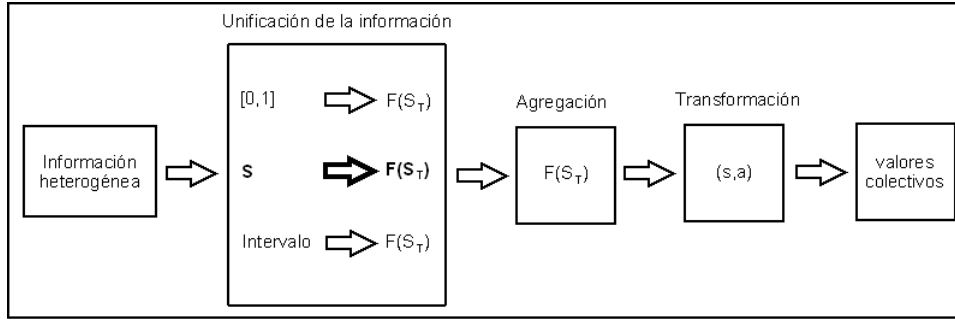


Figura 2.12: Transformación de información lingüística en $F(S_T)$

transformar las etiquetas lingüísticas de los conjuntos de etiquetas usados por los expertos en conjuntos difusos en el CBTL. Además, en el contexto de definición de nuestro problema, podemos utilizar diferentes conjuntos de etiquetas para valorar cada una de las alternativas.

La transformación de un conjunto de valoraciones lingüísticas definidas en S_i en $F(S_T)$ se realiza como sigue:

1. En el caso de que la etiqueta lingüística s_j que vamos a transformar a un $F(S_T)$ pertenezca a un conjunto de etiquetas S_i tal que sea escogido como CBTL, es decir $S_i = S_T$, entonces el conjunto difuso que representará un término lingüístico será en todos 0 excepto en el valor correspondiente al ordinal, j , de la etiqueta lingüística, $s_j \in S_i$ que será uno.
2. En caso contrario, S_i y S_T son distintos, y por tanto no coinciden ni en granularidad, ni en sintaxis, ni en semántica. Para transformar las etiquetas $s_j \in S_i$ en $F(S_T)$ utilizaremos la siguiente función de transformación:

Definición 2.7. Sea $S = \{l_0, \dots, l_p\}$ y $S_T = \{s_0, \dots, s_g\}$ dos conjuntos de términos lingüísticos. Entonces, una función de transformación lingüísticas en conjuntos difusos sobre el CBTL, $\tau_{S S_T}$, se define como:

$$\begin{aligned}
\tau_{SS_T} : S &\rightarrow F(S_T) \\
\tau_{SS_T}(l_i) &= \{(c_k, \gamma_k^i) / k \in \{0, \dots, g\}\}, \forall l_i \in S \\
\gamma_k^i &= \max_y \min \{\mu_{l_i}(y), \mu_{c_k}(y)\}
\end{aligned} \tag{2.4}$$

donde $F(S_T)$ es el conjunto de conjuntos difusos definidos en S_T , y $\mu_{l_i}(\cdot)$ y $\mu_{c_k}(\cdot)$ son funciones de pertenencia de los conjuntos difusos asociados con los términos l_i y c_k respectivamente.

Por lo tanto, el resultado de τ_{SS_T} para cualquier valor lingüístico de S es un conjunto difuso definido en las CBTL, S_T .

Ejemplo

Sea $S = \{l_0, l_1, \dots, l_4\}$ y $S_T = \{s_0, s_1, \dots, s_6\}$ dos conjuntos de términos, con 5 y 7 etiquetas, respectivamente, y con la siguiente semántica asociada:

$$\begin{aligned}
l_0 &= (0, 0, 0.25) & s_0 &= (0, 0, 0.16) \\
l_1 &= (0, 0.25, 0.5) & s_1 &= (0, 0.16, 0.34) \\
l_2 &= (0.25, 0.5, 0.75) & s_2 &= (0.16, 0.34, 0.5) \\
l_3 &= (0.5, 0.75, 1) & s_3 &= (0.34, 0.5, 0.66) \\
l_4 &= (0.75, 1, 1) & s_4 &= (0.5, 0.66, 0.84) \\
& & s_5 &= (0.66, 0.84, 1) \\
& & s_6 &= (0.84, 1, 1)
\end{aligned}$$

El conjunto difuso obtenido después de aplicar τ_{SS_T} para l_1 es (ver figura 2.13):

$$\tau_{SS_T}(l_1) = \{(s_0, 0.39), (s_1, 0.85), (s_2, 0.85), (s_3, 0.39), (s_4, 0), (s_5, 0), (s_6, 0)\}$$

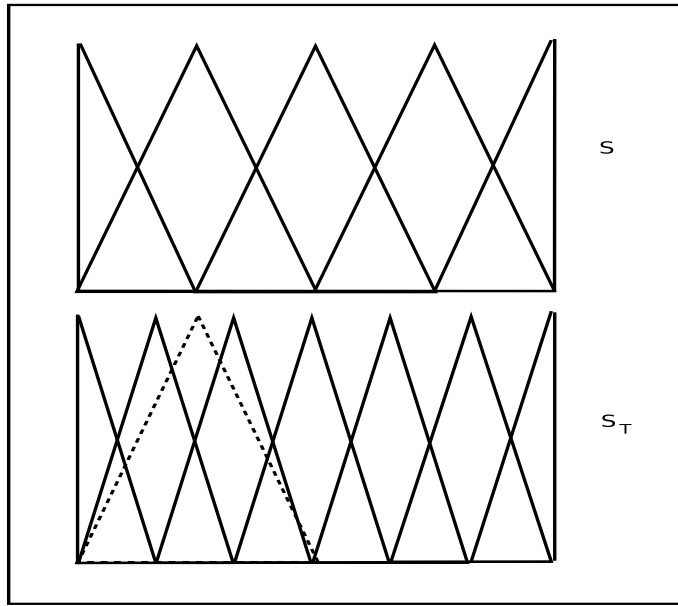


Figura 2.13: Transformación $l_1 \in S$ en un conjunto difuso en S_T

2.2.2.5. Transformación de Conjuntos Difusos sobre el CBTL en 2-tupla Lingüísticas del CBTL

A continuación pasamos a presentar la definición para la función de transformación de conjuntos difusos en una 2-tupla lingüística. En la figura 2.14 podemos ver gráficamente en que paso del proceso de decisión nos encontramos en este momento.

Una vez la información heterogénea está expresada de forma uniforme podemos utilizar un operador de agregación para obtener los valores colectivos para cada una de las alternativas. Los valores colectivos expresados en conjuntos difusos sobre el CBTL son útiles para resolver el problema de decisión, pero presentan dos problemas fundamentalmente:

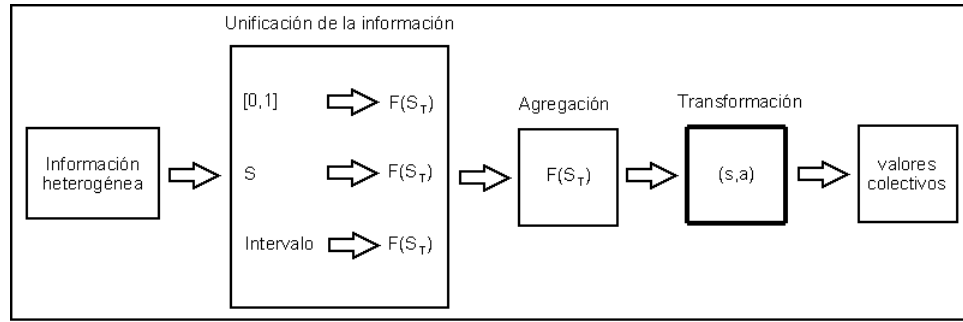


Figura 2.14: Transformación de conjuntos difusos sobre el CBTL en 2-tupla lingüísticas

- i Su orden no es directo por lo que, según el *ranking* para conjuntos difusos utilizado puede dar diferentes resultados, y
- ii No son valores fáciles de entender por los expertos.

Por tanto para mejorar estos dos aspectos presentamos una función que transformará los conjuntos difusos sobre el CBTL, que expresan los valores colectivos para cada una de las alternativas, en 2-tuplas lingüísticas pertenecientes al CBTL. Esta transformación nos permite superar los dos problemas anteriores ya que el modelo de representación de la información por medio de 2-tuplas lingüísticas tiene definido un orden total [68] que ayudará al proceso de explotación a seleccionar una solución. Además esta representación es fácil de entender por cualquier experto independientemente de cuál sea su dominio de expresión inicial.

En [60] se presentó una función que transforma un conjunto difuso en un valor numérico en el intervalo de granularidad de S_T , $[0, g]$. En este capítulo presentaremos una nueva función χ que transformará directamente los conjuntos difusos en 2-tuplas lingüísticas:

Definition 2.8.

$$\begin{aligned}\chi : F(S_T) &\rightarrow S_T \times [-0.5, 0.5) \\ \chi(F(S_T)) &= \chi(\{(s_j, \gamma_j), j = 0, \dots, g\}) = \Delta \left(\frac{\sum_{j=0}^g j \gamma_j}{\sum_{j=0}^g \gamma_j} \right) = \\ &= \Delta(\beta) = (s, \alpha)\end{aligned}$$

Ejemplo

Partimos de un conjunto difuso que representa el valor colectivo obtenido para una alternativa en un problema de TD:

$$((N, 0), (VL, 0), (L, 0), (M, .41), (H, 1), (VH, .19), (P, 0))$$

definido en siguiente CBTL (ver figura 2.15):

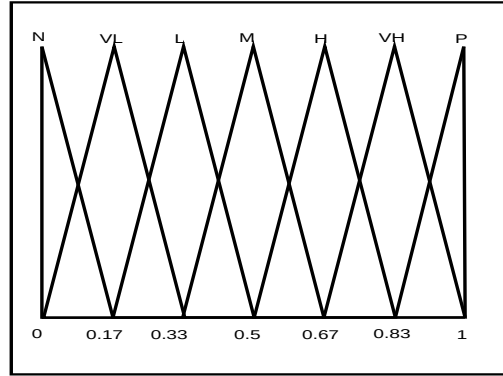


Figura 2.15: CBTL

Y queremos expresarlo mediante una 2-tupla, para ello aplicaremos la función χ .

$$\chi((0, 0, 0, .41, 1, .19, 0)) = \Delta \left(\frac{\sum_{j=0}^6 j \gamma_j}{\sum_{j=0}^6 \gamma_j} \right) = \Delta(4.33) = (H, .33)$$

Por lo tanto, nuestra 2-tupla solución es: $(H, .33)$.

2.3. Conclusiones

En este capítulo hemos definido las funciones necesarias que nos ayudarán a manejar la información heterogénea (numérica, intervalar y lingüística) en problemas de toma de decisión definidos en un contexto heterogéneo. Ahora ya podemos presentar el modelo de decisión para resolver este tipo de problemas en el próximo capítulo de esta memoria.

Capítulo 3

Modelo de Decisión para Problemas de Toma de Decisión en Contextos Heterogéneos

Una vez presentadas, en el capítulo 2 de esta memoria, las funciones necesarias para el tratamiento de la información heterogénea el objetivo de este capítulo es proponer un modelo de decisión para resolver un problema de toma de decisiones con Múltiples Expertos (TDME) o toma de decisión en Grupo (TDG) definidos en contextos heterogéneos donde las valoraciones de los expertos sobre las alternativas pueden estar expresadas en los dominios numérico, intervalar, y lingüístico.

3.1. El Modelo de Decisión

El modelo de decisión que proponemos para resolver este tipo de problemas sigue el esquema presentado en la figura 3.1 presentada anteriormente en el capítulo 2 de esta memoria. Una vez definidas las herramientas necesarias para poder llevar

acabo las operaciones con información heterogénea revisaremos todas las fases y finalmente mostraremos su aplicación a un problema de TDME en la instalación de un sistema ERP para una empresa:

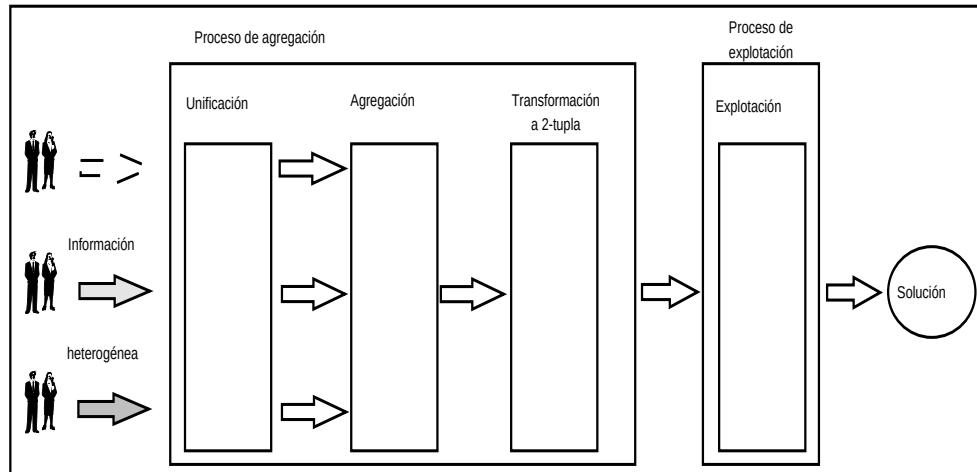


Figura 3.1: Modelo de decisión en contexto heterogéneo

1. *Adquisición de la información:* Los distintos expertos suministrarán sus preferencias en el dominio de información más adecuado a su grado de conocimiento y a la naturaleza de las alternativas del problema.
2. *Proceso de agregación:* Obtiene un valor colectivo para cada alternativa, en problemas definidos en contextos heterogéneos realiza los siguientes pasos:
 - a) *Unificación de información heterogénea:* La información será unificada en conjuntos difusos sobre un conjunto básico de términos lingüísticos (CBTL).
 - b) *Proceso de agregación de los valores de las preferencias individuales:* Obtendremos una valoración colectiva de cada alternativa en conjuntos difusos sobre el CBTL a partir de los valores unificados utilizando

operadores de agregación adecuados al problema de decisión.

- c) *Transformación a 2-tupla*: transformaremos los conjunto difusos valorados en el CBTL a 2-tuplas valoradas en el CBTL [69], para facilitar el proceso de explotación y mejorar la comprensión de los resultados finales.

3. *Proceso de explotación*: A partir de las valoraciones globales evaluadas mediante 2-tuplas construiremos la solución del problema.

3.1.1. Adquisición de la Información

A continuación describiremos cada una de las fases y utilizaremos un ejemplo simple para explicar, de forma práctica, el funcionamiento de éstas. Para ello introduciremos el ejemplo que se utilizará en la explicación donde los expertos suministran sus preferencias en un dominio de información que les resulte más adecuado a su grado de conocimiento sobre las distintas alternativas del problema:

Ejemplo

Supongamos que una compañía quiere renovar sus coches. Existen cuatro modelos de coches disponibles, $\{CAR1, CAR2, CAR3, CAR4\}$ y tres expertos nos proporcionan sus relaciones de preferencia sobre los cuatro coches. El primer experto expresa sus relaciones de preferencia utilizando valores numéricos en $[0, 1]$, P_1^n . El segundo expresa sus preferencias por medio de valores lingüísticos en un conjunto de términos S (ver figura 3.2), P_2^S . Y el tercer experto los expresa utilizando valores de preferencia intervalares en $[0, 1]$, P_3^I . Los tres expertos quieren obtener una solución colectiva.

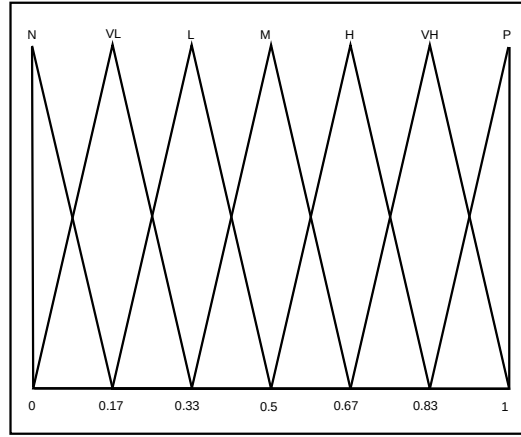


Figura 3.2: Conjunto S utilizado

$$P_1^n = \begin{pmatrix} - & .5 & .8 & .4 \\ .3 & - & .9 & .3 \\ .3 & .2 & - & .4 \\ .9 & .8 & .5 & - \end{pmatrix} \quad P_2^S = \begin{pmatrix} - & H & VH & M \\ L & - & H & VH \\ VL & N & - & VH \\ L & VL & N & - \end{pmatrix}$$

$$P_3^I = \begin{pmatrix} - & [.7, .8] & [.65, .7] & [.8, .9] \\ [.3, .35] & - & [.6, .7] & [.8, .85] \\ [.3, .35] & [.3, .4] & - & [.7, .9] \\ [.1, .2] & [.2, .4] & [.1, .3] & - \end{pmatrix}$$

3.1.2. Proceso de Agregación

Esta fase tiene como objeto obtener un valor colectivo para cada una de las alternativas valoradas por los distintos expertos en los diferentes dominios (numérico, intervalar y lingüístico). Para ello se efectuarán los pasos de:

1. Unificación de la información heterogénea.

2. Proceso de agregación de los valores de las preferencias individuales.
3. Transformación a 2-tupla de los valores colectivos.

3.1.2.1. Unificación de la Información

El objetivo de este proceso es expresar la información inicial aportada por los expertos en un único dominio de expresión. Esta información heterogénea puede ser numérica, intervalar, o lingüística y en nuestro problema se representará mediante vectores de utilidad o relaciones de preferencia. Para poder operar con toda esta información, debemos expresarla en un dominio de expresión unificado. El primer paso para realizar este proceso es definir el CBTL y a continuación utilizando las funciones presentadas en el capítulo anterior. Al final de esta fase, la información de entrada estará expresada mediante conjuntos difusos en un dominio de expresión unificado, el CBTL, $S_T = \{s_0, \dots, s_g\}$.

El proceso se realiza siguiendo el siguiente orden:

1. Selección del CBTL.
2. Transformamos los valores numéricos del intervalo $[0, 1]$ al $F(S_T)$.
3. Transformamos los valores lingüísticos al $F(S_T)$.
4. Transformamos los valores intervalares al $F(S_T)$.

En el siguiente ejemplo podemos ver como transformar la información proporcionada por los expertos al CBTL.

Ejemplo

Ahora realizaremos la unificación de la información de entrada. Para ello realizaremos los siguientes pasos:

1. Selección del CBTL. En este ejemplo, será S , dado que cumple las condiciones necesarias presentadas en la sección 2.2.1 del capítulo anterior (ver figura 3.3).

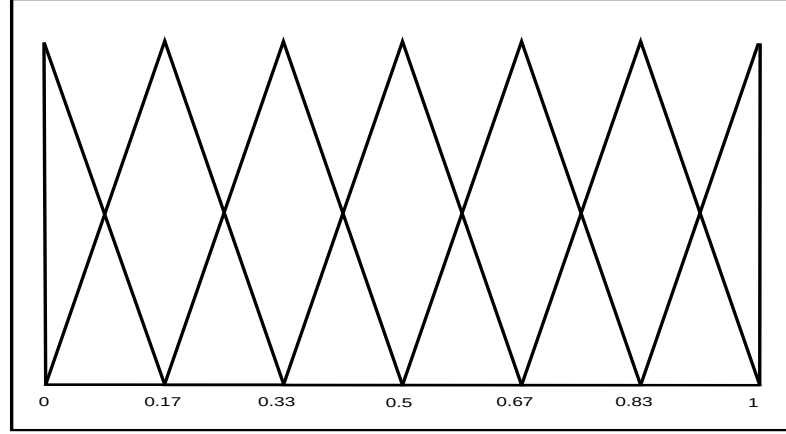


Figura 3.3: CBTL Utilizado

2. Transformación de la información de entrada en $F(S_T)$. Aplicando las funciones de transformación presentadas en el capítulo anterior.

Por simplicidad de la notación los conjuntos difusos $((s_o, \gamma_0), \dots, (s_n, \gamma_n))$ se mostrarán como $(\gamma_0, \dots, \gamma_n)$.

$$P_1^n = \begin{pmatrix} - & (0, 0, 0, 1, 0, 0, 0) & (0, 0, 0, 0, .19, .81, 0) & (0, 0, .59, .41, 0, 0, 0) \\ (0, .19, .81, 0, 0, 0, 0) & - & (0, 0, 0, 0, 0, .59, .41) & (0, .19, .81, 0, 0, 0, 0) \\ (0, .19, .81, 0, 0, 0, 0) & (0, .81, .19, 0, 0, 0, 0) & - & (0, 0, .59, .41, 0, 0, 0) \\ (0, 0, 0, 0, 0, .59, .41) & (0, 0, 0, 0, .19, .81, 0) & (0, 0, 0, 1, 0, 0, 0) & - \end{pmatrix}$$

$$P_2^S = \begin{pmatrix} - & (0, 0, 0, 0, 1, 0, 0) & (0, 0, 0, 0, 0, 1, 0) & (0, 0, 0, 1, 0, 0, 0) \\ (0, 0, 1, 0, 0, 0, 0) & - & (0, 0, 0, 0, 1, 0, 0) & (0, 0, 0, 0, 0, 1, 0) \\ (0, 1, 0, 0, 0, 0, 0) & (1, 0, 0, 0, 0, 0, 0) & - & (0, 0, 0, 0, 0, 1, 0) \\ (0, 0, 1, 0, 0, 0, 0) & (0, 1, 0, 0, 0, 0, 0) & (1, 0, 0, 0, 0, 0, 0) & - \end{pmatrix}$$

$$P_3^I = \begin{pmatrix} - & (0, 0, 0, 0, .81, .81, 0) & (0, 0, 0, .12, 1, .19, 0) & (0, 0, 0, 0, .19, 1, .41) \\ (0, .19, 1, .12, 0, 0, 0) & - & (0, 0, 0, .41, 1, .19, 0) & (0, 0, 0, 0, .19, 1, .12) \\ (0, .19, 1, .12, 0, 0, 0) & (0, .19, 1, .41, 0, 0, 0) & - & (0, 0, 0, 0, .81, 1, .41) \\ (.41, 1, .19, 0, 0, 0, 0) & (0, .81, 1, .41, 0, 0, 0) & (.41, 1, .81, 0, 0, 0, 0) & - \end{pmatrix}$$

3.1.2.2. Proceso de Agregación de los Valores de las Preferencias Individuales

Una vez aplicadas las funciones de transformación hemos logrado que la información de entrada este expresada en un único dominio de expresión por medio de conjuntos difusos en el CBTL, $S_T = \{s_0, \dots, s_g\}$. El siguiente paso en el proceso de decisión es agregar estos valores individuales de preferencia para obtener una valoración colectiva de cada alternativa. Para ello, usaremos una función de agregación para combinar los conjuntos difusos en el CBTL para obtener una valoración colectiva para cada alternativa que sea un conjunto difuso en el CBTL.

Para la toma de decisión en grupo con información heterogénea las relaciones de preferencia individuales unificadas, P_{e_k} , se expresan por medio de un conjunto difuso en el CBTL de la siguiente manera:

$$\begin{pmatrix} p_{11}^k = \{(s_0, \gamma_0)_k^{11}, \dots, (s_g, \gamma_g)_k^{11}\} & \dots & p_{1n}^k = \{(s_0, \gamma_0)_k^{1n}, \dots, (s_g, \gamma_g)_k^{1n}\} \\ \vdots & \dots & \vdots \\ p_{n1}^k = \{(s_0, \gamma_0)_k^{n1}, \dots, (s_g, \gamma_g)_k^{n1}\} & \dots & p_{nn}^k = \{(s_0, \gamma_0)_k^{nn}, \dots, (s_g, \gamma_g)_k^{nn}\} \end{pmatrix},$$

donde p_{ij}^k es el grado de preferencia de la alternativa x_i sobre x_j proporcionada por el experto e_k y donde $(s_v, \gamma_v)_k^{ij}$ representa el grado de pertenencia γ_v sobre el término lingüístico $s_v \in S_T$ del CBTL de la preferencia de la alternativa x_i sobre la alternativa x_j dada por el experto e_k .

Representaremos cada conjunto difuso, p_{ij}^k , como $r_{ij}^k = (\gamma_0, \dots, \gamma_g)_k^{ij}$ siendo los valores de r_{ij}^k su respectivos grados de pertenencia. Entonces, la valoración

colectiva de la relación de preferencia de acuerdo a todas las relaciones de preferencia proporcionados por los expertos $\{r_{ij}^k, \forall e_k\}$ se obtiene agregando estos conjuntos difusos. Esta valoración de preferencia colectiva, se denotará por r_{ij} , está compuesta por conjuntos difusos definidos en S_T :

$$r_{ij} = (\gamma_0, \dots, \gamma_v, \dots, \gamma_g)^{ij},$$

donde γ_v de r_{ij} se caracteriza por la siguiente función de pertenencia:

$$\gamma_v = f(\gamma_{v_1}, \dots, \gamma_{v_k}, \dots, \gamma_{v_m}),$$

donde f es un operador de agregación, m es el número de expertos, y γ_{v_k} representa el grado de pertenencia sobre el término lingüístico $s_v \in S_T$ cuando se está valorando la preferencia de x_i sobre x_j por parte del experto e_k .

Ejemplo

Continuando el ejemplo anterior, ahora realizaremos la agregación de los valores de preferencias individuales. Cuando toda la información se expresa por medio de conjuntos difusos definidos en un CBTL utilizaremos un operador de agregación para combinarlo. En este ejemplo usaremos como operador de agregación, f , la media aritmética obteniéndose la siguiente relación de preferencia colectiva, P_C :

$$\begin{pmatrix} - & (0, 0, 0, .33, .6, .27, 0) & (0, 0, 0, .04, .4, .19, 0) & (0, 0, .2, .47, .06, .33, .04) \\ (0, .13, .94, .04, 0, 0, 0) & - & (0, 0, 0, .14, .67, .26, .14) & (0, .06, .27, 0, .06, .67, .04) \\ (0, .46, .6, .04, 0, 0, 0) & (.33, .33, .4, .14, 0, 0, 0) & - & (0, 0, .2, .14, .27, .67, .14) \\ (.14, .33, .4, 0, 0, .20, .14) & (0, .6, .33, .14, .06, .27, 0) & (.47, .33, .27, .33, 0, 0, 0) & - \end{pmatrix}$$

3.1.2.3. Transformación a 2-tupla

Una vez obtenida la relación de preferencia colectiva, tal que, cada preferencia está expresada en conjuntos difusos en S_T , para mejorar su comprensión y su

manipulación matemática transformaremos los conjuntos difusos en el CBTL en 2-tuplas lingüísticas en el CBTL con el objetivo de facilitar el proceso de ordenación para la fase de explotación del proceso de decisión.

Ejemplo

Partimos de las valoraciones globales obtenidas en el proceso anterior:

$$\begin{pmatrix} - & (0, 0, 0, .33, .6, .27, 0) & (0, 0, 0, .04, .4, .19, 0) & (0, 0, .2, .47, .06, .33, .04) \\ (0, .13, .94, .04, 0, 0, 0) & - & (0, 0, 0, .14, .67, .26, .14) & (0, .06, .27, 0, .06, .67, .04) \\ (0, .46, .6, .04, 0, 0, 0) & (.33, .33, .4, .14, 0, 0, 0) & - & (0, 0, .2, .14, .27, .67, .14) \\ (.14, .33, .4, 0, 0, .20, .14) & (0, .6, .33, .14, .06, .27, 0) & (.47, .33, .27, .33, 0, 0, 0) & - \end{pmatrix}$$

Nuestro objetivo es transformar los valoraciones globales representadas mediante conjuntos difusos en el CBTL a 2-tupla mediante la función de transformación vista en el capítulo anterior. Así, por ejemplo:

$$r_{12} = (\chi((0, 0, 0, .33, .6, .27, 0))) = \Delta \left(\frac{\sum_{j=0}^6 j\gamma_j}{\sum_{j=0}^6 \gamma_j} \right) = \Delta(3.96) = (H, -.04)$$

$$r_{13} = (\chi((0, 0, 0, .04, .4, .19, 0))) = \Delta \left(\frac{\sum_{j=0}^6 j\gamma_j}{\sum_{j=0}^6 \gamma_j} \right) = \Delta(4.24) = (H, .24)$$

$$r_{14} = (\chi((0, 0, .2, .47, .06, .33, 0))) = \Delta \left(\frac{\sum_{j=0}^6 j\gamma_j}{\sum_{j=0}^6 \gamma_j} \right) = \Delta(3.58) = (H, -.42)$$

$$r_{21} = (\chi((0, .13, .94, .04, 0, 0, 0))) = \Delta \left(\frac{\sum_{j=0}^6 j\gamma_j}{\sum_{j=0}^6 \gamma_j} \right) = \Delta(1.92) = (L, -0.08)$$

...

Transformados todos los conjuntos difusos que expresan las preferencias colectivas en 2-tuplas lingüísticas se obtiene como relación de preferencia colectiva, P_C :

$$\begin{pmatrix} - & (H, -.04) & (H, .24) & (H, -.42) \\ (L, -.08) & - & (H, .33) & (H, .03) \\ (L, -.38) & (VL, .29) & - & (H, .29) \\ (L, .45) & (L, .34) & (VL, .33) & - \end{pmatrix}$$

3.1.3. Proceso de Explotación

Llegados a esta fase, el modelo de decisión tiene para cada alternativa una valoración global expresada por medio de 2-tuplas lingüísticas en el CBTL. Para obtener la mejor alternativa o el mejor conjunto de alternativas podemos utilizar algún grado de selección de los presentados en [113, 123]. Sabemos que el modelo de representación de las 2-tuplas lingüísticas tiene definido un orden total sobre si mismo. Si obtenemos el valor global de cada alternativa podremos ordenar los valores globales [63] de las alternativas utilizando este orden y obtener la solución de nuestro problema.

Ejemplo

Para obtener el conjunto de alternativas solución podemos utilizar la función de selección que calcule el grado de dominancia de cada alternativa, x_i , sobre el resto de alternativas.

$$\Lambda(x_i) = \frac{1}{n-1} \sum_{j=0|j \neq i}^n \beta_{ij},$$

donde n es el número de alternativas y $\beta_{ij} = \Delta^{-1}(p_{ij})$ siendo p_{ij} una 2-tupla lingüística que representa el valor colectivo de la preferencia de la alternativa x_i sobre x_j de acuerdo al grupo de expertos. Así:

$$\Lambda(CAR1) = \frac{1}{3} \sum_{j=0|j \neq 0}^3 \beta_{ij} = 3.92$$

$$\Lambda (CAR2) = \frac{1}{3} \sum_{j=0|j \neq 1}^3 \beta_{ij} = 3.43$$

$$\Lambda (CAR3) = \frac{1}{3} \sum_{j=0|j \neq 2}^3 \beta_{ij} = 2.4$$

$$\Lambda (CAR4) = \frac{1}{3} \sum_{j=0|j \neq 3}^3 \beta_{ij} = 2.04$$

La solución al problema será aquella que tenga el grado de dominancia más alto. Expresaremos las soluciones obtenidas con 2-tuplas aplicando el operador Δ al resultado obtenido por el operador Λ .

CAR1	CAR2	CAR3	CAR4
(H,-.08)	(M,.43)	(L,.4)	(L.04)

Cuadro 3.1: Valores globales de las alternativas

Y devolveremos como solución *CAR1* dado que es la alternativa con el valor más alto de dominancia.

3.2. Estudio de TDME sobre la adecuación de la instalación de un ERP

En esta sección aplicaremos el modelo de decisión anterior para un problema de TDME definido en un contexto heterogéneo en el que se estudia la adecuación de instalar un ERP (Enterprise Resource Planing) en una empresa [129].

3.2.1. Introducción: Enterprise Resource Planning

Un ERP es un sistema estructurado para la optimización de una cadena de valor interna de una compañía. El software, está instalado completamente a través de

toda la empresa, conectando los componentes de la empresa a través de transmisiones lógicas y compartiendo datos comunes con un ERP integrado. Cuando datos tales como, una venta llegan a estar disponibles en un punto del negocio, toman su camino a través del software, que automáticamente calcula el efecto de la transacción en otras áreas, tales como fabricación, inventario, consecución, facturación, y reserva de la venta actual en el libro de contabilidad [110, 131].

Lo que realmente hace el ERP es organizar, codificar y estandarizar los procesos y datos de los negocios de la empresa. El software transforma los datos transaccionales en información útil y compagina los datos para que puedan ser analizados. De esta forma, todos los datos transaccionales recogidos se convierten en información que las compañías pueden utilizar para apoyar sus decisiones en los negocios. Cuando un sistema ERP está completamente introducido en la organización de la empresa, puede producir muchos beneficios: reducir el tiempo de ciclo, proporcionar transacciones de información más rápidas, facilitar una mejor gestión financiera, servir como un recurso más de trabajo para el comercio electrónico, y hacer que el conocimiento tácito sea explícito.

El software ERP no es intrínsecamente estratégico; más bien, está proporcionando una tecnología que es un conjunto de módulos de software integrados que componen el núcleo del procesamiento de transacciones internas. La instalación de un ERP, implica una gran inversión, debido a que se requieren cambios importantes en los procesos de organización, culturales y de negocio. Los cambios más importantes son aquellos referidos al papel de los individuos dentro de la organización. Muchos de los productos ERP han obligado a las compañías a rediseñar sus procesos de negocio para evitar realizar tareas inútiles y reasignando a los trabajadores a otras nuevas actividades más adecuadas incrementando la productividad de forma sistemática y por lo tanto sus beneficios.

Estas mejoras han hecho que las organizaciones de todo el mundo y las compa-

ñas de mediano y pequeño tamaño estén interesadas en la instalación de este tipo de producto. Sin embargo, la conveniencia de un ERP no es siempre provechosa. La causa de esto es que los sistemas ERP son muy complejos y tienen un gran impacto en toda la organización. Implementar un sistema ERP es siempre muy caro y lento de instalar, además la productividad y los beneficios de la compañía puede que no aumenten de forma drástica en algunos casos tal y como cabría esperar. Por lo tanto, antes de instalar un ERP debemos de evaluar su conveniencia en cada compañía, analizando un conjunto de parámetros de la organización para decidir la viabilidad de la implementación del ERP [101]. En este ejemplo utilizaremos un modelo de evaluación difuso basado en el modelo de decisión presentado anteriormente que tratará con información heterogénea y que estudiará la conveniencia de la instalación de un ERP de acuerdo a los diferentes parámetros de cada compañía.

3.2.2. Adecuación de un Sistema ERP: Proceso de Decisión

El proceso de decisión sobre la adecuación de instalar un sistema ERP en una compañía consiste en considerar las opiniones proporcionadas por varios expertos sobre una serie de atributos propios de una empresa [101]. De esta forma, este problema puede ser modelado como un problema de toma de decisión multiexperto-multiatributo (TDME-MA). En un problema TDME-MA, un conjunto finito de expertos $E = \{e_1, \dots, e_n\}$ evalúan las m alternativas $X = \{x_1, \dots, x_m\}$ atendiendo a distintos atributos $At = \{at_1, \dots, at_l\}$ mediante vectores de utilidad:

$$e_i \rightarrow \{p_{j1}^i, \dots, p_{jl}^i\},$$

donde p_{jk}^i ($i \in \{1, \dots, n\}$, $j \in \{1, \dots, m\}$, $k \in \{1, \dots, l\}$) será la valoración asignada por el experto e_i al atributo at_k de la alternativa x_j . Debido a la naturaleza de los atributos los expertos podrán proporcionarnos información no homogénea

valorada en diferentes dominios, de tal forma que, los vectores de utilidad pueden ser valorados por medio de valores numéricos, intervalares y lingüísticos. Siendo el vector de utilidad denotado por:

$$\{p_{j1}^{id}, \dots, p_{jl}^{id}\}$$

donde p_{jk}^{id} es la preferencia asignada al atributo at_k de la alternativa x_j por el experto e_i y valorado en el dominio D^d , $d \in \{N, L, I\}$ Numérico, Lingüístico o Intervalar respectivamente.

3.2.3. TD sobre la Instalación de un ERP

Por simplicidad utilizaremos el modelo de decisión para una compañía dada (x_1), que está considerando la posibilidad de instalar un ERP. En este caso, donde se pretende mostrar un ejemplo simple, se tendrán en cuenta nueve atributos para la compañía (aunque en un caso real podrían llegar a ser varias decenas), valorados en diferentes dominios para la evaluación de la conveniencia del sistema ERP:

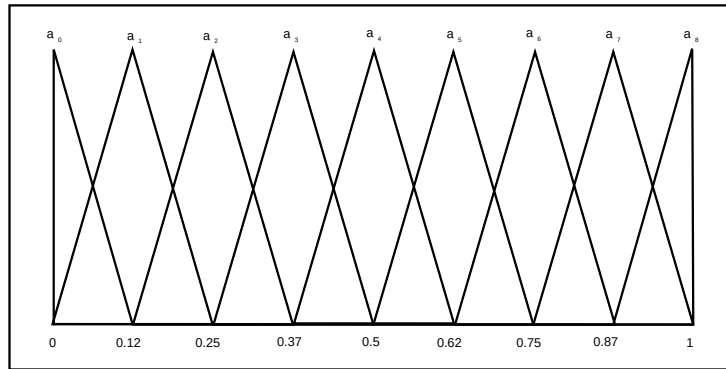
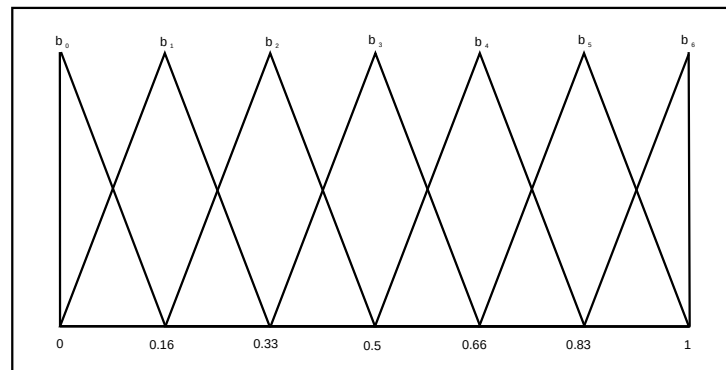
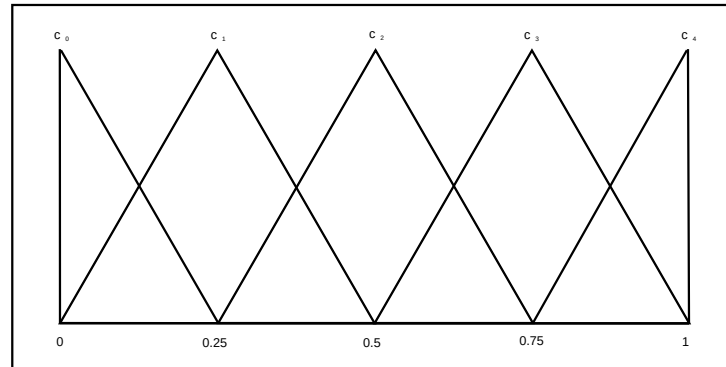
- at_1 : *Inversión en tecnologías de la información para los empleados*. Utilizaremos valores intervalares (el máximo valor será 6000)
- at_2 : *Precio de la implementación*. Utilizaremos un valor numérico (el máximo valor será 240000).
- at_3 : *Urgencia de la implementación*. Está valorado con términos lingüísticos y utilizaremos el conjunto de términos A .
- at_4 : *Grado de estandarización*. Está valorado en el conjunto de términos lingüísticos C .

- at_5 : *Interrelación con otros subsistemas*. Es un valor numérico valorado en $[0, 1]$.
- at_6 : *Capacidad del usuario para concretar*. Está valorado en el conjunto de términos C .
- at_7 : *Peticiones de cambio por el usuario*. Está valorado en el conjunto de términos B .
- at_8 : *Disponibilidad del personal*. Valorado en el conjunto de términos lingüísticos B .
- at_9 : *Capacidad de influencia del cliente en el proveedor*. Está valorado en el conjunto de términos lingüísticos D .

La semántica de los conjuntos de términos lingüísticos se muestra en la tabla 3.2 y en las figuras 3.4, 3.5, 3.6, 3.7:

Conjunto A		Conjunto B		Conjunto C		Conjunto D	
a_0	(0, 0, .12)	b_0	(0, 0, .16)	c_0	(0, 0, .25)	d_0	(0, 0, 0, 0)
a_1	(0, .12, .15)	b_1	(0, .16, .33)	c_1	(0, .25, .5)	d_1	(0, .01, .02, .07)
a_2	(.12, .25, .37)	b_2	(.16, .33, .5)	c_2	(.25, .5, .75)	d_2	(.4, .1, .18, .23)
a_3	(.25, .37, .5)	b_3	(.33, .5, .66)	c_3	(.5, .75, 1)	d_3	(.17, .41, .58, .65)
a_4	(.37, .5, .62)	b_4	(.5, .66, .83)	c_4	(.75, 1, 1)	d_4	(.32, .41, .58, .65)
a_5	(.5, .62, .75)	b_5	(.66, .83, 1)			d_5	(.58, .63, .80, .86)
a_6	(.62, .75, .87)	b_6	(.83, 1, 1)			d_6	(.72, .78, .92, .97)
a_7	(.75, .87, 1)					d_7	(.93, .98, .99, 1)
a_8	(.87, 1, 1)					d_8	(1, 1, 1, 1)

Cuadro 3.2: Semántica de los conjuntos de términos lingüísticos.

Figura 3.4: Semántica del conjunto de etiquetas A Figura 3.5: Semántica del conjunto de etiquetas B Figura 3.6: Semántica del conjunto de etiquetas C

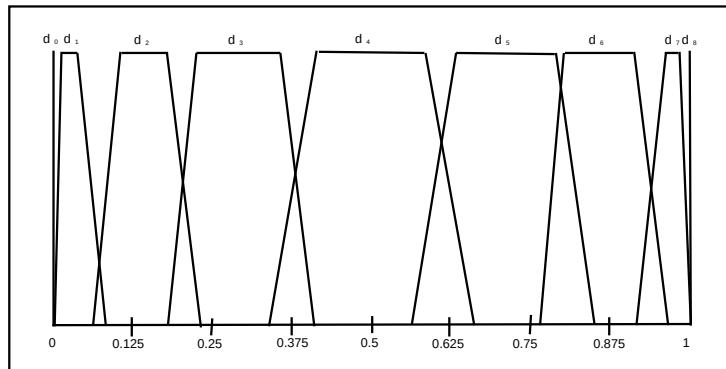


Figura 3.7: Semántica del conjunto de etiquetas D

En este ejemplo, cuatro expertos evalúan la conveniencia del ERP proporcionando sus preferencias para los atributos mediante de vectores de utilidad (ver tabla 3.3):

	e_1	e_2	e_3	e_4
at_1	[3500, 4000]	[2000, 2500]	[3100, 3800]	[4500, 5000]
at_2	12000	18000	10000	16000
at_3	a_5	a_6	a_5	a_4
at_4	c_2	c_2	c_3	c_1
at_5	.2	.35	.75	.3
at_6	c_1	c_1	c_2	c_3
at_7	b_3	b_4	b_3	b_4
at_8	b_4	b_5	b_5	b_3
at_9	d_1	d_6	d_5	d_5

Cuadro 3.3: Vectores de utilidad de los expertos.

Podemos ver que la interpretación de los atributos evaluados es distinta dependiendo de la semántica de cada atributo, así at_2 , at_5 , at_6 , at_7 y at_8 son atributos (precio de la implementación, interrelación con otros subsistemas, capacidad del usuario para concretar,...) que si tienen un valor alto, indican un grado bajo de aceptación (interpretación decreciente) mientras que el resto de atributos tiene una interpretación creciente, a mayor valor, mayor aceptación. Para facilitar las operaciones de agregación, debemos homogeneizar las interpretaciones de los atributos y tratarlos todos de forma creciente o decreciente. Por ello, transformaremos los atributos at_2 , at_5 , at_6 , at_7 y at_8 de forma que sus valoraciones y su interpretación sean crecientes y similar al del resto de los atributos. De esta manera los atributos siempre tendrán una interpretación creciente y podremos operar fácilmente sobre ellos. En la tabla 3.4 podemos ver los vectores de utilidad proporcionados por los expertos después de normalizar la información numérica y después de haber transformado los atributos en una interpretación creciente.

Aplicando el modelo de decisión presentado en la figura 3.1, realizaremos los siguientes pasos:

1. **Adquisición de la información.** El proceso ha sido descrito en las líneas anteriores. Donde hemos visto cómo los expertos suministran sus valoraciones sobre los distintos atributos. Normalizaremos las valoraciones numéricas e intervalares al intervalo $[0,1]$. Por último, y como ya se ha explicado antes, debemos conseguir una interpretación creciente para todos los atributos antes de pasar al siguiente paso, con lo que la información de entrada queda como podemos ver en la tabla 3.4.

2. **Proceso de agregación:**

- a) *Unificación de información:* Escogemos el CBTL. En este caso, de acuerdo a las reglas mostradas en la sección 2.2.1 del capítulo anterior,

	e_1	e_2	e_3	e_4
at_1	[.58, .67]	[.33, .42]	[.52, .63]	[.75, .83]
at_2	.5	.25	.58	.33
at_3	a_5	a_6	a_5	a_4
at_4	c_2	c_2	c_3	c_1
at_5	.8	.65	.25	.7
at_6	c_3	c_3	c_2	c_1
at_7	b_3	b_2	b_3	b_2
at_8	b_2	b_1	b_1	b_3
at_9	d_1	d_6	d_5	d_5

Cuadro 3.4: Interpretación normalizada y creciente. Vectores de utilidad de los expertos.

escogeremos como S_T un conjunto de términos distribuido simétricamente y de forma uniforme con 15 etiquetas (más detalles en [62], ver figura 3.8). Ahora aplicaremos las funciones de transformación a la información de entrada para unificarla (ver tablas 3.5, 3.6, 3.7, y 3.8). Estas funciones de transformación son las presentadas en el capítulo anterior por las ecuaciones (2.2), (2.3) , (2.4).

Por ejemplo, para transformar la información aportada por el experto 1 (ver tabla 3.4), debemos estudiar cada una de los atributos y ver a que dominio pertenece para saber que transformación debemos aplicar.

Así:

- 1) at_1 está expresada por medio de información intervalar, y su valoración actual es [.58, .67]. Por lo tanto, nosotros debemos aplicar:

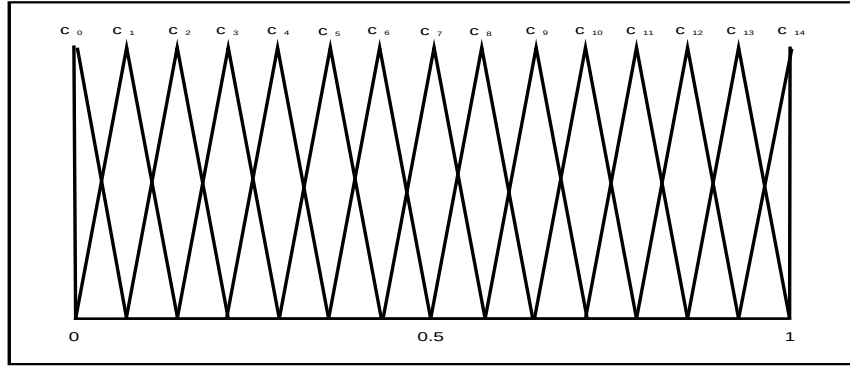


Figura 3.8: CBTL de 15 etiquetas distribuidas simétricamente.

$$\tau_{IS_T}([.58, .67]) = \{(c_k, \gamma_k^i) / k \in \{0, \dots, 14\}\}$$

Donde:

$$\gamma_k^i = \max_y \min \{\mu_{l_i}(y), \mu_{c_k}(y)\}$$

y obtendremos como resultado:

$$(0, 0, 0, 0, 0, 0, 0, .86, 1, .43, 0, 0, 0, 0, 0)$$

2) at_2 y at_5 son valoraciones numéricas. Por ejemplo, para calcular el valor de at_5 realizaremos:

$$\tau(0, 2) = \{(s_0, \gamma_0), \dots, (s_{14}, \gamma_{14})\}, s_i \in S_T \text{ y } \gamma_i \in [0, 1]$$

donde:

$$\gamma_i = \mu_{s_i}(\vartheta) = \begin{cases} 0, & \text{si } \vartheta \notin \text{Soporte}(\mu_{s_i}(x)) \\ \frac{\vartheta - a_i}{b_i - a_i}, & \text{si } a_i \leq \vartheta \leq b_i \\ \frac{c_i - \vartheta}{c_i - d_i}, & \text{si } b_i \leq \vartheta \leq c_i \end{cases}$$

con lo que obtendremos:

$$(0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, .71, .29, 0, 0)$$

De forma análoga resolveríamos at_2 .

- 3) Por último, at_3 , at_4 , at_6 , at_7 , at_8 y at_9 son atributos valorados lingüísticamente. Por ejemplo, para calcular at_3 tendremos que realizar los siguientes cálculos:

$$\tau_{SS_T}(a_5) = \{(c_k, \gamma_k^i) / k \in \{0, \dots, 14\}\}$$

Donde

$$\gamma_k^i = \max_y \min \{\mu_{l_i}(y), \mu_{c_k}(y)\}$$

Con lo que obtendremos:

$$(0, 0, 0, 0, 0, 0, 0, .36, .73, .89, .55, .2, 0, 0, 0)$$

Lo mismo haríamos con el resto de atributos lingüísticos at_4 , at_6 , at_7 , at_8 y at_9 .

Las tablas obtenidas después de la unificación son las 3.5, 3.6, 3.7, 3.8.

	e_1
at_1	(0, 0, 0, 0, 0, 0, 0, .86, 1, .43, 0, 0, 0, 0, 0)
at_2	(0, 0, 0, 0, 0, 0, 0, 1, 0, 0, 0, 0, 0, 0, 0)
at_3	(0, 0, 0, 0, 0, 0, 0, .36, .73, .89, .55, .2, 0, 0, 0)
at_4	(0, 0, 0, .12, .34, .56, .78, 1, .78, .56, .34, .12, 0, 0, 0)
at_5	(0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, .71, .29, 0, 0)
at_6	(0, 0, 0, 0, 0, 0, 0, .21, .43, .65, .87, .9, .68, .45, .21)
at_7	(0, 0, 0, 0, .12, .41, .7, 1, .69, .39, .08, 0, 0, 0, 0)
at_8	(0, 0, .24, .54, .83, .87, .58, .29, 0, 0, 0, 0, 0, 0, 0)
at_9	(0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, .58, .87)

Cuadro 3.5: Información unificada por el experto 1.

	e_2
at_1	(0, 0, 0, 0, .43, 1, .86, 0, 0, 0, 0, 0, 0, 0, 0)
at_2	(0, 0, 0, .57, .43, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0)
at_3	(0, 0, 0, 0, 0, 0, 0, .1, .45, .79, .84, .47, .09, 0)
at_4	(0, 0, 0, .12, .34, .56, .78, 1, .78, .56, .34, .12, 0, 0, 0)
at_5	(0, 0, 0, 0, 0, 0, 0, 0, 0, .86, .14, 0, 0, 0, 0)
at_6	(0, 0, 0, 0, 0, 0, 0, .21, .43, .65, .87, .9, .68, .45, .21)
at_7	(0, 0, .24, .54, .83, .87, .58, .29, 0, 0, 0, 0, 0, 0, 0)
at_8	(.3, .97, .95, .75, .45, .16, 0, 0, 0, 0, 0, 0, 0, 0, 0)
at_9	(0, 0, 0, 0, 0, 0, 0, 0, 0, .35, .76, 1, .92, .33)

Cuadro 3.6: Información unificada por el experto 2.

	e_3
at_1	(0, 0, 0, 0, 0, 0, 0, .71, 1, .86, 0, 0, 0, 0, 0)
at_2	(0, 0, 0, 0, 0, 0, 0, .86, .14, 0, 0, 0, 0, 0)
at_3	(0, 0, 0, 0, 0, 0, 0, .36, .73, .89, .55, .2, 0, 0, 0)
at_4	(0, 0, 0, 0, 0, 0, 0, .21, .43, .65, .87, .9, .68, .45, .21)
at_5	(0, 0, 0, 0, .57, .43, 0, 0, 0, 0, 0, 0, 0, 0)
at_6	(0, 0, 0, .12, .34, .56, .78, 1, .78, .56, .34, .12, 0, 0, 0)
at_7	(0, 0, 0, 0, .12, .41, .7, 1, .69, .39, .08, 0, 0, 0, 0)
at_8	(.3, .97, .95, .75, .45, .16, 0, 0, 0, 0, 0, 0, 0, 0)
at_9	(0, 0, 0, 0, 0, 0, 0, .5, 1, 1, 1, .61, .07, 0)

Cuadro 3.7: Información unificada por el experto 3.

	e_4
at_1	(0, 0, 0, 0, 0, 0, 0, 0, 0, .43, 1, .71, 0, 0)
at_2	(0, 0, 0, 0, .43, .57, 0, 0, 0, 0, 0, 0, 0, 0)
at_3	(0, 0, 0, 0, 0, .29, .65, 1, .63, .26, 0, 0, 0, 0)
at_4	(.21, .65, .68, .9, .87, .65, .43, .21, 0, 0, 0, 0, 0, 0)
at_5	(0, 0, 0, 0, 0, 0, 0, 0, 0, .14, .86, 0, 0, 0)
at_6	(.21, .65, .68, .9, .87, .65, .43, .21, 0, 0, 0, 0, 0, 0)
at_7	(0, 0, .24, .54, .83, .87, .58, .29, 0, 0, 0, 0, 0, 0)
at_8	(0, 0, 0, 0, .12, .41, .7, 1, .69, .39, .08, 0, 0, 0)
at_9	(0, 0, 0, 0, 0, 0, 0, .5, 1, 1, 1, .61, .07, 0)

Cuadro 3.8: Información unificada por el experto 4.

r_{11}	(0, 0, 0, 0, .10, .25, .21, .39, .5, .32, .10, .25, .17, 0, 0)
r_{12}	(0, 0, 0, .14, .21, .14, 0, .25, .21, .03, 0, 0, 0, 0, 0)
r_{13}	(0, 0, 0, 0, 0, .07, .16, .43, .54, .62, .47, .31, .11, .02, 0)
r_{14}	(.05, .16, .17, .28, .38, .44, .49, .60, .49, .44, .38, .28, .17, .11, .05)
r_{15}	(0, 0, 0, 0, .14, .10, 0, 0, 0, .25, .25, .17, .07, 0, 0)
r_{16}	(.05, .16, .17, .25, .30, .30, .30, .40, .41, .46, .52, .48, .34, .22, .10)
r_{17}	(0, 0, .12, .27, .47, .64, .64, .64, .34, .19, .04, 0, 0, 0, 0)
r_{18}	(.15, .48, .53, .51, .46, .4, .32, .32, .17, .09, .02, 0, 0, 0, 0)
r_{19}	(0, 0, 0, 0, 0, 0, 0, .25, .5, .58, .69, .55, .41, .3)

Cuadro 3.9: Información agregada

- b) *Agregación de los vectores de utilidad individuales.* Aquí aplicaremos como operador agregación la media aritmética (ver tabla 3.9), pero podemos usar otros operadores dependiendo de si consideramos todos los parámetros igualmente importantes o no. Obteniéndose un valor colectivo agregado, r_{jk} , que expresa el valor colectivo del atributo k para la alternativa j .
- c) *Transformar el vector de utilidad colectivo a 2-tupla:* Una vez obtenida una valoración colectiva, para cada atributo expresado mediante conjuntos difusos r_{jk} en el S_T , estos serán transformados a 2-tuplas lingüísticas en el CBTL tal y como sigue:

$$\begin{aligned}
 r_{11} &= \Delta(\chi((0, 0, 0, 0, .10, .25, .21, .39, .5, .32, .10, .25, .17, 0, 0))) = \\
 &= \Delta\left(\frac{\sum_{j=0}^{14} j\gamma_j}{\Sigma_{j=0}^{14}}\right) = \Delta(8) = (s_8, 0) \\
 r_{12} &= \Delta(\chi((0, 0, 0, .14, .21, .14, 0, .25, .21, .03, 0, 0, 0, 0, 0))) = \\
 &= \Delta\left(\frac{\sum_{j=0}^{14} j\gamma_j}{\Sigma_{j=0}^{14}}\right) = \Delta(5.8) = (s_6, -0.2) \\
 r_{13} &= \Delta(\chi((0, 0, 0, 0, 0, .07, .16, .43, .54, .62, .47, .31, .11, .02, 0))) = \\
 &= \Delta\left(\frac{\sum_{j=0}^{14} j\gamma_j}{\Sigma_{j=0}^{14}}\right) = \Delta(8.8) = (s_9, -0.2) \\
 r_{14} &= \Delta(\chi((.05, .16, .17, .28, .38, .44, .49, .60, .49, .44, .38, .28, .17, .11, .05))) = \\
 &= \Delta\left(\frac{\sum_{j=0}^{14} j\gamma_j}{\Sigma_{j=0}^{14}}\right) = \Delta(6.9) = (s_7, -0.1) \\
 r_{15} &= \Delta(\chi((0, 0, 0, 0, .14, .10, 0, 0, 0, .25, .25, .17, .07, 0, 0))) = \\
 &= \Delta\left(\frac{\sum_{j=0}^{14} j\gamma_j}{\Sigma_{j=0}^{14}}\right) = \Delta(8.7) = (s_9, -0.3) \\
 r_{16} &= \Delta(\chi((.05, .16, .17, .25, .30, .30, .30, .40, .41, .46, .52, .48, .34, .22, .10))) = \\
 &= \Delta\left(\frac{\sum_{j=0}^{14} j\gamma_j}{\Sigma_{j=0}^{14}}\right) = \Delta(7.8) = (s_8, -0.2) \\
 r_{17} &= \Delta(\chi((0, 0, .12, .27, .47, .64, .64, .64, .34, .19, .04, 0, 0, 0, 0))) = \\
 &= \Delta\left(\frac{\sum_{j=0}^{14} j\gamma_j}{\Sigma_{j=0}^{14}}\right) = \Delta(6) = (s_6, 0) \\
 r_{18} &= \Delta(\chi((.15, .48, .53, .51, .46, .4, .32, .32, .17, .09, .02, 0, 0, 0, 0))) = \\
 &= \Delta\left(\frac{\sum_{j=0}^{14} j\gamma_j}{\Sigma_{j=0}^{14}}\right) = \Delta(3.9) = (s_4, -.1) \\
 r_{19} &= \Delta(\chi((0, 0, 0, 0, 0, 0, 0, 0, .25, .5, .58, .69, .55, .41, .3))) = \\
 &= \Delta\left(\frac{\sum_{j=0}^{14} j\gamma_j}{\Sigma_{j=0}^{14}}\right) = \Delta(11) = (s_{11}, 0)
 \end{aligned}$$

at_1	at_2	at_3	at_4	at_5
$(s_8, 0)$	$(s_6, -.2)$	$(s_9, -.2)$	$(s_7, -.1)$	$(s_9, -.3)$
at_6	at_7	at_8	at_9	
$(s_8, -.2)$	$(s_6, 0)$	$(s_4, -.1)$	$(s_{11}, 0)$	

Cuadro 3.10: Vector de utilidad colectivo expresado por medio de una 2-tupla lingüística.

obteniendo la tabla 3.10 que mostrará los valores colectivos para cada atributo expresado mediante 2-tuplas lingüísticas.

3. Proceso de explotación:

En la fase de explotación queremos obtener una valoración global de todos los atributos que nos proporcionará un grado de adecuación [101], de la alternativa estudiada. Este grado de adecuación lo compararemos con los valores de una tabla de empírica *ad hoc* 3.11. Esta tabla de recomendación ha sido obtenida de forma experimental y expresa la actuación recomendada a partir del grado de adecuación obtenido en esta fase. Nos recomendará la conveniencia de la instalación de un sistema ERP en la empresa. Utilizaremos como operador la media aritmética para 2-tuplas [68] para obtener el grado de adecuación para la instalación del ERP, suponiendo que todos los atributos son igualmente importantes en la instalación del sistema ERP:

$$AM^*((s_0, \alpha_0), \dots, (s_g, \alpha_g)) = \Delta\left(\sum_{i=0}^g \frac{1}{g+1} \Delta^{-1}(s_i, \alpha_i)\right) = \Delta\left(\frac{1}{g+1} \sum_{i=0}^g \beta_i\right) = (s_7, -0.7)$$

Si consultamos la tabla 3.11, para ese grado de adecuación la instalación del ERP es **factible** y por lo tanto podemos inferir que no es totalmente conveniente.

Grado de adecuación	Recomendación
$\leq s_4$	No instalar
$> s_4 \text{ y } \leq s_6$	La instalación no es conveniente
$> s_6 \text{ y } \leq s_9$	La instalación es factible
$> s_9 \text{ y } \leq s_{11}$	La instalación es conveniente
$> s_{11}$	La instalación es muy conveniente

Cuadro 3.11: Ejemplo de tabla de recomendación.

3.3. Conclusiones

En este capítulo se ha presentado el modelo de decisión para problemas de toma de decisión en contextos heterogéneos. El cual ha mostrado la necesidad, utilidad y viabilidad de las funciones definidas en el capítulo 2 de esta memoria para trabajar con problemas definidos en contextos con información no homogénea.

Capítulo 4

Especificación UML para el Tratamiento de Información Heterogénea y el Modelo de Decisión Heterogéneo

Una vez presentadas las herramientas necesarias para el tratamiento de información heterogénea en el capítulo 2 de esta memoria, así como el modelo de decisión para problemas de toma de decisión en contextos heterogéneos en el capítulo 3, hemos considerado conveniente hacer un diseño de dichas herramientas y modelo de decisión para que puedan utilizarse de forma directa, sencilla y reusable en cualquier implementación de sistemas soporte a la decisión o sistemas de evaluación donde su uso pueda flexibilizar el modelado de la información del problema a resolver.

Para alcanzar dicho objetivo aplicaremos los principios de diseño de la ingenie-

ría del software [20] que nos permite crear un modelo para un sistema de soporte a la decisión o un sistema de evaluación que utilicen información heterogénea.

En este capítulo no se pretende presentar de forma exhaustiva los diferentes pasos presentes en la ingeniería del software para construir sistemas software, sino más bien en hacer un diseño completo, correcto, claro y verificable de las clases principales, para tratar con información heterogénea, por medio de diagramas UML [13, 79] (Lenguaje Unificado de Modelado).

El UML es la herramienta de diseño más utilizada actualmente en el paradigma de la programación orientada a objetos [20], el cual nos permitirá la reutilización del software ya implementado de forma simple y directa. UML proporciona distintos tipos de diagramas, como son, diagramas de casos de uso, diagramas de clase, diagramas de secuencia, diagramas de casos de estado, diagramas de actividad, etc. Cada uno de ellos nos ayudará al diseño de nuestro sistema, pero nosotros en esta memoria sólo describiremos brevemente los diagramas de clase ya que el resto son dependiente de cada sistema en particular.

El diseño de los diagramas de clase del sistema soporte a la decisión o sistema de evaluación que utiliza información heterogénea estará orientado para una posterior implementación mediante lenguaje Java [6] y el diseño de clases se dividirá en dos bibliotecas las cuales son:

1. *Biblioteca para el tratamiento de la información heterogénea.* Esta biblioteca incluirá todas las clases necesarias para un correcto tratamiento de la información heterogénea que será utilizada por nuestro modelo de toma de decisión.
2. *Biblioteca del modelo de toma de decisión heterogénea.* Aquí se incluirán las clases necesarias para el diseño del modelo presentado en el capítulo 3 que quedó presentado en la figura 3.1. En esta biblioteca también se incluirán

los operadores de agregación más comunes así como algunos procesos de explotación habituales que suelen utilizarse a la hora de resolver distintos problemas de toma de decisión.

Estas bibliotecas serán comunes de uso habitual para el desarrollo de cualquier sistema soporte a la decisión o sistema de evaluación que trate con información heterogénea. De esta forma disponemos de la base necesaria para poder desarrollar productos software que ayuden a resolver problemas de decisión de forma simple, rápida y eficaz. Para cada nuevo problema de decisión tendremos que diseñar e implementar una biblioteca con las clases propias donde se recogerán las particularidades de dicho problema. Al final de este capítulo se presentará un diseño de clases UML para el ejemplo del proceso de decisión sobre la instalación de un sistema ERP en una empresa presentado en la sección 3.2 de esta memoria. Además en el capítulo 5 se presentará el diseño e implementación de un prototipo para la evaluación de la calidad docente en las Universidades Andaluzas [106] que fue parte de un proyecto realizado para la convocatoria de Grupos de Estudio y Análisis Específicos sobre Calidad en las Universidades Andaluzas de la Unidad para la Calidad de las Universidades Andaluzas (UCUA) en los años 2005 y 2006.

La decisión de utilizar Java, en vez de otro lenguaje orientado a objetos, está determinada porque este lenguaje es independiente de la plataforma donde se ejecutarán los programas y sólo es necesario la instalación de una máquina virtual. De esta manera una vez desarrolladas las clases no necesitan ningún cambio para poder utilizarse en distintos sistemas informáticos. Reduciendo de esta forma el esfuerzo necesario para el desarrollo del software.

4.1. Lenguaje Unificado de Modelado, UML

Aunque no pretendemos hacer una revisión en profundidad del UML, nos parece adecuada hacer un breve resumen y presentar aquellos elementos necesarios para la comprensión de los diagramas que se presentarán a lo largo de este capítulo.

UML es una notación que se produjo como resultado de la unificación de la técnica de modelado de objetos [12] e ingeniería del software orientada a objetos [80]. El UML ha sido diseñado para un amplio rango de aplicaciones. Por lo tanto, proporciona construcciones válidas para sistemas y actividades de distinto tipo (sistemas de tiempo real, sistemas distribuidos, etc.).

El desarrollo de sistemas se enfoca en tres modelos diferentes del sistema:

- El **modelo funcional**, representado en UML con diagramas de caso de uso, describe la funcionalidad del sistema desde el punto de vista del usuario.
- El **modelo de objetos**, representado en UML con diagramas de clase, describe la estructura de un sistema desde el punto de vista de objetos, atributos, asociaciones y operaciones.
- El **modelo dinámico**, representado en UML con diagramas de secuencia, diagramas de gráfica de estado y diagramas de actividad, describe el comportamiento interno del sistema. Los diagramas de secuencia describen el comportamiento como una secuencia de mensajes intercambiados entre un *conjunto de objetos*, mientras que los diagramas de gráfica de estado describen el comportamiento desde el punto de vista de estados de un *objeto individual* y las transacciones posibles entre estados. Por último los diagramas de actividad describe el sistema desde el punto de vista de las actividades que realiza.

Dado que el objetivo principal de este capítulo es presentar el diseño de clases que serán agrupadas en bibliotecas, para posteriormente ser utilizadas en la implementación de un sistema soporte apoyo a la decisión o un sistema de evaluación, nos centraremos sólo en el **modelo de objetos** para el cual UML, utiliza como herramienta los *diagramas de clase*.

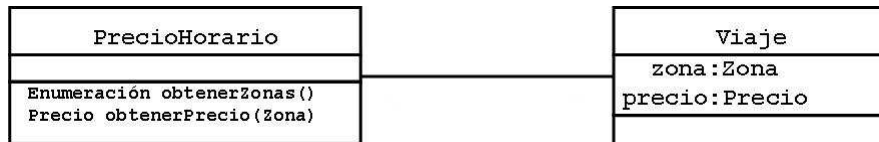


Figura 4.1: Ejemplo de diagrama de clase UML

Diagrama de clase

Se utilizan para describir la estructura del sistema, un ejemplo de esta tipo de diagrama lo podemos ver en la figura 4.1. Las clases son abstracciones que especifican la estructura y el comportamiento común de un conjunto de objetos. Los objetos son instancias de las clases que se crean, modifican y destruyen durante la ejecución del sistema. Los objetos tienen estados que incluyen valores de sus atributos y sus relaciones con otros objetos.

Los diagramas de clases describen el sistema desde el punto de vista de clases, atributos, operadores y sus relaciones:

- *Las clases*, están representadas en el diagrama por medio de un rectángulo con tres divisiones. En la superior es donde se coloca el nombre que le corresponde a la clase.
- *Los atributos*, representan las características propias de los objetos de una clase. Cuando se crea un objeto, éstos atributos tomarán unos valores espe-

cíficos que definirán las características propias de un objetos que lo diferenciarán del resto de objetos de la misma clase. En el diagrama se encuentran en la segunda división, justo debajo del nombre de la clase.

- *Los operadores*, permitirán alterar los valores de los atributos de un objeto o comunicarse con otros objetos presentes en el sistema. Su posición corresponde a la última división.
- *Las relaciones*, representan el modo en que las distintas clases se relacionan entre sí dentro del sistema. Los dos tipos de relaciones principales son:
 - *De asociación*, cuando una clase forma parte de un atributo de otra clase. En el diagrama se representa por medio de una línea. La línea puede terminar en una flecha simple que indica la clase que es el atributo de la otra.
 - *De herencia*, cuando una clase es una especificación de otra más general. En el diagrama se representa por medio de una línea terminada en una flecha con la punta hueca. La flecha con la punta hueca apunta a la clase desde la que se hereda.

4.2. Biblioteca para el tratamiento de la información heterogénea

En primer lugar presentaremos el diseño de clases mediante un diagrama UML para la biblioteca que nos permitirá manejar información heterogénea en problemas definidos en contextos de este tipo, es decir, el diseño de las clases para el modelado de información heterogénea. El diagrama se presenta en la figura 4.2.

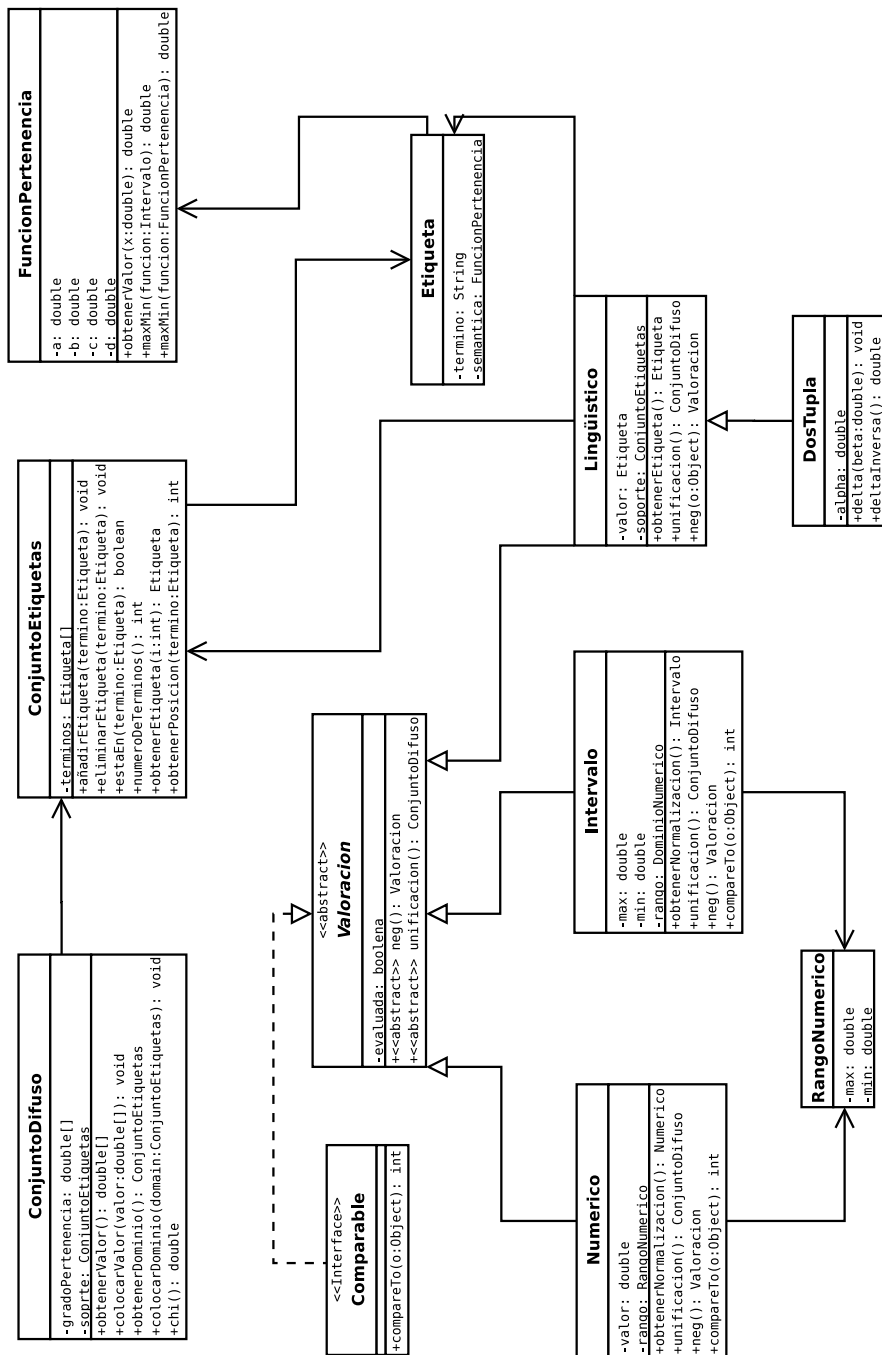


Figura 4.2: Diagrama de clases para el tratamiento de la información heterogénea

En dicho diagrama sólo se han incluido aquellos elementos que consideramos más importantes para la comprensión del diseño. Se han omitido detalles más específicos que serán necesarios para su posterior implementación en Java.

Detallamos a continuación cada una de las clases así como los componentes más importantes de las mismas:

<<abstract>> Valoracion	
-evaluada: booleana	
+<<abstract>> neg(): Valoracion	
+<<abstract>> unificacion(): ConjuntoDifuso	

Valoracion Es una clase abstracta que permite generalizar los elementos necesarios para que los expertos puedan expresar el conocimiento que tienen de las alternativas presentes en un problema de toma de decisión. Hasta el momento de la especificación del problema no se definirá el dominio de expresión en el que será valorada una alternativa. Esta es la razón para declarar la clase como abstracta, de esta forma todos los dominios de expresión tendrán unas propiedades comunes que por medio del polimorfismo¹ permite un tratamiento uniforme de la información (numérica, intervalar, lingüística), a través de sus métodos particulares de cada tipo de información. Además se pueden añadir nuevos dominios de expresión sin que ello altere el desarrollo de los problemas de toma de decisión.

Una propiedad necesaria que tienen que implementar todos los dominios de expresión para poder utilizarlos para modelar información en los problemas que nos atañen es un orden. Esto se consigue por medio de la interfaz **Comparable** y su método *compareTo()* presente en la biblioteca de utilidades de Java. A continuación detallamos los elementos más importantes de esta

¹Característica de programación orientada a objetos presente en el lenguaje Java [6]

clase:

ATRIBUTOS

- *evaluada*. Variable de instancia que nos permitirá saber cuando una alternativa ha sido evaluada por un experto.

MÉTODOS

- *neg()*. Método abstracto que tendrán que implementar cada uno de los dominios de expresión (numérico, intervalar y lingüístico) y que devolverá una **Valoracion** que será la negación de ella misma.
- *unificacion()*. Método abstracto que tendrán que implementar cada uno de los dominios de expresión y que devolverá un **ConjuntoDifuso** expresado en el conjunto básico de términos lingüísticos (CBTL). Para cada una de las clases, que representa a los dominios de expresión, se implementará este método para que calcule la función correspondiente presentada anteriormente en el capítulo 2 (definiciones 2.5, 2.6 y 2.8) de esta memoria.

Numerico
-valor: double
-rango: RangoNumerico
+obtenerNormalizacion(): Numerico
+unificacion(): ConjuntoDifuso
+neg(): Valoracion
+compareTo(o:Object): int

Numerico Clase que nos permitirá trabajar con el dominio de expresión numérico.

ATRIBUTOS

- *valor*. Variable de instancia donde se almacenará el valor que un experto asigna a la valoración de una alternativa numérica.

- *rango*. Variable de instancia que permite al experto una mayor flexibilidad a la hora de suministrar su valoración. Se permitirá definir el rango en el que podrá ser valorada una alternativa numérica que no tiene por que estar comprendido en el intervalo [0,1].

MÉTODOS

- *obtenerNormalizacion()*. Dado que no siempre se valorará una alternativa numérica en el intervalo [0,1] es necesario definir un método que nos permita transformar esa información al intervalo [0,1].

Intervalo
-max: double
-min: double
-rango: DominioNumerico
+obtenerNormalizacion(): Intervalo
+unificacion(): ConjuntoDifuso
+neg(): Valoracion
+compareTo(o:Object): int

Intervalo Clase que nos permitirá trabajar con el dominio de expresión intervalar.

ATRIBUTOS

- *max*. Variable de instancia que almacenará el valor superior para una valoración intervalar de una alternativa.
- *min*. Variable de instancia que almacenará el valor inferior para una valoración intervalar de una alternativa.
- *rango*. Variable de instancia que permite al experto una mayor flexibilidad a la hora de suministrar su valoración. Se permitirá definir el rango en el que podrá ser valorado el intervalo de una alternativa que no tiene por que estar comprendido en el intervalo [0,1].

MÉTODOS

- *obtenerNormalizacion()*. Dado que no siempre se valorará una alternativa intervalar en el intervalo $[0,1]$ es necesario definir un método que nos permita transformar esa información al intervalo $[0,1]$.

RangoNumerico
-max: double
-min: double

RangoNumerico Clase que nos permite asociar un rango de valoraciones diferentes al $[0,1]$ para alternativas valoradas en los dominios de expresión numérico y/o intervalar. De esta forma el experto tendrá una mayor flexibilidad para expresar sus valoraciones y la escala en la que tendrá que hacer su valoración estará más cercana al tipo de conocimiento que posee sobre una alternativa.

Lingüístico
-valor: Etiqueta
-soporte: ConjuntoEtiquetas
+obtenerEtiqueta(): Etiqueta
+unificacion(): ConjuntoDifuso
+neg(o:Object): Valoracion

Lingüístico Clase que nos permitirá trabajar con el dominio de expresión lingüístico ya que a partir de ella se instanciarán variables lingüísticas cuyos valores serán etiquetas lingüísticas.

ATRIBUTOS

- *valor*. Variable de instancia que representa la etiqueta con la que ha sido valorada una alternativa expresada en el dominio lingüístico por parte del experto.

- *soporte*. Variable de instancia que representa el conjunto de etiquetas lingüísticas disponibles por parte del experto para realizar su valoración para una alternativa expresada por medio de un dominio lingüístico. De esta forma se ofrece la posibilidad de que no todas las variables lingüísticas estén valoradas en un único conjunto de etiquetas lingüísticas tal y como propone el modelo presentado en el capítulo 3 de esta memoria.

DosTupla
-alpha: double
+delta(beta:double): void
+deltaInversa(): double

DosTupla Clase que implementa las características de una 2-tupla lingüística presentada en [68].

Etiqueta
-termino: String
-semantica: FuncionPertenencia

Etiqueta Clase que representa las características presentes en una etiqueta lingüística:

ATRIBUTOS

- *termino*. Variable de instancia que nos permite representar la sintaxis de una etiqueta lingüística.
- *semantica*. Variable de instancia que define la semántica para un término lingüístico por medio de la clase **FuncionPertenencia**.

FuncionPertenencia
-a: double -b: double -c: double -d: double
+obtenerValor(x:double): double +maxMin(funcion:Intervalo): double +maxMin(funcion:FuncionPertenencia): double

FuncionPertenencia Clase que define la semántica para una etiqueta lingüística y que estará parametrizada por medio de cuatro valores permitiendo definir funciones de pertenencia trapezoidales y triangulares [10].

ConjuntoEtiquetas
-terminos: Etiqueta[]
+añadirEtiqueta(termino:Etiqueta): void +eliminarEtiqueta(termino:Etiqueta): void +estaEn(termino:Etiqueta): boolean +numeroDeTerminos(): int +obtenerEtiqueta(i:int): Etiqueta +obtenerPosicion(termino:Etiqueta): int

ConjuntoEtiquetas Clase que define conjuntos de etiquetas lingüísticas. Un conjunto de etiquetas será una colección de la clase **Etiqueta** que representarán cada una de las etiquetas lingüísticas (sintaxis y semántica) presentes en el conjunto y los métodos necesarios para la inserción, borrado, eliminación y consulta propios de una colección.

ConjuntoDifuso
-gradoPertenencia: double[] -soporte: ConjuntoEtiquetas
+obtenerValor(): double[] +colocarValor(valor:double[]): void +obtenerDominio(): ConjuntoEtiquetas +colocarDominio(domain:ConjuntoEtiquetas): void +chi(): double

ConjuntoDifuso Clase que permite representar un conjunto difuso valorado en un conjunto de etiquetas lingüísticas. Esta clase será necesaria para la unificación de la información de entrada, previa a cualquier proceso de agregación de información heterogénea.

ATRIBUTOS

- *gradoPertenencia*. Variable de instancia que almacena, por medio de una colección del tipo de dato *double* de Java, los grados de pertenencia de las etiquetas lingüísticas del CBTL para el objeto a unificar (numérico, intervalar o lingüístico) mediante un conjunto difuso [73].
- *soporte*. Conjunto de etiquetas lingüísticas donde se unifica la información por medio de conjuntos difusos. Generalmente este conjunto de etiquetas será el CBTL, dado que esta clase se utilizará para la unificación de la información de entrada de un problema de toma de decisión heterogénea.

MÉTODOS

- *chi()*. Método que nos permitirá transformar un conjunto difuso valorado en el CBTL a su representación en una 2-tupla lingüística. Implementa la función, χ , presentada anteriormente en el capítulo 2 de esta memoria.

4.3. Biblioteca del modelo de toma de decisión heterogénea

Una vez que hemos diseñado las clases para implementar las herramientas necesarias para manejar información heterogénea, presentadas en el capítulo 2, pasaremos a presentar, tal como indicamos al inicio de este capítulo, el diseño del diagrama de clases para implementar el modelo de decisión con información heterogénea expuesto en el capítulo 3 de esta memoria. Este diagrama de clase UML se muestra gráficamente en la figura 4.3.

Al igual que en el diagrama anterior sólo se muestran aquellos elementos más significativos del diseño de las clases. El diagrama incluye únicamente aquellas clases mínimas necesarias para el diseño de esta biblioteca. Por tanto no se incluyen clases para operadores de agregación ni procesos de explotación genéricos presentes en los problemas de decisión dado que no son necesarias para comprender el diseño y sólo serán necesarias una vez que la biblioteca esté implementada en Java. A continuación pasamos a exponer detalladamente cada una de ellas:



Experto Clase que nos permite modelar los elementos necesarios para mantener la información relativa a un experto. Dispondrá de una única variable de instancia de la clase **String** que nos permitirá identificar a los expertos presentes en un problema de decisión.



ConjuntoAlternativas Clase que nos permitirá modelar las alternativas presentes en un problema de toma de decisión. Dispondrá de una variable de instancia

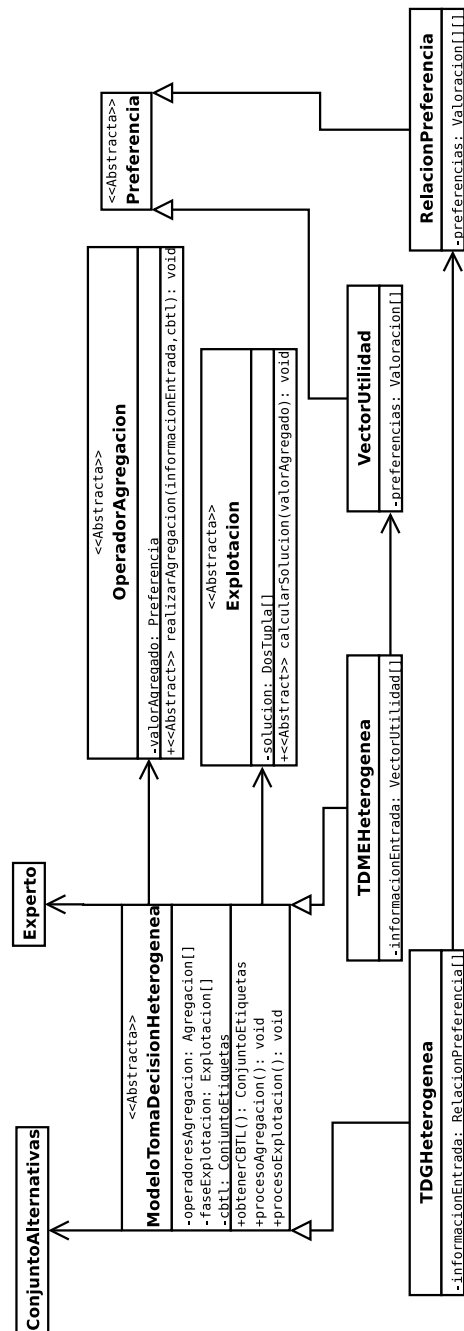
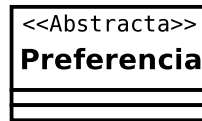
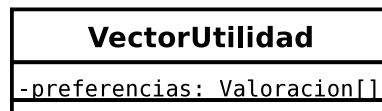


Figura 4.3: Diagrama de clases del modelo de decisión heterogéneo

que será un conjunto de la clase **String** donde cada elemento del conjunto representará a una alternativa.



Preferencia Clase abstracta. Esta clase tendrá utilidad sólo para el diseño de la biblioteca. De ella deberán heredar todas las estructuras de representación de información para los problemas de decisión presentes en la biblioteca. Obtenemos una estructura general para representar los distintas estructuras de representación de información que posteriormente serán utilizados por los operadores de agregación y procesos de explotación presentes en esta biblioteca.



VectorUtilidad Clase que nos permite implementar la estructura de representación de información como son los vectores de utilidad. Donde cada experto proporciona un valor de preferencia para cada alternativa del problema. Esta estructura de representación de información es común en los problemas de toma de decisión multiexperto (TDME).

ATRIBUTOS

- *preferencias*. Variable de instancia que nos permite almacenar las preferencias suministradas por un experto para el conjunto de alternativas presente en el problema de decisión. Las preferencias se almacenan como una colección de la clase **Valoracion** descrita en la biblioteca anterior.

RelacionPreferencia
-preferencias: Valoracion[][]

RelacionPreferencia Clase que nos permite implementar la estructura de representación de información como son las relaciones de preferencia. Donde cada experto proporciona un valor de preferencia de una alternativa sobre las demás presentes en el problema. Esta estructura de representación de información es utilizada habitualmente en los problemas de toma de decisión en grupo (TDG).

ATRIBUTOS

- *preferencias*. Variable de instancia que nos permite almacenar las preferencias suministradas por un experto para el conjunto de alternativas presente en el problema de decisión. Las preferencias se almacenan como una colección de la clase **Valoracion** descrita en la biblioteca anterior.

<<Abstracta>>
ModeloTomaDecisionHeterogenea
-operadoresAgregacion: Agregacion[] -faseExplotacion: Explotacion[] -cbtl: ConjuntoEtiquetas
+obtenerCBTL(): ConjuntoEtiquetas +procesoAgregacion(): void +procesoExplotacion(): void

ModeloTomaDecisionHeterogenea Es una clase abstracta que nos define los elementos comunes presentes en todos los problemas de decisión que estén definidos en un contexto heterogéneo y sigue el modelo de decisión cuya estructura puede verse en la figura 3.1. De esta forma sólo se tienen que diseñar e implementar los elementos particulares según el problema que se

esté tratando de resolver. Los elementos más importantes del diseño son los siguientes:

ATRIBUTOS

- *operadoresAgregacion*. Variable de instancia donde se guardan los distintos operadores de agregación que podremos utilizar dentro de nuestro problema de decisión. De esta forma podremos utilizar distintos operadores para una misma información de entrada del problema, para poder comparar resultados sin tener que volver a pedir a los expertos que nos vuelvan a suministrar sus preferencias sobre las alternativas presentes en el problema de decisión.
- *faseExplotacion*. Variable de instancia donde se almacenan distintos procesos de explotación que nos permiten obtener distintos resultados para un mismo problema sin tener que repetir todo el proceso de decisión. De esta forma para un mismo problema se pueden aplicar tantos procesos de explotación como se tengan ya implementados dentro de la biblioteca o los que se definan específicamente para un problema dado.
- *cbtl*. Variable de instancia donde se almacenará el conjunto básico de etiquetas lingüísticas para realizar el proceso de unificación de la información de entrada presente en las diferentes clases sobre las que los expertos pueden dar sus valoraciones (numérico, intervalar y lingüístico).

MÉTODOS

- *obtenerCBTL()*. Método de instancia que nos implementa el proceso de obtención del CBTL descrito en la sección 2.2.1 del capítulo 2 de esta memoria.

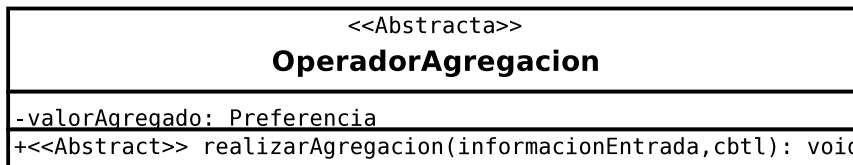
- *procesoAgregacion()*. Método que nos permite aplicar los distintos operadores de agregación presentes en la variable de instancia *informacionAgregada*. Como paso previo se tendrá que calcular el CBTL invocando el método *obtenerCBTL()*. Posteriormente se invocarán los distintos métodos *realizarAgregacion()* presentes en la clase **Agregación** que representan cada uno de los operadores de agregación presentes en el problema de decisión.
- *procesoExplotacion()*. Método que nos permite aplicar los distintos procesos de explotación presentes en la variable de instancia *valorFinal*. Se invocará el método *calcularSolucion()* presente en la clase **Explotacion** que representa a cada uno de los procesos de explotación presentes en el problema de decisión.

TDMEHeterogenea
-informacionEntrada: VectorUtilidad[]

TDMEHeterogenea Clase que hereda de la clase **ModeloTomaDecisionHeterogenea** que nos permite modelar los problemas de TDME con información heterogénea. Dependiendo del problema de toma de decisión la estructura de representación de la información no está definida y habrá que añadirla a la clase. Por tanto esta clase añade, con respecto a **ModeloTomaDecisionHeterogenea**, una variable de instancia para poder almacenar la información de entrada al problema. La estructura de representación de la información para este tipo de problemas son los vectores de utilidad.



TDGHeterogenea Clase que hereda de la clase **ModeloTomaDecisionHeterogenea** que nos permite modelar los problemas de TDG con información heterogénea. Al igual que en la clase anterior se añade una variable de instancia para almacenar la información de entrada. La estructura de representación de la información para este tipo de problemas son las relaciones de preferencia.



OperadorAgregacion Clase abstracta que nos permitirá definir las características comunes para cada uno de los operadores de agregación que queramos definir dentro de esta biblioteca o para un problema de decisión específico. Las características principales son las siguientes:

ATRIBUTOS

- *valorAgregado*. Variable de instancia donde se almacenará el valor obtenido de la aplicación del operador de agregación a la información de entrada unificada.

MÉTODOS

- *realizarAgregacion()*. Método abstracto que deberá ser implementado para cada operador de agregación que esté presente en la biblioteca o para el operador de agregación específico para un problema de toma de decisión. Una vez ejecutado este método la variable de instancia

valorAgregado contendrá la información relativa a la aplicación de ese operador para la información de entrada del problema de decisión.

<<Abstracta>> Explotacion	
-solucion: DosTupla[]	
+<<Abstract>> calcularSolucion(valorAgregado): void	

Explotacion Clase abstracta que nos permitirá definir las características comunes para cada uno de los procesos de explotación que queramos definir dentro de esta biblioteca o para un problema de decisión específico. Las características principales son las siguientes:

ATRIBUTOS

- *solucion*. Variable de instancia donde se almacenarán los valores del conjunto solución una vez aplicado el proceso de explotación. El valor solución vendrá determinado por una colección de la clase **DosTupla**. Como ya se ha justificado anteriormente en el capítulo 3 de esta memoria, las soluciones para nuestros problemas de decisión vendrán expresadas por medio de 2-tuplas lingüísticas.

MÉTODOS

- *calcularSolucion()*. Método abstracto que deberá ser implementado para cada proceso de explotación que esté presente en la biblioteca o para el proceso de explotación específico para un problema de toma de decisión. Una vez ejecutado este método la variable de instancia *solucion* contendrá el valor o valores solución del problema de decisión.

4.4. Biblioteca para la instalación de un sistema ERP

Una vez introducidos los diagramas generales para manejar información heterogénea en problemas de decisión, en esta sección presentaremos el diseño UML de clases para un problema específico. El diseño de clases que se presenta en la figura 4.4 pretende adecuarse de una forma fiel al ejemplo presentado en el capítulo 3 de esta memoria, problema de decisión para la instalación de un sistema ERP en una empresa. El objetivo de este diseño es presentar la metodología que deberemos seguir en el diseño e implementación de problemas de decisión específicos y cómo se integran en el diseño e implementación las bibliotecas anteriormente definidas. En el próximo capítulo de esta memoria presentaremos el diseño e implementación de un prototipo para la evaluación de la calidad docente en las Universidades Andaluzas [106].

A continuación pasamos a describir los elementos principales de las clases presentes en el diagrama UML:

OperadorAgregacionMedia
+realizarAgregacion(informacionEntrada,cbtl): void

OperadorAgregacionMedia Clase que implementa el operador de agregación media aritmética. Esta clase se encuentra implementada dentro de la biblioteca *modeloTomaDecision*, pero se presenta en este momento porque resulta más claro a la hora de entender el diseño de este problema de decisión. Este será el único operador que se aplicará para la implementación presentada de este problema de toma de decisión, aunque podría haberse dejado sin redefinir la variable de instancia *informaciónAgregada* y de esta forma se podrían aplicar distintos operadores de agregación presentes en la biblioteca **modeloTomaDecisionHeterogenea**. Este diseño responde al ejemplo presentado

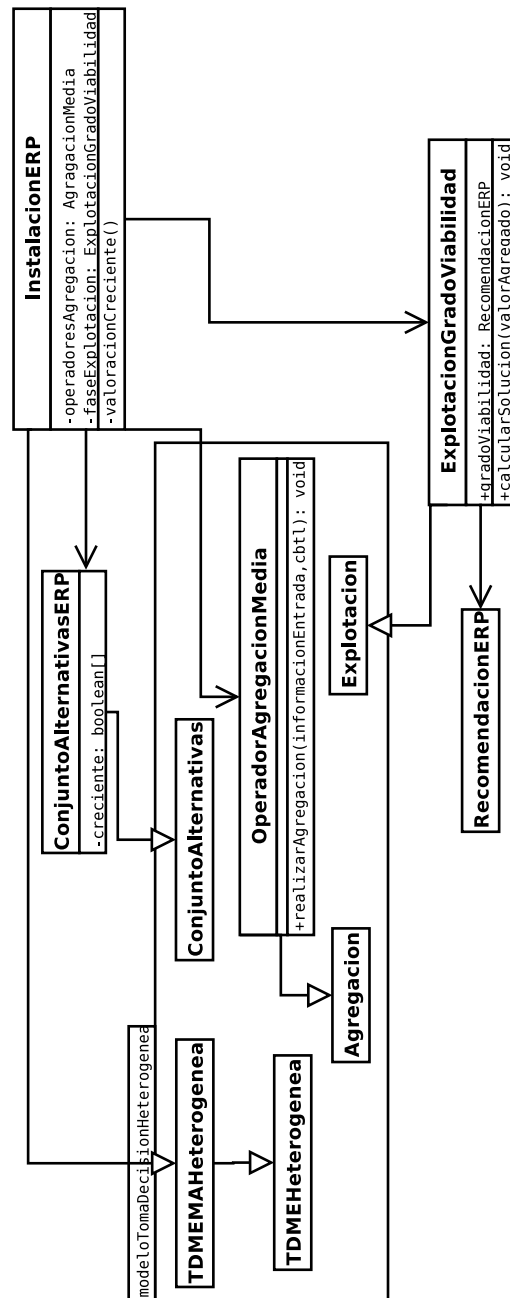


Figura 4.4: Diagrama de clases para la instalación de un sistema ERP

en el capítulo 3.

MÉTODOS

- *realizarAgregacion()*. Como ya se ha explicado anteriormente se ha de implementar este método para cada una de las clases que implementen operadores de agregación. En nuestro caso una media aritmética de todos los valores presentes para una alternativa valorados por cada uno de los expertos. Como resultado obtendremos un nuevo vector de utilidad que representa estos valores medios para cada una de las alternativas presentes en el problema.

RecomendacionERP

RecomendacionERP Clase que nos permitirá implementar una tabla de recomendación para la instalación de un sistema ERP en una empresa como la presentada en el capítulo 3 en la tabla 3.11.

ExplotacionGradoViabilidad
+gradoViabilidad: RecomendacionERP
+calcularSolucion(valorAgregado): void

ExplotacionGradoViabilidad Clase que implementa el proceso de explotación presentado en el capítulo 3 para el proceso de decisión sobre la instalación de un sistema ERP. Este es un proceso de explotación propio de implementación para la instalación de un ERP.

ATRIBUTOS

- *gradoViabilidad*. Variable de instancia que vendrá representada por una

instancia de la clase **RecomendacionERP**. De esta forma se podrá interpretar la solución obtenida en el proceso de evaluación.

MÉTODOS

- *calcularSolucion()*. Como el problema de decisión es un TDME la información de entrada agregada viene expresada por medio de un vector de utilidad. El proceso de explotación es la implementación del presentado en el ejemplo del capítulo 3 y como solución se obtiene un grado de viabilidad expresado por medio de una 2-tupla lingüística y que será interpretado por medio del objeto almacenado en la variable de instancia *gradoViabilidad*.

ConjuntoAlternativasERP
-creciente: boolean[]

ConjuntoAlternativasERP Clase que hereda de **ConjuntoAlternativas**. Como se presentó en el ejemplo del capítulo 3 de esta memoria, no todas las alternativas de un problema de decisión para la instalación de un sistema ERP en una empresa tienen una interpretación creciente. Por tanto es necesario identificar qué alternativas tienen una interpretación creciente y cuáles no. Para ello es necesario añadir una nueva variable de instancia, *creciente* para cada una de las alternativas donde poder almacenar esa información.

TDMEMAHeterogenea

TMDEMAHeterogenea Clase que hereda de **TDMEHeterogenea** donde cada una de las alternativas del problema de decisión es valorada mediante un

conjunto de de atributos. Todos los expertos proporcionan un vector de utilidad para cada alternativa que contendrá un valor de los atributos de la misma.

InstalacionERP
-operadoresAgregacion: AgragacionMedia
-faseExplotacion: ExplotacionGradoViabilidad
-valoracionCreciente()

InstalacionERP Clase que hereda de **TDMEMAHeterogenea** y por tanto sólo deberán definirse, o redefinirse, aquellos elementos específicos para el problema de instalación de un sistema ERP.

ATRIBUTOS

- *operadoresAgregacion*. Esta variable de instancia se ha redefinido puesto que sólo se utilizará un único operador de agregación para este problema. De esta forma se representa de forma fiel el ejemplo presentado en el capítulo 3 de esta memoria. Pero si no se redefine esta variable se podrían utilizar distintos operadores de agregación definidos dentro de la biblioteca **modeloTomaDecisionHeterogenea** u otros definidos específicamente para este problema de decisión.
- *faseExplotacion*. Se ha redefinido esta variable de instancia puesto que el proceso de explotación de este problema de decisión se ha definido expresamente para el mismo y ya ha sido descrito en el capítulo 3 de esta memoria.

MÉTODOS

- *valoracionCreciente()*. Método privado que deberá ser invocado como primera acción en el método *procesoAgregacion* de esta clase para que todas las alternativas del problema de decisión tengan una interpreta-

ción no creciente de sus valores como ya se explicó en el ejemplo del capítulo 3 de esta memoria.

4.4.1. Conclusiones

En este capítulo hemos presentado un diseño correcto, claro, completo y verificable, siguiendo los pasos descritos en la ingeniería del software [20], para el diseño de las distintas bibliotecas que nos permitirán implementar sistemas de soporte a la decisión o sistemas de evaluación con un esfuerzo mínimo. Además hemos presentado el diseño de la biblioteca necesario para la implementación de un problema de decisión que permite decidir si se instalará un sistema ERP en una empresa. Con este diseño hemos pretendido ilustrar los pasos a seguir para cualquier problema de decisión definido en un contexto heterogéneo.

En el próximo capítulo presentaremos un ejemplo sobre un problema de evaluación donde además detallaremos la implementación de un prototipo para el diseño propuesto.

Capítulo 5

Evaluación de la Calidad Docente en las Universidades Andaluzas

En este capítulo se presenta el diseño de un prototipo para un *Sistema de Evaluación de la Calidad Docente en las Universidades Andaluzas* que pretende medir la calidad de los profesores, los servicios que tiene la Universidad para docencia y aspectos generales del alumnado [106]. Actualmente las encuestas, que se utilizan para la evaluación de la calidad docente universitaria, están compuestas por distintas preguntas que son rellenadas por los alumnos que cursan una asignatura en la Universidad. Cada alumno rellenará una encuesta por profesor que le imparta la asignatura, tanto en teoría como en prácticas. Las respuestas de todas las preguntas se evalúan en una escala numérica discreta de 1 a 5 (ver figura 5.1). A pesar de que muchas de las preguntas están relacionadas con percepciones propias del alumno que son difíciles de cuantificar de forma precisa ya que, implican incertidumbre en su conocimiento. Además los indicadores evaluados en estas encuestas no tienen siempre la misma naturaleza por lo que, unos pueden ser evaluados fácilmente mediante información precisa (indicadores cuantitativos) mientras que otros

ENCUESTA DE OPINIÓN DEL ALUMNADO SOBRE LA ACTUACIÓN DOCENTE DEL PROFESORADO

Profesor/a

Si estás **completamente de acuerdo** con que las afirmaciones describen el comportamiento del(la) profesor(a), marca la casilla "5". Si estás **completamente en desacuerdo**, marca la casilla "1". Utiliza los "puntos intermedios de la escala (2,3,4) para graduar tu respuesta". Si el enunciado **no procede** o no tienes información para responder, marca la opción "NS-NC". Procura responder todas las preguntas.

marca así

así no marcas

1. El profesor informa del programa de la asignatura cuando comienza a impartirla	1	2	3	4	5	NS-NC
2. Informa de los objetivos del programa de la asignatura						
3. El programa contiene información bibliográfica útil para el desarrollo de la asignatura	1	2	3	4	5	NS-NC
4. El profesor comienza las clases a la hora fijada en el horario						
5. Imparte clases los días establecidos	1	2	3	4	5	NS-NC
6. Cuando falta a clase da las razones de su ausencia						
7. Cuando falta a clase la recupera otro día	1	2	3	4	5	NS-NC
8. Cuando asistes a sus tutorías en el horario establecido, te atiende						
9. Explica los contenidos de la asignatura con seguridad	1	2	3	4	5	NS-NC
10. Inicia cada tema exponiendo los objetivos del mismo						
11. Explica con claridad, facilitando, en su caso, la toma de notas	1	2	3	4	5	NS-NC
12. Destaca los aspectos fundamentales de cada tema						
13. Pregunta durante el desarrollo de las clases para averiguar si los alumnos tienen dificultades	1	2	3	4	5	NS-NC
14. Motiva a los alumnos para que se interesen por la asignatura						
15. Expone ejemplos o situaciones en las que se utilizan los contenidos de la asignatura	1	2	3	4	5	NS-NC
16. En general, hace interesantes las clases						
17. Propone actividades para favorecer el aprendizaje autónomo (búsqueda de información complementaria, trabajos, investigaciones, etc.)	1	2	3	4	5	NS-NC
18. Favorece que los alumnos desarrollen una actividad reflexiva						
19. Utiliza recursos didácticos (transparencias, pizarra, medios audiovisuales, informáticos, etc.) que ayudan a comprender los contenidos	1	2	3	4	5	NS-NC
20. Utiliza una metodología de enseñanza adecuada a las características del grupo y de la asignatura						
21. Utiliza un lenguaje claro e inteligible	1	2	3	4	5	NS-NC
22. Informa del sistema de evaluación al principio del curso						
23. Utiliza diferentes procedimientos para evaluar el aprendizaje de los alumnos	1	2	3	4	5	NS-NC
24. Toma en consideración las propuestas de los alumnos sobre el desarrollo de la asignatura						
25. Tiene un trato igualitario con todos los alumnos	1	2	3	4	5	NS-NC
26. Es respetuoso en el trato con los alumnos						
27. Responde con interés a las intervenciones de los alumnos	1	2	3	4	5	NS-NC
28. Valora globalmente al profesor en esta asignatura						

Valoración de Servicios e Infraestructuras

29. Valora las condiciones de espacio, material, equipamientos, etc. en que se desarrolla la asignatura	1	2	3	4	5	NS-NC
30. Valora tu grado de satisfacción con el servicio de secretaría de la Escuela/Facultad						
31. Valora tu grado de satisfacción con el servicio de reprografía	1	2	3	4	5	NS-NC
32. Valora tu grado de satisfacción con el servicio de cafetería/comedor						
33. Valora tu grado de satisfacción con el servicio de Atención al Estudiante	1	2	3	4	5	NS-NC
34. Valora tu participación en actividades culturales						
35. Valora tu participación en actividades deportivas de la Universidad	1	2	3	4	5	NS-NC
36. Valora tu satisfacción con las actividades culturales y/o deportivas en las que participas						

Datos Personales

Para poder interpretar los resultados, necesitamos que respondas a las siguientes preguntas:

37. Sexo: ☐ Hombre ☐ Mujer

38. ¿En qué convocatoria te encuentras en esta asignatura?

1ª 2ª 3ª 4ª 5ª 6ª

☐ ☐ ☐ ☐ ☐ ☐

39. ¿Cuál es tu nivel de asistencia con este profesor a esta asignatura?

☐ Menos del 20% ☐ Entre el 20% y el 40% ☐ Entre el 40% y el 60%
☐ Entre el 60% y el 80% ☐ Más del 80%

Teléfono (voluntario) a efectos de validación del cuestionario

0	0	0	0	0	0	0	0
1	1	1	1	1	1	1	1
2	2	2	2	2	2	2	2
3	3	3	3	3	3	3	3
4	4	4	4	4	4	4	4
5	5	5	5	5	5	5	5
6	6	6	6	6	6	6	6
7	7	7	7	7	7	7	7
8	8	8	8	8	8	8	8
9	9	9	9	9	9	9	9

Figura 5.1: Encuesta para la evaluación de la calidad docente universitaria

se adaptarán mejor a una evaluación cualitativa. Por lo que los objetivos principales de nuestra propuesta son:

1. Introducir metodologías de tratamiento de la incertidumbre en los procesos de evaluación de calidad docente, tales como, la lógica difusa y el modelado de preferencias heterogéneo [61] que nos permitan modelar las percepciones de los alumnos de una forma más adecuada que la actual. Y así ofrecer a los alumnos un marco de expresión más flexible que el utilizado hasta el momento. Nuestro objetivo es permitir que los alumnos puedan expresar su evaluación mediante información *numérica*, *intervalar* y *lingüística* (contexto heterogéneo), dependiendo de la naturaleza del indicador a evaluar.
2. Automatizar el proceso evaluativo que es inherente al uso de un sistema informático, aunque esta automatización puede llevarse a cabo a distintos niveles en cada actividad de la evaluación.

Para alcanzar los objetivos propuestos utilizaremos las herramientas presentadas en el capítulo 2 y nuestro modelo de evaluación se basará en el modelo de decisión presentado en el capítulo 3. Para las herramientas y el modelo de decisión ya se ha presentado un diseño de bibliotecas en el capítulo 4 de esta memoria. En este capítulo presentaremos el diseño de los elementos necesarios para alcanzar el objetivo de automatización del proceso de evaluación. Para ello presentaremos la arquitectura de nuestro sistema de evaluación y el diseño del sistema de evaluación así como un prototipo que hemos desarrollado del mismo.

5.1. Arquitectura del Sistema de Evaluación

Como primer paso se presentará la arquitectura del *Sistema de Evaluación de la Calidad Docente* la cual se presenta en la figura 5.2. En esta arquitectura se

recogen los elementos fundamentales presentes en el sistema de evaluación.

Nuestro sistema de evaluación de calidad docente será utilizado por dos tipos de usuarios bien diferenciados con diferentes atribuciones cada uno:

- **El administrador.** Este usuario es el encargado de diseñar las encuestas, es decir, será el encargado de definir las preguntas así como los dominios de información en las que serán valoradas (numérico, lingüístico o intervalar). Además también será el encargado de gestionar la información presente en la base de datos. En ella se almacenarán los diseños de las encuestas, datos sobre profesores, asignaturas, alumnos y además todas las encuestas completadas por los alumnos. En la base de datos también se almacenará la información extraída de las encuestas como parte del proceso de evaluación.
- **Los alumnos.** Serán los encargados de completar las encuestas que se encuentran a su disposición en el sistema de evaluación. Antes que el alumno pueda completar una encuesta deberá registrarse en el sistema. De esta forma se puede presentar la encuesta apropiada a un alumno dependiendo de las encuestas que estén presentes en el sistema y que no haya completado ya.

Una vez presentados los usuarios que utilizarán el sistema de evaluación pasaremos a describir brevemente los módulos principales que componen la arquitectura del mismo:

- **Interfaz de Alumno.** Cuando un alumno se identifica en el sistema sólo tendrá acceso a este módulo del sistema donde se le presentarán las encuestas que puede rellenar. Este módulo tendrá las siguientes funcionalidades:
 - *Información sobre la encuesta.* Al alumno se le presentará toda la información que sea precisa para ayudarlo a comprender qué se le está pidiendo que responda en cada momento.

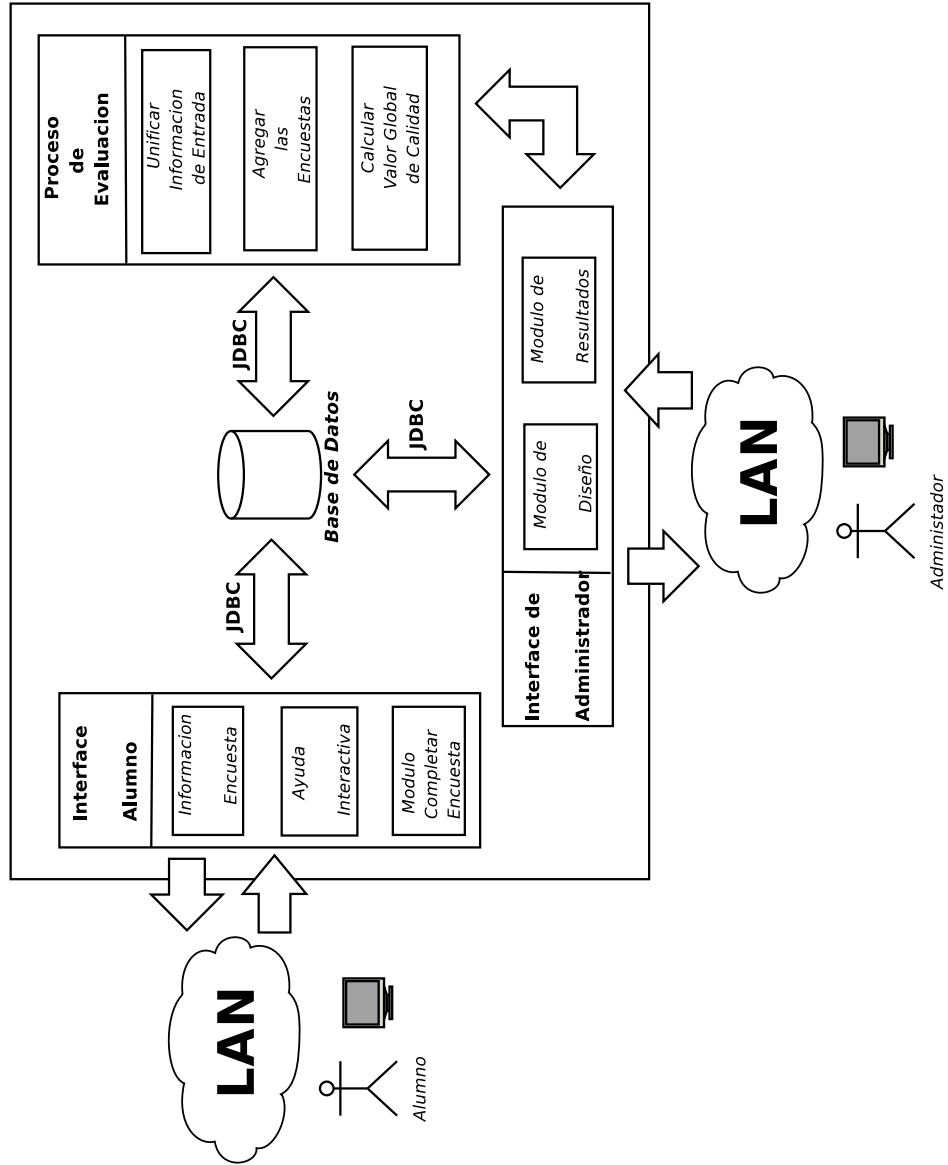


Figura 5.2: Arquitectura Sistema de Evaluación de la Calidad Docente

- *Ayuda interactiva.* El alumno tendrá ayuda a su disposición sobre la interfaz que está utilizando. Esta ayuda podrá ser solicitada por el alumno en cualquier momento.
 - *Módulo para completar la encuesta.* La principal funcionalidad de esta interfaz es permitir al alumno poder evaluar las preguntas presentes en la encuesta valorándolas cada una de ellas en su dominio y una vez finalizada enviarla a la base de datos para su almacenamiento (anónimo) y posterior análisis.
- **Interfaz de Administrador.** El administrador del sistema tendrá a su disposición una interfaz que le permitirá tener acceso a los siguientes módulos dentro del sistema:
- *Módulo de diseño.* Con este módulo el administrador podrá establecer las preguntas de una encuesta así como los dominios de expresión donde serán valorados por los alumnos. Una vez completado el diseño de una encuesta se almacenará en la base de datos y desde ese momento estará disponible para que todos los alumnos puedan evaluarla.
 - *Módulo de resultados.* Una vez que todas las encuestas sobre una asignatura y el profesor que la imparte se encuentran almacenadas, o el administrador considera que ha habido tiempo suficiente para su evaluación, este módulo recupera la información de la base de datos para poder calcular el índice de calidad. Para ello se utilizará el módulo que implementa el proceso de evaluación. También se almacenarán los cálculos realizados en la base de datos, de esta forma se calculan los resultados una vez y se pueden consultar muchas veces sin necesidad de volver a calcular los resultados. Recordar que una vez iniciado el

proceso de cálculo para una encuesta ésta deja de estar disponible en el sistema para los alumnos.

■ **Proceso de Evaluación.** Este módulo es llamado por el *Módulo de resultados* del administrador. En él se implementan las acciones necesarias para realizar el proceso de evaluación de la calidad docente. Este proceso de evaluación se implementa como un proceso de toma de decisiones en un contexto heterogéneo (ver figura 2.2) con los siguientes pasos:

- *Unificar la información de entrada.* Como las distintas preguntas de la encuesta pueden ser valoradas por los alumnos en distintos dominios de información, según el diseño que haya seguido el administrador, deben ser unificadas todas ellas en un único dominio de expresión para poder aplicar el paso siguiente.
- *Agregar las encuestas.* Una vez que todas las preguntas están unificadas en un único dominio de expresión se aplica un operador de agregación para obtener un valor colectivo para cada una de las preguntas de la encuesta.
- *Calcular el valor global de calidad.* Por último se aplica un proceso de explotación a los valores colectivos de la encuesta para obtener un valor global de la calidad docente.

5.2. Biblioteca de clases para el Proceso de Evaluación Docente

Una vez que se ha descrito la arquitectura del sistema de evaluación pasaremos a presentar el diseño de clases para el *Proceso de Evaluación Docente*. El siste-

ma de evaluación sigue la estructura del modelo de TD definido en un contexto heterogéneo que podemos ver en la figura 2.2. A continuación detallaremos las diferencias que presenta este sistema de evaluación respecto al modelo de decisión anteriormente presentado:

- *Adquisición de la información.* En este paso del modelo sólo hay que hacer notar que son los alumnos los que toman el papel de los expertos a la hora de suministrar sus preferencias sobre las distintas alternativas a evaluar dentro del sistema.
- *Proceso de Agregación.* Este paso del modelo queda inalterado dentro del sistema de evaluación que estamos definiendo.
- *Proceso de Evaluación.* Este paso del modelo sufre una alteración. Nuestro sistema de evaluación una vez ha conseguido el valor colectivo para cada una de las alternativas que se están evaluando no pretende obtener la mejor alternativa entre todas. En este caso lo que se pretende es calcular un valor global que nos indique el grado de la evaluación que se está realizando.

Utilizando las bibliotecas descritas en el capítulo anterior de esta memoria, y teniendo en cuenta las particularidades de nuestro sistema de evaluación, pasamos a presentar el diseño propuesto para la biblioteca que posteriormente implementará el módulo **Proceso Evaluacion** de la arquitectura presentada anteriormente en este capítulo. El diseño se presenta por medio de un diagrama UML de clase y se muestra en la figura 5.3.

A continuación pasamos a describir los elementos principales de las clases presentes en el diagrama UML sin entrar en los detalles específicos que deberán ser tenidos en cuenta para realizar la implementación. De esta forma se presenta un ejemplo más de cómo se pueden utilizar las bibliotecas definidas en el capítulo an-

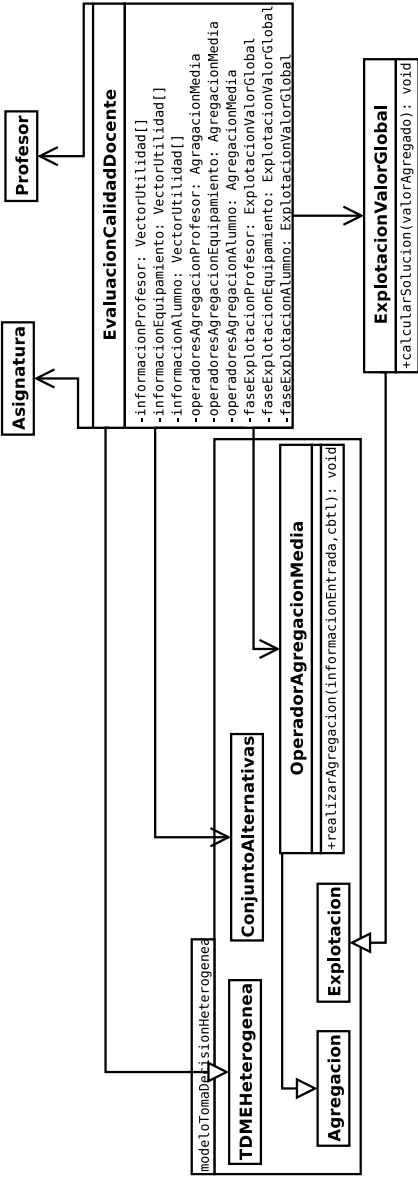


Figura 5.3: Diagrama de clases para el Proceso de Evaluación Docente

terior de la memoria reduciendo el esfuerzo necesario para el desarrollo del sistema de evaluación de la calidad docente. Las clases son las siguientes:

OperadorAgregacionMedia
<code>+realizarAgregacion(informacionEntrada,cbtl): void</code>

OperadorAgregacionMedia Clase que implementa el operador de agregación media aritmética. Esta clase se encuentra implementada dentro de la biblioteca *modeloTomaDecision*, ya ha sido presentada en el capítulo anterior pero volvemos a detallarla para una mejor comprensión del sistema de evaluación. En esta etapa del desarrollo del sistema de evaluación consideramos que todas las alternativas presentes en la evaluación tienen la misma importancia.

MÉTODOS

- *realizarAgregacion()*. Como ya se ha explicado anteriormente se ha de implementar este método para cada una de las clases que implementen operadores de agregación. En nuestro caso una media aritmética de todos los valores presentes para una alternativa valorados por cada uno de los expertos. Como resultado obtendremos un nuevo vector de utilidad que representa estos valores medios para cada una de las alternativas presentes en el problema.

Asignatura

Asignatura Clase que nos permite almacenar la información relativa a la asignatura sobre la que se está realizando la evaluación.



Profesor Clase que nos permite almacenar la información relativa al profesor que se está realizando la evaluación.



ExplotacionValorGlobal Clase que hereda de **Explotacion** y que implementa el proceso de explotación para el sistema de evaluación. Como ya hemos comentado antes, este proceso no persigue obtener la mejor alternativa y en su lugar obtendrá un valor global para la evaluación que se está realizando por parte del sistema.

MÉTODOS

- *calcularSolucion()*. Este método debe ser implementado para cada clase que hereda de **Explotacion**. En el caso del sistema de evaluación de la calidad docente se le suministra la valoración global de las alternativas a evaluar y el método almacenará en la variable de instancia *solucion* una 2-tupla lingüística que representa el valor global de la evaluación.

EvaluacionCalidadDocente
-informacionProfesor: VectorUtilidad[] -informacionEquipamiento: VectorUtilidad[] -informacionAlumno: VectorUtilidad[] -operadoresAgregacionProfesor: AgregacionMedia -operadoresAgregacionEquipamiento: AgregacionMedia -operadoresAgregacionAlumno: AgregacionMedia -faseExplotacionProfesor: ExplotacionValorGlobal -faseExplotacionEquipamiento: ExplotacionValorGlobal -faseExplotacionAlumno: ExplotacionValorGlobal

EvaluacionCalidadDocente Clase que hereda de **TDMEHeterogenea** dado que los alumnos suministrarán sus preferencias sobre las distintas alternativas a evaluar por medio de vectores de utilidad. Esta clase deberá definir aquellos elementos necesarios para poder llevar a cabo el proceso de evaluación.

ATRIBUTOS

- *informacionProfesor, informacionEquipamiento, informacionAlumno.* Como la actual encuesta que se le suministra a los alumnos para realizar el proceso de evaluación, ver figura 5.1, se encuentra dividida en tres secciones diferenciadas debemos almacenar las preguntas para cada una de las partes de la encuesta en una variable distinta. De esta forma podremos obtener la evaluación para cada una de ellas dado que son independientes entre sí.
- *operadoresAgregacionProfesor, operadoresAgregacionEquipamiento, operadoresAgregacionAlumno.* Dado que la información de entrada se encuentra dividida en tres partes diferenciadas, realizaremos el proceso de agregación para cada una de ellas de forma independiente. Como en esta fase del diseño consideramos que todas las alternativas presentes en la evaluación tienen el mismo grado de importancia sólo utilizaremos un único operador de agregación para cada una de ellas y será la

media aritmética. Este operador está implementado por la clase **OperadorAgregacionMedia** definido en biblioteca *modeloTomaDecisionHeterogeneo* presentada en el capítulo anterior de esta memoria. En una posible revisión del diseño podremos considerar la necesidad de añadir más operadores de agregación o utilizar distintos operadores de agregación para cada una de las partes del sistema de evaluación (profesorado, equipamiento y alumnado)

- *faseExplotacionProfesor, faseExplotacionEquipamiento, faseExplotacionAlumno*. Dado que se ha calculado un valor colectivo para cada una de las tres partes presentes en la información de entrada, también debemos obtener un valor global de evaluación para cada una de ellas. De esta forma necesitaremos tres variables de instancia que almacenarán cada una su propia instancia de la clase **ExplotacionValorGlobal** para poder calcular y almacenar su propio valor global de evaluación.

5.3. Prototipo del Sistema de Evaluación Docente

Una vez mostrado el diseño UML pasamos a presentar el prototipo para el *Sistema de Evaluación Docente*. El prototipo ha sido implementado en JAVA por tratarse de un lenguaje orientado a objetos, independiente de la plataforma. Esta característica del lenguaje JAVA nos permite ejecutar nuestro prototipo virtualmente en cualquier sistema informático. De esta forma no se necesita adaptar el prototipo para distintas plataformas sólo hay que disponer de una máquina virtual JAVA que es la encargada de ejecutar el prototipo y además esta máquina virtual es gratuita y se encuentra disponible para los distintos sistemas informáticos más habituales del mercado.

El prototipo es una primera implementación de la arquitectura presentada en

la figura 5.2. Vamos a presentar en detalle cada uno de los elementos principales del prototipo desde el punto de vista de los dos tipos de usuarios presentes en el sistema:

- El Administrador
- Los Alumnos

Como el objetivo principal de este capítulo es presentar el funcionamiento del prototipo para un proceso de evaluación basado en nuestro modelo de decisión en contexto heterogéneo no presentaremos en esta memoria el diseño de la base de datos necesaria para almacenar la información generada por nuestro sistema de evaluación.

5.3.1. Interfaz del Administrador

Primero pasaremos a describir la interfaz del administrador. Esta interfaz estará compuesta por los módulos de diseño y de resultados. En el prototipo actual ambos módulos se encuentran separados en interfaces independientes. Una vez que se avance más en el desarrollo del sistema de evaluación se estudiará el aspecto final de la interfaz para acceder a los distintos módulos presentes en el sistema.

El primer módulo que presentaremos será el del diseño de una encuesta de evaluación dado que será necesario tener alguna encuesta presente en el sistema para que los alumnos puedan dar su evaluación en las distintas preguntas presentes en la misma.

5.3.1.1. Módulo de Diseño

Este módulo permitirá crear una encuesta al personal experto para la realización de las preguntas relativas a la evaluación de la calidad docente de los profe-

sores de las Universidades Andaluzas. La ventana de la interfaz de este módulo puede verse en la figura 5.4.

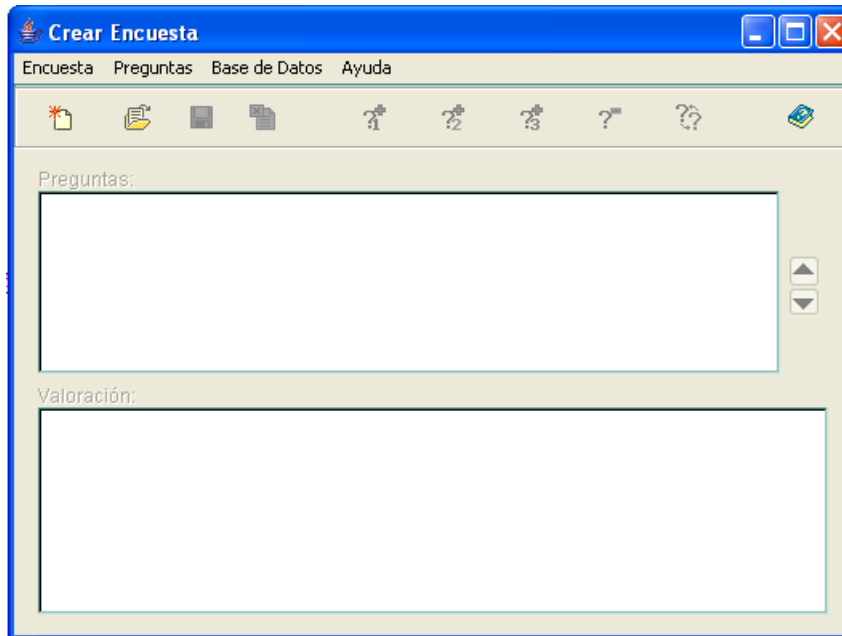


Figura 5.4: Interfaz para la creación de una encuesta

Una vez que tenemos lanzado el módulo para crear una encuesta pulsamos el botón  o si ya tenemos almacenada una en la base de datos y queremos completarla o modificarla pulsamos en el botón . Si lo que hemos hecho es pulsar en el botón  aparecerá una ventana parecida a la mostrada en la figura 5.5.

En el caso de haber cargado una ya existente se nos mostraría la ventana que podemos ver en la figura 5.6, donde podemos seleccionar cualquier encuesta que esté almacenada en la base de datos y la ventana del módulo mostrará algo parecido a la figura 5.7. Una vez que hemos empezado a crear una encuesta o a modificar una ya existente podremos añadir preguntas o modificarlas. Hay que hacer notar

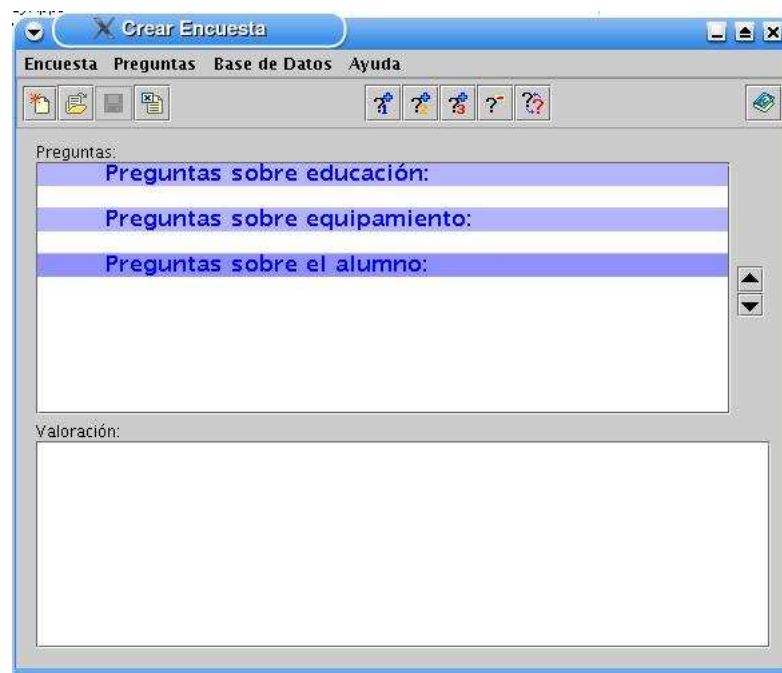


Figura 5.5: Creación de una encuesta



Figura 5.6: Seleccionar encuesta

que la encuesta está dividida en 3 partes, una para preguntas relacionadas con el profesorado, otra para preguntas relacionadas con el equipamiento a disposición del profesorado puesto por la Universidad y una última para preguntas relacionadas con el alumnado. De esta forma se mantiene la estructura presentada por la figura 5.1 dado que es el formato actual de las encuestas de evaluación del profesorado de las Universidades Andaluzas.

Para insertar las preguntas en alguna de estas categorías tenemos tres botones:

Crear Encuesta - ejemplo

Encuesta Preguntas Base de Datos Ayuda

Preguntas:

Preguntas sobre educación:

- 1: El profesor informa del programa de la asignatura cuando comienza a impartirla: Valoración lingüística sin valorar
- 2: Informa de los objetivos del programa de la asignatura: Valoración lingüística sin valorar
- 3: Cuando asistes a sus tutorías en el horario establecido, te atiende: Valoración enumerada
- 4: Motiva a los alumnos para que se interesen en la asignatura: Valoración lingüística sin valorar

Preguntas sobre equipamiento:

- 1: Utiliza recursos didácticos (transparencias, pizarra, medios audiovisuales, informática): Valoración lingüística sin valorar
- 2: El programa contiene información bibliográfica útil para el desarrollo de la asignatura: Valoración lingüística sin valorar

Preguntas sobre el alumno:

- 1: Valora globalmente al profesor de esta asignatura: Valoración lingüística sin valorar
- 2: En general hace interesantes las clases: Valoración lingüística sin valorar
- 3: Utiliza una metodología de enseñanza adecuada a las características del grupo y de la asignatura: Valoración lingüística sin valorar

Valoración:

Nada Poco Normal Bastante Mucho

Figura 5.7: Encuesta ya creada con anterioridad

para añadir preguntas relacionadas con el profesorado,
 para añadir preguntas relacionadas con el equipamiento,
 para añadir preguntas relacionadas con el alumnado. También tenemos la posibilidad de borrar una pregunta o de modificarla .

Añadir preguntas

Si lo que deseamos es añadir una pregunta pulsaremos sobre el botón correspondiente al grupo de preguntas en donde deseamos añadir algo y aparecerá la pantalla que se muestra en la figura 5.8. Como el problema que estamos implementando permite a los alumnos expresar su conocimiento en distintos dominios de información, éste es el punto donde se definirá el dominio apropiado para cada alternativa que posteriormente será evaluada por los alumnos.

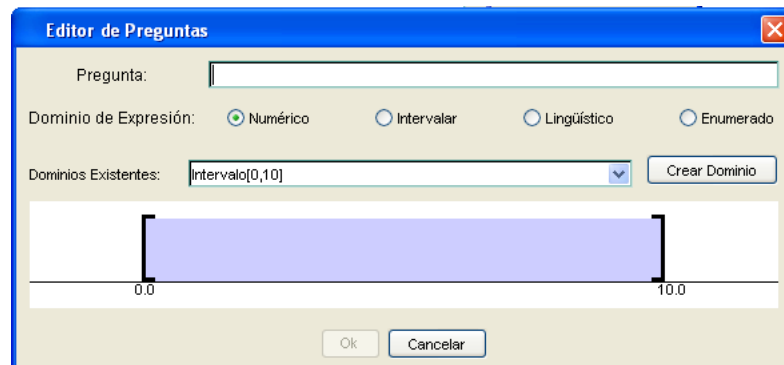
The image shows a software window titled "Editor de Preguntas". It contains a text field labeled "Pregunta:". Below it are four radio buttons for "Dominio de Expresión": "Numérico" (selected), "Intervalar", "Lingüístico", and "Enumerado". Underneath is a dropdown menu labeled "Dominios Existentes" with the value "Intervalo[0,10]". To the right of the dropdown is a button labeled "Crear Dominio". Below these elements is a horizontal axis with tick marks at "0.0" and "10.0". A light blue shaded rectangular area is positioned between these two points, representing the interval. At the bottom of the window are two buttons: "Ok" and "Cancelar".

Figura 5.8: Añadir preguntas a la encuesta

Introduciremos la pregunta, el dominio de información donde será evaluada (numérico, intervalar, lingüístico o enumerado); o podremos crear uno nuevo si los existentes no se ajustan a nuestras necesidades. Todo ello se muestra en la figura 5.9.

Si no tenemos el dominio de información donde una pregunta será evaluada podremos crearlo pulsando el botón (Crear Dominio) y se nos mostrará una ventana como en las figuras 5.10, 5.11, ó 5.12 dependiendo del dominio de expresión que tengamos seleccionado.

Para el caso de un dominio intervalar o numérico la ventana que se nos muestra nos permite introducir los límites superior e inferior para este tipo de valores en la

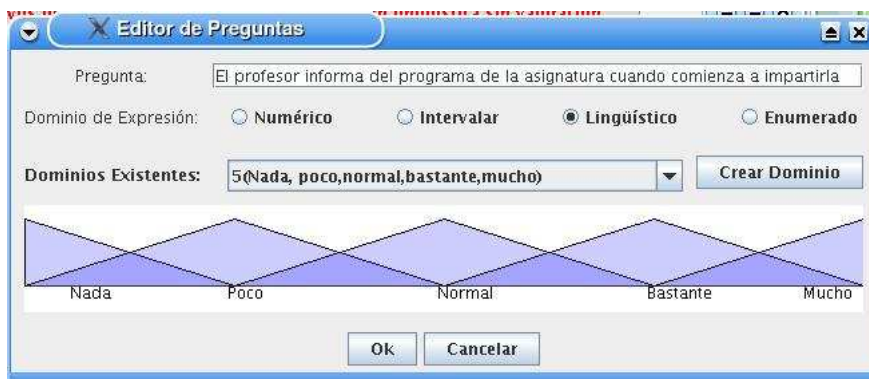


Figura 5.9: Complementando una pregunta de la encuesta

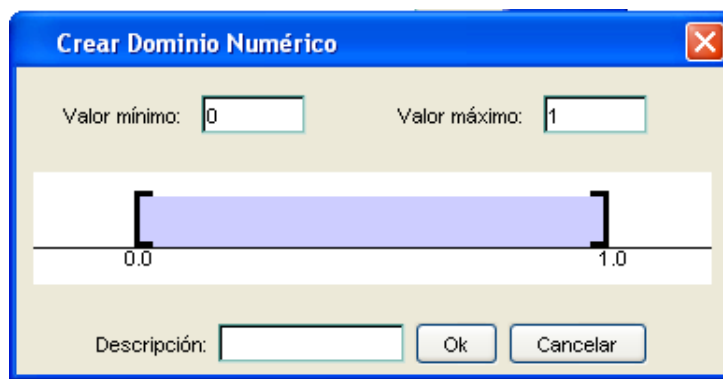


Figura 5.10: Dominio numérico

pregunta correspondiente. Deberemos dar un nombre al dominio creado y éste se almacenará en la base de datos.

Si queremos crear un dominio lingüístico indicamos el número de etiquetas que lo componen, siempre un número impar, y luego pulsaremos en el botón **Crear**. Una vez hecho esto se nos mostrará una representación gráfica para las etiquetas que serán triangulares y uniformemente distribuidas. Podemos dar un nombre a cada etiqueta seleccionándola con el ratón y además podremos dar valores a los distintos parámetros hasta adaptar el conjunto difuso que representa a esa etiqueta

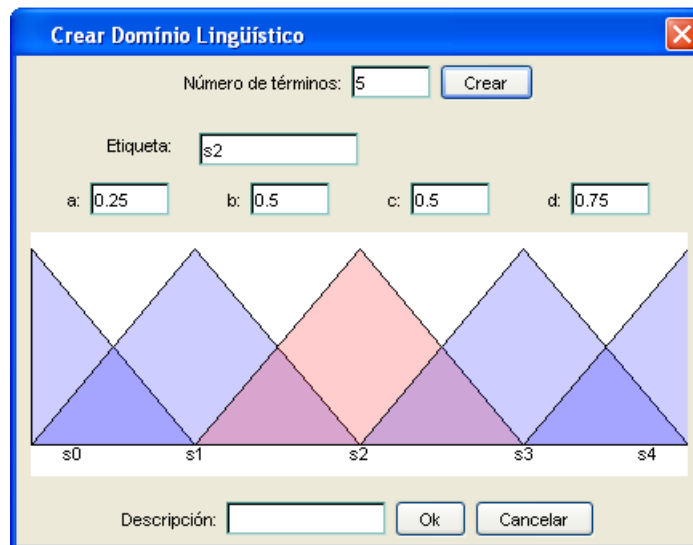


Figura 5.11: Dominio lingüístico

según nuestras necesidades. Para terminar pondremos nombre al dominio para su almacenamiento en el fichero de base de datos.



Figura 5.12: Dominio enumerado

Para valores que se representan sólo por un conjunto enumerado de etiquetas a las cuales no se le asignará una interpretación lingüística, por ejemplo, hombre o mujer; necesitamos crear un dominio donde se recojan estas etiquetas. Basta con

escribir la etiqueta y añadirla al conjunto. Cuando terminemos le daremos nombre al dominio para su almacenamiento en la base de datos.

Borrar preguntas

Si queremos borrar una pregunta la seleccionamos y pulsamos sobre el botón de borrar .

Modificar preguntas

Para modificar una pregunta, la seleccionamos y pulsamos el botón  que nos mostrara una pantalla similar a la figura 5.13.

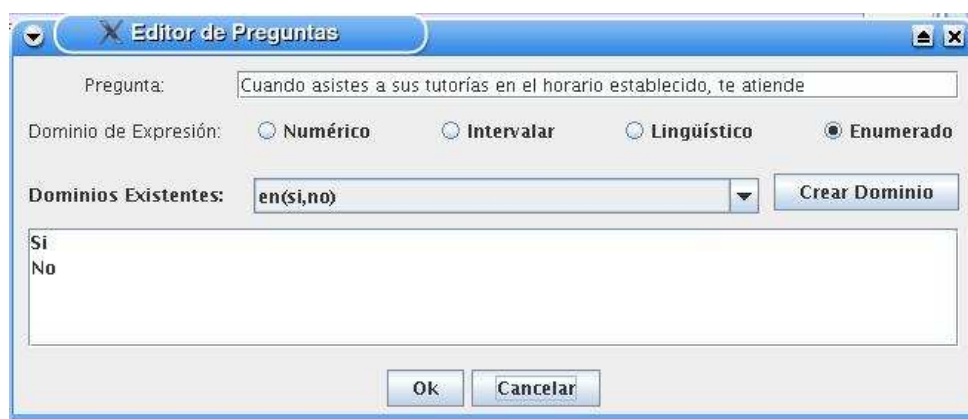


Figura 5.13: Modificar una pregunta

La pantalla es similar a la que se nos muestra cuando estamos creando una pregunta de la encuesta salvo que ya tiene valores para todos sus elementos. Podemos modificar cualquiera de ellos y si pulsamos en el botón **OK** los cambios quedarán registrados.

Una vez que hemos completado toda la encuesta nos quedará una pantalla similar a la que se muestra en la figura 5.14.

Crear Encuesta - ejemplo

Encuesta Preguntas Base de Datos Ayuda

Preguntas:

Preguntas sobre educación:

- 1: El profesor informa del programa de la asignatura cuando comienza a impartirla: Valoración lingüística sin valoración
- 2: Informa de los objetivos del programa de la asignatura: Valoración lingüística sin valoración
- 3: Cuando asistes a sus tutorías en el horario establecido, te atiende: Valoración enumerada
- 4: Motiva a los alumnos para que se interesen en la asignatura: Valoración lingüística sin valoración

Preguntas sobre equipamiento:

- 1: Utiliza recursos didácticos (transparencias, pizarra, medios audiovisuales, informáticos, etc.): Valoración lingüística sin valoración
- 2: El programa contiene información bibliográfica útil para el desarrollo de la asignatura: Valoración lingüística sin valoración

Preguntas sobre el alumno:

- 1: Valora globalmente al profesor de esta asignatura: Valoración lingüística sin valoración
- 2: En general hace interesantes las clases: Valoración lingüística sin valoración
- 3: Utiliza una metodología de enseñanza adecuada a las características del grupo y de la asignatura: Valoración lingüística sin valoración
- 4: Tiene un trato igualitario con todos los alumnos: Valoración lingüística sin valoración

Valoración:

Nada Poco Normal Bastante Mucho

Figura 5.14: Encuesta terminada

Sólo resta guardar la encuesta con un nombre y quedará almacenada en la base de datos para que posteriormente los alumnos puedan rellenarla.

5.3.2. Módulo de resultados

El objetivo de este módulo de la aplicación es el cálculo de los resultados de las encuestas realizadas por los alumnos. Es donde se implementa el modelo de decisión para problemas definidos en contexto heterogéneo presentado en el capítulo 3 de esta memoria, adaptándolo al proceso de evaluación que nos ocupa. Como

se verá a continuación, existen diferentes vistas o formas de presentación de los resultados. Se puede analizar el conjunto de resultados de todas las preguntas o bien centrarse en los resultados parciales de cada pregunta y encuesta. Por lo tanto, la persona o grupos de personas encargadas de analizar los resultados de las encuestas disponen de una herramienta que mostrará los resultados en un dominio de expresión que es fácil de entender e interpretar a la vez, además, nos permitirá un análisis gráfico para conocer de forma individual o colectiva los resultados obtenidos en el proceso de evaluación.

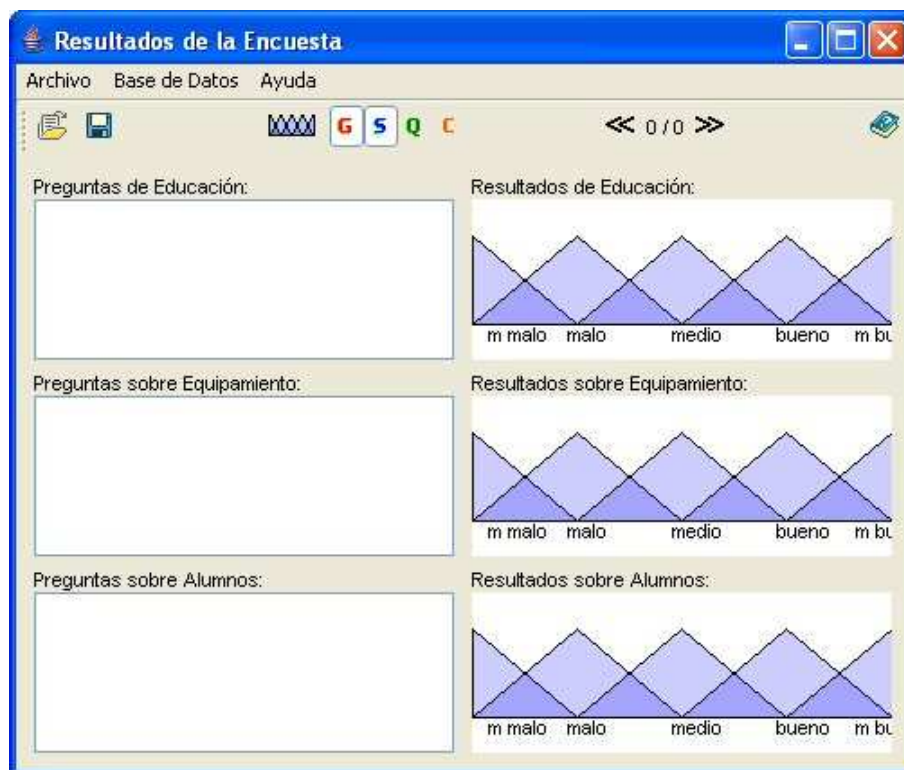


Figura 5.15: Resultados de la encuesta

A continuación describiremos las diferentes opciones y operaciones relacionadas con los resultados de las encuestas. En la figura 5.15 podemos ver la interfaz

inicial que se nos presenta para este módulo.

5.3.2.1. Abriendo una Encuesta

La primera operación a realizar es la apertura del modelo de encuesta que deseemos consultar. Hemos de tener en cuenta que cada modelo de encuesta se identifica a partir de tres campos que son: la asignatura, el profesor y el nombre de la encuesta. Al pulsar sobre el botón , la aplicación visualiza una ventana (figura 5.16) donde el usuario debe seleccionar los valores de los anteriores campos.

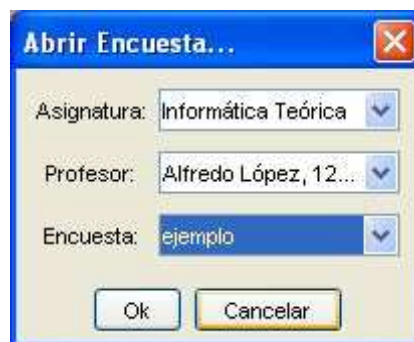


Figura 5.16: Selección de la encuesta a consultar

Si para la combinación de los valores de estos tres campos existen encuestas contestadas por los alumnos, se mostrará la ventana de la figura 5.17 con toda la información referente a esta encuesta.

De esta ventana podemos destacar dos partes:

1. Selección de criterios a visualizar: En la parte superior de la ventana aparecen varios botones que permiten seleccionar el tipo de información que se quiere consultar.
2. Presentación de resultados: El bloque central de la ventana lo componen por un lado las preguntas que forma parte del modelo de encuesta que se está

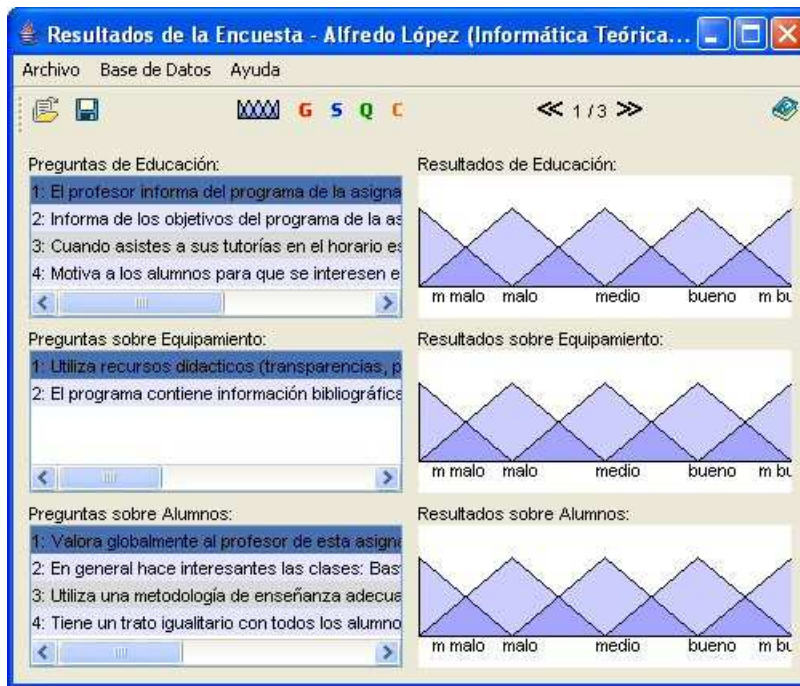
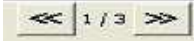


Figura 5.17: Encuesta seleccionada

consultando junto a una representación gráfica en la que aparecerán visibles los diferentes valores obtenidos a partir de los criterios seleccionados.

Los detalles de cada una de estas partes serán comentadas en los siguientes apartados.

5.3.2.2. Encuesta actual/Número de encuestas totales

En la parte superior de la ventana, aparecen los botones  que indican el número de encuesta que actualmente está visualizándose en la ventana así como el número total de encuestas disponibles para la asignatura-profesor que estamos evaluando. Pulsando en los respectivos botones de avance y retroceso el usuario puede consultar los resultados de la encuesta deseada.

5.3.2.3. Selección de criterios

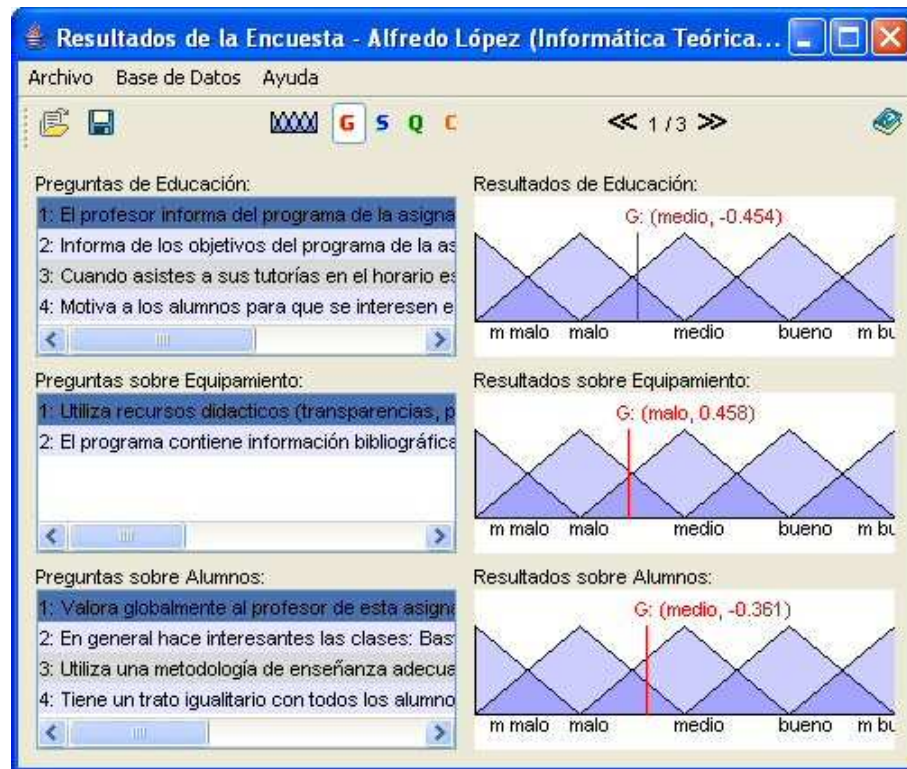


Figura 5.18: Valoración calidad docente

El conjunto de botones  permite visualizar los resultados de las encuestas desde diferentes puntos de vista.

- Valoración global de la encuesta : Representa el valor global de cada uno de los apartados que componen el modelo de encuesta obtenido a partir de todas las encuestas. No depende de la encuesta actual ni de la pregunta seleccionada en ese momento (figura 5.18). Este botón implementa el modelo de evaluación que se basa en nuestro modelo de decisión que ha sido presentado en el capítulo 3 de esta memoria.

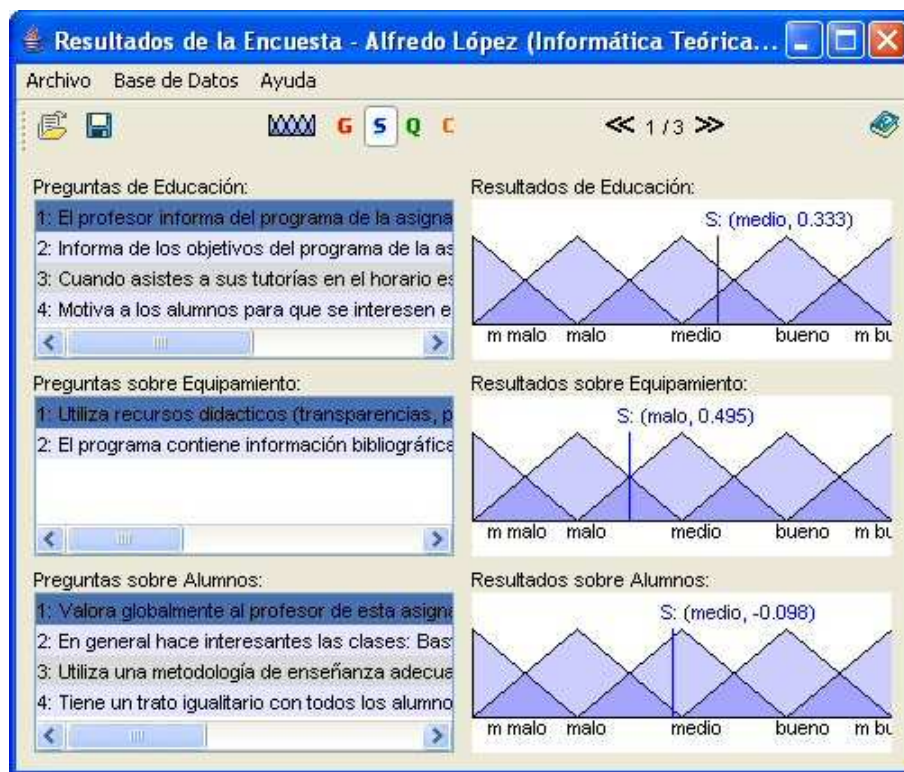


Figura 5.19: Valoración encuesta actual

- S
 Valoración de la encuesta actual: Muestra la medida global pero sólo teniendo en cuenta una encuesta, la que se presenta en ese momento. No depende de la pregunta seleccionada en ese momento. Para consultar el resto de encuestas contestadas, puede usar los botones $\ll$ y $\gg$ para avanzar y retroceder en las encuestas presentes en la base de datos (figura 5.19). Implementa la misma acción que en el caso anterior pero sólo teniendo en cuenta los valores de la encuesta que se está mostrando en ese momento.



Figura 5.20: Valores globales por pregunta

- Valoración global de cada pregunta (agregadas todas las encuestas) : Representa el valor global de cada pregunta que componen los diferentes apartados obtenido a partir de todas las encuestas. No depende de la encuesta actual. Seleccionando cada pregunta se irán visualizando los correspondientes valores globales (figura 5.20). Este botón implementa las herramientas presentadas para la realización del proceso de agregación con información heterogénea presentadas en el capítulo 2 de esta memoria.

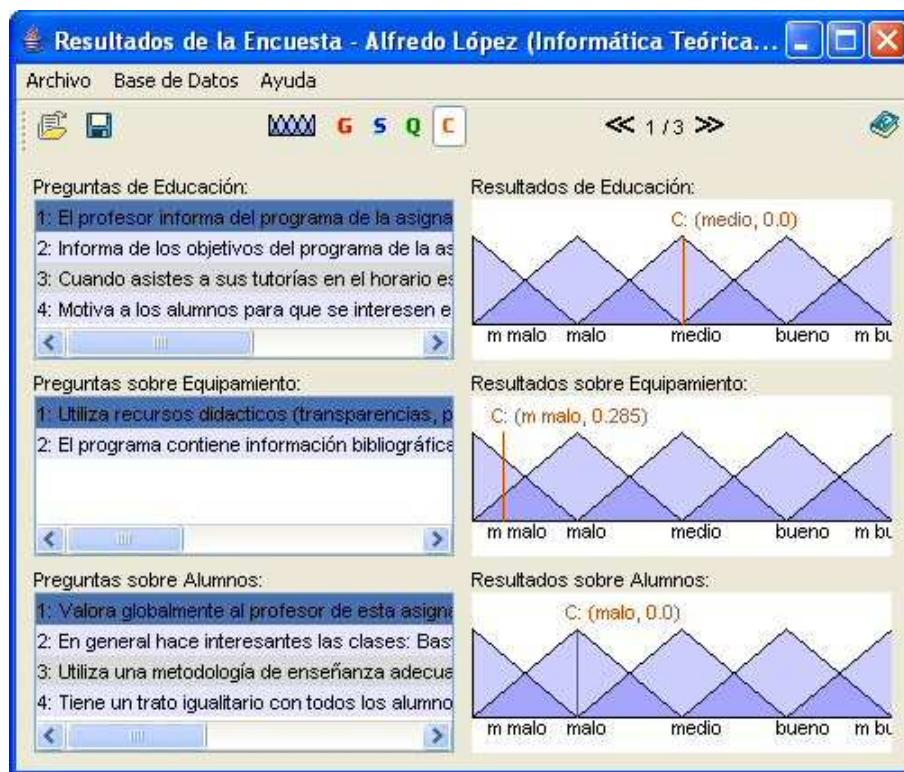


Figura 5.21: Valores en cada pregunta

- Valoración local de una pregunta : Representa el valor de cada una de las preguntas de la encuesta actual. Seleccionando cada una de las preguntas se irán visualizando los correspondientes valores de esa pregunta (figura 5.21).

5.3.3. Interfaz del Alumno

Una vez concluida la presentación de la interfaz del administrador pasaremos a presentar en detalle la interfaz del alumno. Esta interfaz permitirá, a los alumnos presentes en la base de datos, acceder a los módulos correspondientes del sistema de evaluación para poder completar una encuesta. Una vez que el alumno se ha

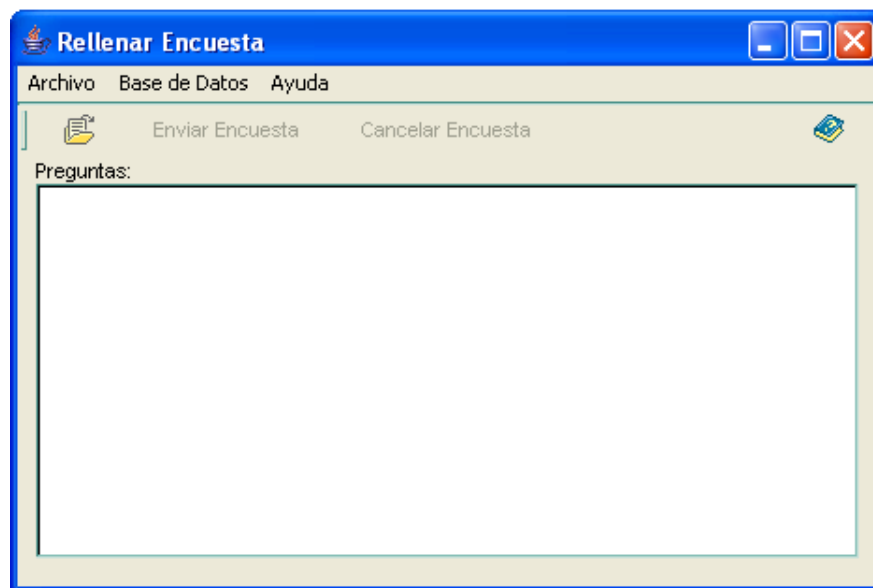


Figura 5.22: Entrada de opiniones

identificado en el sistema de evaluación deberá elegir una de las encuestas presentes y se le mostrará una ventana similar a la figura 5.22.

Para evitar que un alumno pueda contestar a la misma encuesta varias veces, la aplicación no permite a un alumno evaluar más de una vez la encuesta para el mismo profesor-asignatura. De esta forma estamos asegurando la fiabilidad de los resultados obtenidos así como evitamos el uso mal intencionado de la misma. Además, aunque el alumno deba identificarse para acceder al sistema, **no se almacenará información relativa al alumno** que ha rellenado una encuesta una vez que se almacena ésta en la base de datos, dado que el sistema garantiza su anonimato. Hechas estas importantes aclaraciones, pasamos a describir las operaciones a realizar desde esta interfaz:

Abrir una encuesta

El primer paso que cualquier alumno ha de realizar es la apertura de la encuesta que va a rellenar. Para esto, ha de pulsar el botón . Inmediatamente aparecerá una nueva ventana en la que el usuario deberá elegir una serie de opciones relacionadas con la asignatura y profesor a evaluar. El alumno irá seleccionando sucesivamente la asignatura a evaluar, el profesor que la imparte y el nombre de la encuesta que va a responder (figura 5.23).



Figura 5.23: Selección del perfil de la encuesta

Como se ha comentado anteriormente, en la lista desplegable que hace referencia al nombre de la encuesta a rellenar no aparecen las encuestas que hayan sido contestadas previamente por el alumno. El último paso es pulsar el botón **Ok**, y aparecerá una ventana similar a la que podemos ver en la figura 5.24.

Respondiendo a la Encuesta

En este paso el alumno pasará a completar la encuesta y según podemos observar en la figura 5.24 la interfaz muestra el contenido de la encuesta así como una serie de botones que se detallan a continuación.

En esta ventana podemos identificar tres zonas diferentes:

Archivo Base de Datos Ayuda

Enviar Encuesta Cancelar Encuesta

Preguntas:

Preguntas sobre educación:

- 1: El profesor informa del programa de la asignatura cuando comienza a impartirla: Valoración lingüística sin valoración
- 2: Informa de los objetivos del programa de la asignatura: Valoración lingüística sin valoración
- 3: Cuando asistes a sus tutorías en el horario establecido, te atiende: Valoración enumerada sin valoración
- 4: Motiva a los alumnos para que se interesen en la asignatura: Valoración lingüística sin valoración

Preguntas sobre equipamiento:

- 1: Utiliza recursos didácticos (transparencias, pizarra, medios audiovisuales, informáticos, etc.) que ayudan a comprender los contenidos: Valoración lingüística sin valoración
- 2: El programa contiene información bibliográfica útil para el desarrollo de la asignatura: Valoración invertebral sin valoración

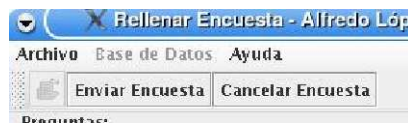
Preguntas sobre el alumno:

- 1: Valora globalmente al profesor de esta asignatura: Valoración lingüística sin valoración
- 2: En general hace interesantes las clases: Valoración lingüística sin valoración
- 3: Utiliza una metodología de enseñanza adecuada a las características del grupo y de la asignatura: Valoración numérica sin valoración
- 4: Tiene un trato igualitario con todos los alumnos: Valoración lingüística sin valoración

Valor: No evalúa... No evaluado Aceptar

Figura 5.24: Preguntas de la encuesta a rellenar

1. Botones para enviar (almacenarla en la base de datos) o cancelar una encuesta. Se encuentran en la parte superior de la ventana y por su nombre indican la acción que tienen asignada una vez contestadas todas las preguntas del cuestionario, se puede enviar la encuesta que quedará almacenada en la base de datos del sistema de evaluación. También se da la posibilidad al alumno de cancelar las respuestas de la encuesta, no serán almacenadas en la base de datos, y de esta forma podrá rellenarla en otro momento.



2. Preguntas a rellenar. En la parte central de la ventana se enumeran las diferentes preguntas sometidas a evaluación y agrupadas en sus diferentes categorías. En este ejemplo concreto y siguiendo el modelo de encuesta llevado a cabo por el Centro Andaluz de Prospectiva del curso académico 2003/04, se han presentado 10 preguntas encuadradas en tres categorías: educación, equipamiento y alumnado.
3. Evaluación de las preguntas. Situados en la parte inferior de la ventana, aparece una lista desplegable y dos botones para responder a cada una de las preguntas. Tras seleccionar el valor adecuado según el tipo de pregunta a la que esté contestando, el usuario obligatoriamente ha de pulsar en el botón **Aceptar** para que su valoración tenga efecto. Si desea cambiar de opinión y no evaluar una pregunta, ha de pulsar en el botón **No evaluado**.



En la encuesta pueden utilizarse hasta cuatro tipos diferentes de valoraciones

para las preguntas, según el diseño realizado por el administrador del sistema de evaluación: *numérica, intervalar, lingüística o enumerada*. Cuando el usuario seleccione una pregunta que aún no ha sido contestada (aparece destacada en color rojo), la aplicación identifica el tipo de pregunta y presenta en la lista desplegable **Valor** el dominio adecuado a esta pregunta, que pueden ser uno de los siguientes:

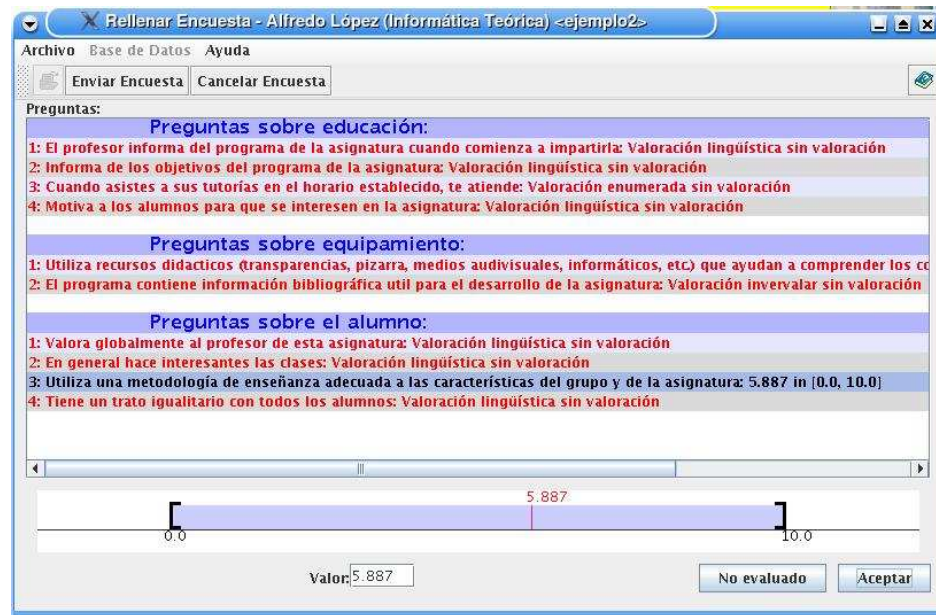


Figura 5.25: Respuesta numérica

1. Respondiendo una Pregunta Evaluada Numéricamente.

En la parte inferior de la ventana aparecerá una representación gráfica del dominio de la pregunta. Concretamente aparecerá un intervalo donde los límites representan los valores mínimo y máximo permitidos para la respuesta de la pregunta en cuestión. El usuario utilizará el ratón para marcar un valor específico dentro del intervalo o el teclado para introducir un valor numérico válido en el campo **Valor** (figura 5.25). El botón **No evaluado** nos permitirá

eliminar el valor para esa pregunta que aparece en el campo **Valor**. Finalmente, tendrá que pulsar el botón **Aceptar** para asignar el valor seleccionado a la pregunta. Si observamos la pregunta que hemos contestado, se puede comprobar como ha cambiado de color y como al final del enunciado de la misma se han añadido el valor seleccionado y el rango válido.

Archivo Base de Datos Ayuda

Enviar Encuesta Cancelar Encuesta

Preguntas:

Preguntas sobre educación:

1: El profesor informa del programa de la asignatura cuando comienza a impartirla: Valoración lingüística sin valoración

2: Informa de los objetivos del programa de la asignatura: Valoración lingüística sin valoración

3: Cuando asistes a sus tutorías en el horario establecido, te atiende: Valoración enumerada sin valoración

4: Motiva a los alumnos para que se interesen en la asignatura: Valoración lingüística sin valoración

Preguntas sobre equipamiento:

1: Utiliza recursos didácticos (transparencias, pizarra, medios audiovisuales, informáticos, etc) que ayudan a comprender los cc

2: El programa contiene información bibliográfica útil para el desarrollo de la asignatura: [4.934, 6.521] in [0.0, 10.0]

Preguntas sobre el alumno:

1: Valora globalmente al profesor de esta asignatura: Valoración lingüística sin valoración

2: En general hace interesantes las clases: Valoración lingüística sin valoración

3: Utiliza una metodología de enseñanza adecuada a las características del grupo y de la asignatura: 5.887 in [0.0, 10.0]

4: Tiene un trato igualitario con todos los alumnos: Valoración lingüística sin valoración

0.0 4.934 6.521 10.0

Valor mínimo: 4.934 Valor máximo: 6.521 No evaluado Aceptar

Figura 5.26: Respuesta intervalar

2. Respondiendo una Pregunta Evaluada con Intervalos.

La respuesta de este tipo de preguntas se caracteriza porque el usuario establece un rango de valores, caracterizado por su valor mínimo y máximo, como respuesta a una pregunta. Este tipo de respuesta es menos precisa que la anterior pero más adecuada para preguntas donde el diseñador de la encuesta considera que por algún motivo (por ejemplo la falta de información) el usuario se puede sentir más cómodo o bien sería más apropiado utilizar un

intervalo de valores en vez de un único valor para expresar su opinión.

En este caso, en la parte inferior de la ventana se presentará una representación gráfica similar a la anterior pero con la diferencia de que ahora se podrá establecer un subrango de valores caracterizado por su límite inferior y superior. Para indicar ambos valores se puede utilizar el ratón pulsando primero para el rango inferior y luego en el rango superior. O utilizando el teclado para rellenar los campos **Valor mínimo:** y **Valor máximo:** (figura 5.26). Por último se pulsa el botón **Aceptar** para asignar el valor seleccionado a la pregunta. El botón **No evaluado** nos permitirá eliminar el valor para esa pregunta.

Rellenar Encuesta - Alfredo López (Informática Teórica) <ejemplo2>

Archivo Base de Datos Ayuda

Enviar Encuesta Cancelar Encuesta

Preguntas:

Preguntas sobre educación:

1: El profesor informa del programa de la asignatura cuando comienza a impartirla: Valoración lingüística sin valoración

2: Informa de los objetivos del programa de la asignatura: Valoración lingüística sin valoración

3: Cuando asistes a sus tutorías en el horario establecido, te atiende: Valoración enumerada sin valoración

4: Motiva a los alumnos para que se interesen en la asignatura: Valoración lingüística sin valoración

Preguntas sobre equipamiento:

1: Utiliza recursos didácticos (transparencias, pizarra, medios audiovisuales, informáticos, etc) que ayudan a comprender los contenidos: [4.934, 6.521] in [0.0, 10.0]

2: El programa contiene información bibliográfica útil para el desarrollo de la asignatura: [4.934, 6.521] in [0.0, 10.0]

Preguntas sobre el alumno:

1: Valora globalmente al profesor de esta asignatura: Valoración lingüística sin valoración

2: En general hace interesantes las clases: Valoración lingüística sin valoración

3: Utiliza una metodología de enseñanza adecuada a las características del grupo y de la asignatura: 5.887 in [0.0, 10.0]

4: Tiene un trato igualitario con todos los alumnos: Valoración lingüística sin valoración

Valor: Normal

No evaluado

Nada

Poco

Normal

No evaluado Aceptar

Figura 5.27: Respuesta lingüística

3. Respondiendo una Pregunta Evaluada Lingüísticamente.

Este tipo está pensado para preguntas que por su contenido o naturaleza la respuesta más adecuada es un valor lingüístico. Existen aspectos o indicadores a valorar donde el uso de un cuantificador lingüístico del tipo bueno, malo, mejor, etc, es más apropiado que el uso de un cuantificador numérico.

Para este tipo de respuestas, la ventana mostrará en su parte inferior una lista desplegable con los posibles valores lingüísticos válidos para la pregunta seleccionada (figura 5.27). El usuario seleccionará el valor deseado y pulsará el botón **Aceptar** para que quede registrado. También podrá establecer la respuesta como no evaluada pulsando los botones **No evaluado**.

4. Respondiendo una Pregunta Enumerada.

La última clase del tipo de respuestas que se pueden utilizar en un cuestiona-

rio son las enumeradas. Este caso es similar a las preguntas lingüísticas con la diferencia de que sólo admiten dos posibles respuestas del tipo Verdadero/Falso o Sí/No. En la parte inferior de la ventana se muestran los posibles valores, seleccionándose uno de ellos. Al igual que para el resto de respuestas, es necesario pulsar el botón **Aceptar** para que la contestación sea registrada. También se puede dejar la pregunta como no evaluada pulsando el botón **No evaluado**.

Enviar o cancelar la encuesta

Una vez que se haya respondido a todas las preguntas se podrán guardar los resultados en la base de datos, pulsando el botón **Enviar Encuesta**. Por otro lado, si no se desea guardar los resultados y descartar las respuestas, se pulsará el botón **Cancelar Encuesta**.

Si por algún motivo el usuario intenta abandonar este módulo sin haberlo enviado, aparecerá un mensaje informando de esta situación y solicitando al alumno si se desea o no enviar la encuesta (figura 5.28)

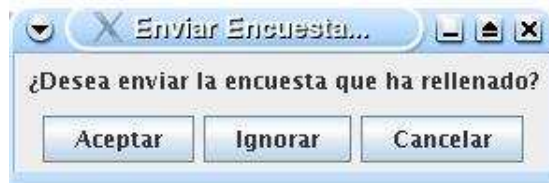


Figura 5.28: Aviso de encuesta no enviada

5.4. Conclusiones

En este capítulo se ha mostrado el diseño e implementación de un prototipo en Java para un sistema de evaluación de la calidad docente universitaria que forma

parte de un proyecto realizado para la convocatoria de Grupos de Estudio y Análisis Específicos sobre Calidad en las Universidades Andaluzas de la Unidad para la Calidad de las Universidades Andaluzas (UCUA) en los años 2005 y 2006 [106]. Este sistema de evaluación está basado en el modelo de decisión para problemas de TD en contextos heterogéneos donde permite a los encuestados tener una mayor flexibilidad a la hora de expresar su evaluación y no sólo restringirse a una escala entre el 1 y el 5.

Conclusiones

A continuación, revisamos cuáles han sido los resultados más destacables que han sido obtenidos a lo largo de esta memoria de investigación, y posteriormente propondremos cuáles serán las propuestas futuras que pensamos desarrollar a partir de estos resultados.

Resultados Obtenidos

El modelado de preferencias se ha aplicado a un gran número de áreas, en esta memoria nuestro interés se ha centrado en el uso del modelado de preferencias en la toma de decisiones y problemas de evaluación, en los que, las preferencias podrían estar modeladas en contextos heterogéneos mediante diferentes dominios de expresión, tales como, *numérico*, *intervalar* y *lingüístico*. Dado que la resolución de un problema de decisión tiene dos procesos principales, tales como:

1. *Proceso de Agregación.*
2. *Proceso de Explotación.*

Estos procesos en contextos heterogéneos implican operar con este tipo de información. Para cada modelo de representación existen operadores y herramientas que nos permiten trabajar con dicho modelado, pero el uso conjunto de los distintos

tipos de información implican la utilización de nuevas herramientas y operadores que nos permitan manejarlos de forma conjunta. Los modelos previos nos permitían operar con contextos heterogéneos de forma limitada, ya que no existían herramientas u operadores diseñados de forma específica para tratar con información lingüística, numérica e intervalar de forma conjunta. Y los modelos para operar con algunos tipos de contextos heterogéneos nos proporcionaban resultados que eran complejos de entender para los decisores que formaban parte del problema de decisión. Atendiendo a estos aspectos, los resultados en esta memoria pueden resumirse en los siguientes apartados:

A. Un Nuevo Modelo de Decisión para Problemas Definidos en Contextos Heterogéneos

El principal problema que presenta el manejar información heterogénea compuesta por datos numéricos, intervalares y lingüísticos es cómo operar de forma conjunta sobre ella. El modelo de decisión propuesto en esta memoria soluciona este problema fundamentándose en:

- Un proceso de unificación de la información heterogénea en un único dominio de expresión como son, los conjuntos difusos sobre un conjunto básico de términos lingüísticos. Para realizar este proceso de unificación se han definido diferentes funciones de transformación para cada modelado de representación de la información (numérico, lingüístico e intervalar) en conjuntos difusos sobre el conjunto básico de términos lingüísticos.
- El uso de la aritmética difusa para realizar las operaciones que implica el proceso de agregación sobre la información unificada.
- Finalmente, los procesos de explotación en los problemas de TD suelen im-

plicar operaciones de ordenación, comparación, etc., lo que para facilitar estos procesos, como para mejorar la comprensión de los resultados obtenidos en esta fase, se ha definido una transformación de la información unificada (conjuntos difusos) a 2-tuplas lingüísticas.

Este modelo de decisión nos permitirá resolver problemas de decisión definidos en contextos de información heterogénea.

B. Un Diseño UML para Implementar un Sistema de Soporte a la Decisión para Problemas Definidos en Contextos Heterogéneos

En la literatura existen un gran número de modelos de decisión que nos permiten abordar problemas de TD, pero la resolución de este tipo de problemas desde un punto de vista informático se lleva a cabo implementando estos modelos en un sistema automático. En esta memoria hemos realizado un diseño utilizando UML como herramienta para el diseño de dos bibliotecas de clases que nos ayudarán a diseñar e implementar sistemas de ayuda a la decisión. Estas bibliotecas son:

1. Biblioteca para el tratamiento de la información heterogénea. Presenta el diseño de las clases necesarias que nos ayudarán a trabajar con información heterogénea para problemas que tratan con información numérica, intervalar y lingüística. Además en el diseño se ha tenido en cuenta la posible incorporación de nuevos dominios de información para que sólo deban ser añadidos a la biblioteca con los mínimos cambios necesarios.
2. Biblioteca para el modelo de decisión heterogéneo. Esta biblioteca se ha diseñado con las clases necesarias que nos permitan implementar el modelo de decisión para problemas definidos en contextos heterogéneos presentado en esta memoria.

Con estas dos bibliotecas se facilitará el diseño e implementación de cualquier sistema de ayuda a la decisión, de forma simple y eficaz.

C. Un Prototipo de un Sistema de Evaluación

Utilizando el análisis y diseño del sistema de ayuda a la decisión anterior, hemos implementado un prototipo para un sistema de evaluación definido en un contexto heterogéneo con información numérica, intervalar y lingüística y cuyo proceso de evaluación se basa en el modelo de decisión presentado en la memoria. Este prototipo proporciona las siguientes funcionalidades:

1. Diseño de encuestas con un modelado de preferencias heterogéneas.
2. Automatización del proceso de rellenado de la encuesta.
3. Aplicaciones de distintos modelos de resolución al proceso de evaluación, dependientes del problema específico o las necesidades de los evaluadores.

Este prototipo ha sido desarrollado en colaboración con la Unidad de Calidad de las Universidades Andaluzas en un proyecto de la Convocatoria de Grupos Específicos y Análisis de la Calidad de las Universidades Andaluzas, con el objetivo de estudiar su implantación o mejoras que estos métodos pueden aportar sobre los métodos actuales.

Trabajos Futuros

En el mundo real hay un gran número de situaciones que pueden considerarse dentro del ámbito de la toma de decisiones, por ejemplo, *el diagnóstico clínico, la gestión comercial, la planificación, predicciones económicas, etc.* En todos estos casos la utilización de la toma de decisiones asistida por ordenador sería altamente beneficiosa, tanto en aspectos económicos como sociales. Debido a esto, nuestros trabajos futuros se encaminan en las siguientes líneas de acción:

- *Teórica.*

1. Aumentar y mejorar el número de modelos, herramientas para el tratamiento de información heterogénea, proporcionando la capacidad de operar con otros tipos de información o estructuras de representación.
2. Desarrollo de un marco claro para definir el modelado de preferencias según su naturaleza. Ésta es una necesidad que demandan los usuarios una vez que se han desarrollado e implementado herramientas para manejar estos contextos. En Psicología y Estadística hay estudios sobre el modelado de la información atendiendo a la in/certidumbre existente sobre los aspectos valorados, pero no crean una metodología precisa para hacerlo. Por tanto, el uso de modelos empíricos de la Psicología que ayuden a definir las valoraciones y rangos posibles para valorar las alternativas, es una línea de estudio imprescindible para utilizar procesos de TD en contextos heterogéneos en problemas reales.

- *Práctica.*

Desarrollo de nuevos Sistemas de Ayuda a la Decisión basados en el diseño presentado en esta memoria, aplicados a problemas diferentes a los que

nos enfrentamos diariamente en el mundo real y que se adapten al modelo propuesto como:

1. Problemas de evaluación sensorial, donde la información proporcionada por los decisores está basado en conocimiento adquirido por sus sentidos.
2. Problemas de TD multi-experto multi-atributo relacionadas con la consultoría tecnológica.
3. Problemas de evaluación de desempeño donde participan expertos con distinto grado de conocimiento sobre los elementos a evaluar.

Bibliografía

- [1] G.I. Adamopoulos and G.P. Pappis. A fuzzy linguistic approach to a multicriteria sequencing problem. *European Journal of Operational Research*, 92:628–636, 1996.
- [2] C. Alcalde, A. Burusco, and R. Fuentes-Gonzalez. A constructive method for the definition of interval-valued fuzzy implication operators. *Fuzzy Sets and Systems*, 153(2):211–227, 2005.
- [3] B. Arfi. Fuzzy decision making in politics: A linguistic fuzzy-set approach (lfsa). *Political Analysis*, 13(1):23–56, 2005.
- [4] B. Arfi. Linguistic fuzzy-logic game theory. *Journal of Conflict Resolution*, 50(1):28–57, 2006.
- [5] W. Armstrong. Uncertainty and utility function. *Economics Journal*, 58:1–10, 1948.
- [6] K. Arnold, J. Gosling, and D. Holmes. *El Lenguaje de Programación JAVA*. 3^a Ed. Addison Wesley, 2001.
- [7] K. Atanassov and G. Gargov. Interval valued intuitionistic fuzzy-sets. *Fuzzy Sets and Systems*, 31(3):343–349, 1989.

- [8] D. Ben-Arieh and C. Zhifeng. Linguistic labels aggregation and consensus measure for autocratic decision-making using group recommendations. *IEEE Transactions on Systems, Man, and Cybernetics Part A - Systems and Humans*, In Press, 2006.
- [9] P.P. Bonissone. *A fuzzy sets based linguistic approach: Theory and applications*, pages 329–339. Approximate Reasoning in Decision Analysis, North-Holland, 1982.
- [10] P.P. Bonissone. *A fuzzy sets based linguistic approach: theory and applications*. Approximate Reasoning in Decision Analysis, North-Holland, 1982. 329-339.
- [11] P.P. Bonissone and K.S. Decker. *Selecting Uncertainty Calculi and Granularity: An Experiment in Trading-Off Precision and Complexity*. In L.H. Kanal and J.F. Lemmer, Editors., *Uncertainty in Artificial Intelligence*. North-Holland, 1986.
- [12] G. Booch. *Object-Oriented Analysis and Design with Applications*, 2a. ed. Benjamin/Cummings, 1994.
- [13] G. Booch, J. Rumbaugh, and I. Jacobson. *El Lenguaje Unificado de Modelado*. Addison Wesley, 2003.
- [14] G. Bordogna, M. Fedrizzi, and G. Pasi. A linguistic modeling of consensus in group decision making based on OWA operators. *IEEE Transactions on Systems, Man and Cybernetics, Part A: Systems and Humans*, 27:126–132, 1997.

- [15] G. Bordogna and G. Pasi. A fuzzy linguistic approach generalizing boolean information retrieval: A model and its evaluation. *Journal of the American Society for Information Science*, 44:70–82, 1993.
- [16] P. Bosc, D. Kraft, and F. Petry. Fuzzy sets in database and information systems: Status and opportunities. *Fuzzy Sets and Systems*, 3(156):418–426, 2005.
- [17] B. Bouchon-Meunier and M. Rifqi. Resemblance in database utilization. *6th IFSA World Congress*, 1995.
- [18] B. Bouchon-Meunier, M. Rifqi, and S. Bothorel. Towards general measures of comparison of objects. *Fuzzy Sets and Systems*, (84):143–153, 1996.
- [19] D. Bouyssou, T. Marchant, M. Pirlot, P. Perny, and A. Tsoukia's. *Evaluation and Decision Models: A critical perspective*. Kluwer Academic Publishers, 2000.
- [20] B. Bruegge and A.H. Dutoit. *Ingeniería de Software Orientado a Objetos*. Addison Wesley, 2002.
- [21] Z. Bubnicki. *Analysis and Decision Making in Uncertain Systems*. Springer-Verlag, 2004.
- [22] T.X. Bui. *A Group Decison Support System for Cooperative Multiple Criteria Group Decison-Making*. Springer-Verlag, 1987.
- [23] H. Bustince, E. Barrenechea, and V. Mohedano. Intuitionistic fuzzy implication operators: An expression and main properties. *International Journal of Uncertainty Fuzziness and Knowledge-Based Systems*, 12(3):387–406, 2004.

- [24] H. Bustince and P. Burillo. Perturbation of intuitionistic fuzzy relations. *International Journal of Uncertainty Fuzziness and Knowledge-Based Systems*, 9(1):81–103, 2001.
- [25] R. Caballero and G.M. Fernández. *Toma de decisiones con criterios múltiples*. Rect@. Tirant lo Blanch, 2002.
- [26] E. Capurso and A. Tsoukiàs. Decision aiding and psychotherapy. *Bulletin of the EURO Working Group on MCDA*, 2003.
- [27] C. Carlsson and R. Fuller. *Fuzzy Reasoning in Decision Making and Optimization*, volume 82 of *Studies in Fuzziness and Soft Computing*. Studies in Fuzziness and Soft Computing Series, 2001.
- [28] S.-L. Chang, R.-C. Wang, and S.-Y. Wang. Applying a direct multi-granularity linguistic and strategy-oriented aggregation approach on the assessment of supply performance. *European Journal of Operational Research*, 172(2):1013–1025, 2007.
- [29] C.T. Chen. Applying linguistic decision-making method to deal with service quality evaluation problems. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 9(Suppl.):103–114, 2001.
- [30] S.J. Chen and C.L. Hwang. *Fuzzy multiple attribute decision-making methods and applications*. Springer-Verlag, 1992.
- [31] C.H. Cheng and Y. Lin. Evaluating the best main battle tank using fuzzy decision theory with linguistic criteria evaluation. *European Journal of Operational Research*, 142:174–186, 2002.

- [32] C.H. Cheng, K.L. Yang, and C.L. Hwang. Evaluating attack helicopters by ahp based on linguistic variable weight. *European Journal of Operational Research*, 116(2):423–435, 1999.
- [33] H. Chernoff. *Elementary Decision Theory*. Dover Publications, 1987.
- [34] S.-J. Chuu. Fuzzy multi-attribute decision-making for evaluating manufacturing flexibility. *Production Planning and Control*, 16(3):323–335, 2005.
- [35] R.T. Clemen. *Making Hard Decisions. An Introduction to Decision Analysis*. Duxbury Press, 1995.
- [36] C. Coombs and J. Smith. On the detection of structures in attitudes and developmental processes. *Psychological Reviews*, 80(5):337–351, 1973.
- [37] R. López de Mántaras and L. Godó. From intervals to fuzzy truth-values: Adding flexibility to reasoning under uncertainty. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 5(3):251–260, 1997.
- [38] G. Debreu. *Theory of Value: An Axiomatic Analysis of Economic Equilibrium*. John Wiley and Sons Inc., 1959.
- [39] R. Degani and G. Bortolan. The problem of linguistic approximation in clinical decision making. *International Journal of Approximate Reasoning*, 2:143–162, 1988.
- [40] M. Delgado, F. Herrera, E. Herrera-Viedma, and L. Martínez. Combining numerical and linguistic information in group decision making. *Information Sciences*, 107:177–194, 1998.
- [41] M. Delgado, J.L. Verdegay, and M.A Vila. Linguistic decision making models. *International Journal of Intelligent Systems*, 7:479–492, 1992.

- [42] M. Delgado, J.L. Verdegay, and M.A. Vila. Linguistic decision making models. *International Journal of Intelligent Systems*, 7:479–492, 1993.
- [43] M. Delgado, M.A. Vila, and W. Voxman. On a canonical representation of fuzzy numbers. *Fuzzy Sets and Systems*, 94:125–135, 98.
- [44] L. Dombi. *Fuzzy Logic and Soft Computing*, chapter A General Framework for the Utility-Based and Outranking Methods, pages 202–208. World Scientific, 1995.
- [45] J. Doyle. Prospects for preferences. *Computational Intelligence*, 20(2):111–136, 2004.
- [46] D. Dubois and H. Prade. *Fuzzy Sets and Systems: Theory and Applications*. Kluwer Academic, New York, 1980.
- [47] D. Dubois and H. Prade. Fuzzy set-theoretic differences and inclusions and their use in fuzzy arithmetics and analysis. Linz, Austria, 1983. In 5th International Seminar on Fuzzy Set Theory.
- [48] D. Dubois and H. Prade. Rough fuzzy-sets and fuzzy rough sets. *International Journal of General Systems*, 13(2-3):191–209, 1990.
- [49] R. Duncan and H. Raiffa. *Games and Decision. Introduction and Critical Survey*. Dover Publications, 1985.
- [50] Z-P. Fan, J. Ma, and Q. Zhang. An approach to multiple attribute decision making based on fuzzy preference information alternatives. *Fuzzy Sets and Systems*, 131(1):101–106, 2002.

- [51] Z.P. Fan and X. Chen. Consensus measures and adjusting inconsistency of linguistic preference relations in group decision making. *Lecture Notes in Artificial Intelligence*, 3613:130–139, 2005.
- [52] Z.P. Fan, S.H. Xiao, and G.F. Hu. An optimization method for integrating two kinds of preference information in group decision-making. *Computers & Industrial Engineering*, 46(2):329–335, 2004.
- [53] J. Fodor and M. Roubens. *Fuzzy preference modelling and multicriteria decision support*. Kluwer Academic Publishers, Dordrecht, 1994.
- [54] P. Fortemps and R. Slowinski. A graded quadrivalent logic for ordinal preference modelling: Loyola-like approach. *Fuzzy Optimization and Decision Making*, 1:93–111, 2002.
- [55] J.L. Garcia-Lapresta. A general class of simple majority decision rules based on linguistic opinions. *Information Sciences*, 176(4):352–365, 2006.
- [56] R.A. Gheorghe, A. Bufardi, and P. Xirouchakis. Fuzzy multicriteria decision aid method for conceptual design. *Cirp Annals-Manufacturing Technology*, 54(1):151–154, 2005.
- [57] D. Gomez and J. Montero. A discussion on aggregation operators. *Kybernetika*, 40(1):107–120, 2004.
- [58] S. Greco, B. Matarazzo, and R. Slowinski. Rough sets theory for multicriteria decision analysis. *European Journal of Operational Research*, 129(1):1–47, 2001.
- [59] F. Herrera and E. Herrera-Viedma. Linguistic decision analysis: Steps for solving decision problems under linguistic information. *Fuzzy Sets and Systems*, 115:67–82, 2000.

- [60] F. Herrera, E. Herrera-Viedma, and L. Martínez. A fusion approach for managing multi-granularity linguistic term sets in decision making. *Fuzzy Sets and Systems*. 114 43-58, 114:43–58, 2000.
- [61] F. Herrera, E. Herrera-Viedma, L. Martinez, F. Mata, and P.J. Sanchez. *Combining Heterogeneous Information in GDM*. Intelligent Systems for Information Processing: From Representation to Applications. Elsevier, Bernadette Bouchon-Meunier, Laurent Foulloy, Ronald R. Yager (Eds.), 2003.
- [62] F. Herrera, E. Herrera-Viedma, L. Martinez, F. Mata, and P.J. Sanchez. *A Multi-Granular Linguistic Decision Model for Evaluating the Quality of Network Services*. Intelligent Sensory Evaluation: Methodologies and Applications. Springer, Ruan Da, Zeng Xianyi (Eds.), 2004.
- [63] F. Herrera, E. Herrera-Viedma, L. Martínez, and P.J. Sánchez. A linguistic decision process for evaluating the installation of an ERP system. In *9th International Conference on Fuzzy Theory and Technology*, pages 164–167, Cary (North Carolina) USA, 2003.
- [64] F. Herrera, E. Herrera-Viedma, and J.L. Verdegay. A sequential selection process in group decision making with linguistic assessment. *Information Sciences*, 85:223–239, 1995.
- [65] F. Herrera, E. Herrera-Viedma, and J.L. Verdegay. A sequential selection process in group decision making with linguistic assessment. *Information Sciences*, 85:223–239, 1995.
- [66] F. Herrera, E. Herrera-Viedma, and J.L. Verdegay. Direct approach processes in group decision making using linguistic OWA operators. *Fuzzy Sets and Systems*, 79:175–190, 1996.

- [67] F. Herrera, E. Herrera-Viedma, and J.L. Verdegay. A model of consensus in group decision making under linguistic assessments. *Fuzzy Sets and Systems*, 79:73–87, 1996.
- [68] F. Herrera and L. Martínez. A 2-tuple fuzzy linguistic representation model for computing with words. *IEEE Transactions on Fuzzy Systems*, 8(6):746–752, 2000.
- [69] F. Herrera and L. Martínez. An approach for combining linguistic and numerical information based on 2-tuple fuzzy representation model in decision-making. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 8(5):539–562, 2000.
- [70] F. Herrera and L. Martínez. The 2-tuple linguistic computational model. Advantages of its linguistic description, accuracy and consistency. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 9(Suppl.):33–49, 2001.
- [71] F. Herrera and L. Martínez. A model based on linguistic 2-tuples for dealing with multigranularity hierarchical linguistic contexts in multiexpert decision-making. *IEEE Transactions on Systems, Man and Cybernetics. Part B: Cybernetics*, 31(2):227–234, 2001.
- [72] F. Herrera, L. Martínez, and P.J. Sánchez. Integration of heterogeneous information in decision-making problems. In *Proceedings of the Eighth International Conference on Information Processing and Management of Uncertainty in Knowledge-Bases Systems. IPMU'2000*, volume Vol I, pages 633–640, Madrid (Spain), 2000.

- [73] F. Herrera, L. Martínez, and P.J. Sánchez. Managing non-homogeneous information in group decision making. *European Journal of Operational Research*, 166(1):115–132, 2005.
- [74] E. Herrera-Viedma, F. Herrera, and F. Chiclana. A consensus model for multiperson decision making with different preference structures. *IEEE Transactions on Systems, Man and Cybernetics-A*, 32(3):394–402, 2002.
- [75] E. Herrera-Viedma, L. Martínez, F. Mata, and F. Chiclana. A consensus support system model for group decision-making problems with multi-granular linguistic preference relations. *IEEE Transactions on Fuzzy Systems*, 13(5):644–658, 2005.
- [76] U. Hohle. *Mathematics of fuzzy sets*. Kluwer Academic Publishers, 1998.
- [77] V.N. Huynh and Y. Nakamori. A satisfactory-oriented approach to multi-expert decision-making with linguistic assessments. *IEEE Transactions On Systems Man And Cybernetics Part B-Cybernetics*, 35(2):184–196, 2005.
- [78] M. Inuiguchi. Generalizations of rough sets: From crisp to fuzzy cases. *Lecture Notes in Artificial Intelligence*, 3066:26–37, 2004.
- [79] I. Jacobson, G. Booch, and J. Rumbaugh. *El Proceso Unificado de Desarrollo de Software*. Addison Wesley, 2001.
- [80] I. Jacobson, M. Christerson, P. Jonsson, and G. Overgaard. *Object-Oriented Software Engineering. A Use Case Driven Approach*. Addison Wesley, 1992.
- [81] A. Jiménez, S. Ríos-Insua, and A. Mateos. A decision support system for multiattribute utility evaluation based on imprecise assignments. *Decision Support Systems*, 36(1):65–79, 2003.

- [82] J. Kacprzyk. Group decision making with a fuzzy linguistic majority. *Fuzzy Sets and Systems*, 18:105–118, 1986.
- [83] J. Kacprzyk. *The Analysis of Fuzzy Information*, chapter On Some Fuzzy Cores and “Soft” Consensus Measures in Group Decision Making, pages 119–130. In: J. Bezdek. (Ed.). CRC Press, 1987.
- [84] J. Kacprzyk and M. Fedrizzi. *Multi-Person Decision Making Using Fuzzy Sets and Possibility Theory*. Kluwer Academic Publishers, Dordrecht, 1990.
- [85] J. Kacprzyk, M. Fedrizzi, and H. Nurmi. Group decision making and consensus under fuzzy preferences and fuzzy majority. *Fuzzy Sets and Systems*, 49:21–31, 1992.
- [86] J. Kacprzyk, H. Nurmi, and M. Fedrizzi. *Consensus under Fuzziness*. Kluwer Academic Publishers, 1997.
- [87] D. Kahneman, D. Ruan, and I. Doğan. Fuzzy group decision-making for facility location selection. *Information Sciences*, 157(1-4):135–153, 2003.
- [88] D. Kahneman, P. Slovic, and A. Tversky. *Judgement under uncertainty: Heuristics and biases*. Cambridge University Press, 1981.
- [89] A. Kaufman. *Introduction À la Théorie Des Sous-Ensembles Flous*, volume 3. 1973.
- [90] R.L. Keeney and H. Raiffa. *Decisions with Multiple Objectives: Prefereces and Value Tradeoffs*. Cambridge University Press, 1993.
- [91] G.J. Klir and B. Yuan. *Fuzzy sets an fuzzy logic: Theory and Applications*. Prentice-Hall PTR, 1995.

- [92] S. Kundu. Min-transitivity of fuzzy leftness relationship and its application to decision making. *Fuzzy Sets and Systems*, 86:357–367, 1997.
- [93] S. Kundu. Preference relation on fuzzy utilities based on fuzzy leftness relation on intervals. *Fuzzy Sets and Systems*, (97):183–191, 1998.
- [94] H.M. Lee. Applying fuzzy sets theory for evaluating the rate of aggregative risk in software development. *Fuzzy Sets and Systems*, 79:323–336, 1996.
- [95] H.M. Lee. Generalization of the group decision making using fuzzy sets theory for evaluating the rate of aggregate risk in software development. *Information Sciences*, 113:301–311, 1999.
- [96] HS Lee and M. O’Mahony. Sensory evaluation and marketing: measurement of a consumer concept. *Food Quality And Preference*, 16(3):227–235, 2005.
- [97] E. Levrat, A. Voisin, S. Bombardier, and J. Bremont. Subjective evaluation of car seat comfort with fuzzy set techniques. *International Journal of Intelligent Systems*, 12:891–913, 1997.
- [98] D. Li and J.B. Yang. A multiattribute decision making approach using intuitionistic fuzzy sets. In *Proceedings Eusflat 2003*, pages 183–186, Zitaú, 2003.
- [99] R.D. Luce and P. Suppes. *Handbook of Mathematical Psychology*, chapter Preferences, Utility and Subject Probability, pages 249–410. Wiley, 1965.
- [100] J. Ma, D. Ruan, Y. Xu, and G. Zhang. A fuzzy-set approach to treat determinacy and consistency of linguistic terms in multi-criteria decision making. *International Journal of Approximate Reasoning*, 44(2):165–181, 2007.

- [101] P. Maestre. Business intelligence: del erp y kim, al asp y crm. In *I Observatorio Dintel*, Madrid, 2002.
- [102] Marimin, M. Umano, I. Hatono, and H. Tamura. Linguistic labels for expressing fuzzy preference relation in fuzzy group decision making. *IEEE Transaction on Systems, Man and Cybernetics, Part B: Cybernetics*, 28:205–218, 1998.
- [103] L. Martínez. Sensory evaluation based on linguistic decision analysis. *International Journal of Approximated Reasoning*, 44(2):148–164, 2007.
- [104] L. Martínez, J. Liu, and J.B. Yang. A fuzzy model for design evaluation based on multiple-criteria analysis in engineering systems. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 14(3):317–336, 2006.
- [105] L. Martínez, J. Liu, J.B. Yang, and F. Herrera. A multi-granular hierarchical linguistic model for design evaluation based on safety and cost analysis. *International Journal of Intelligent Systems.*, 20(12):1161–1194, 2005.
- [106] L. Martínez, P.J. Sánchez, C. García, B. Montes, F. Mata, and L.G. Pérez. *Un Sistema de Evaluación Basado en Técnicas de Difusión Difusas*. UCUA, Sevilla, 2005.
- [107] G.A. Miller. The magical number seven plus or minus two: Some limits on our capacity of processing information. *Psychological Review*, 63:81–97, 1956.
- [108] J.N. Morderson and P.S. Nair. *Fuzzy mathematics*. Physica Verlag, 1998.
- [109] H. T. Nguyen, Vladik Kreinovich, and Qiang Zuo. Interval-valued degrees of belief: Applications of interval computations to expert systems

- and intelligent control. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 5(3):317–358, 1997.
- [110] G. Norris, J.R. Hurley, and et al. *E-Business and ERP. Transforming the Enterprise*. John Wiley & Sons Inc., 2000.
- [111] H. Nurmi. *Assumptions of Individual Preferences in the Theory of Voting Procedures*, pages 142–155. In: J. Kacprzyk and M. Roubens, Eds., *Non conventional Preference Relations in Decision Making*. Springer-Verlag, 1988.
- [112] M-N Omri. *Système Interactif Flou D’aide A L’utilisation de Dispositifs Techniques: Le Système SIFADE*. PhD thesis, Universit’e Pierre et Marie Curie, Paris, France, 1994.
- [113] S.A. Orlovsky. Decision-making with a fuzzy preference relation. *Fuzzy Sets Systems*, 1:155–167, 1978.
- [114] M. Oztürk, A. Tsoukiàs, and Ph. Vincke. *Preference Modelling*, pages 27–72. In: *State of the Art in Multiple Criteria Decsioin Analysis*, M. Ehrgott, S. Greco and J. Figueira (Ed.). Wiley Series on Intelligent Systems. Springer-Verlag, 2005.
- [115] W. Pedrycz. *Fuzzy modeling: Paradigms and practice*. Kluwer Academic, 1996.
- [116] W. Pedrycz and F. Gomide. *An introduction to fuzzy sets*. The MIT Press, 1998.
- [117] P. Perny and A. Tsoukias. On the continuous extension of a four Valued logic for preference modelling. pages 302–309, Paris, 1998. IPMU.

- [118] M.H. Rasmy, S.M. Lee, W.F. Abd El-Wahed, A.M. Ragab, and M.M. El-Sherbiny. An expert system for multiobjective decision making: Application of fuzzy linguistic preferences and goal programming. *Fuzzy Sets and Systems*, 127:209–220, 2002.
- [119] S. Rios, C. Bielza, and A. Mateos. *Fundamentos de los Sistemas de Ayuda a la Decisión*. Ra-Ma, 2002.
- [120] D.W. Roberts. An anticommutative difference operator for fuzzy sets and relations. *Fuzzy Sets and Systems*, (21):35–42, 1987.
- [121] C. Romero. *Teoría de la Decisión Multicriterio: Conceptos, Técnicas, Aplicaciones*. Alianza Universidad, 1993.
- [122] J-P Rossazza. *Utilisation de Hierarchie de Classes Floues Pour la Representation Des Connaissances Imprécises et Sujettes À Exception: Le Systeme SORCIER*. PhD thesis, Université Paul Sabatier, Toulouse, France, 1990.
- [123] M. Roubens. Fuzzy sets and decision analysis. *Fuzzy Sets and Systems*, 90:199–206, 1997.
- [124] M. Roubens and Ph. Vincke. *Preference modelling*. Springer-Verlag, 1985.
- [125] E.H. Ruspini. A new approach to clustering. *Inform. Control*, 15:22–32, 1969.
- [126] E.H. Ruspini. On the semantics of fuzzy logic. *International Journal Approximate Reasoning*, (5):45–88, 1991.
- [127] T.L. Saaty. *The Analytic Hierarchy Process*. Mc-Graw-Hill, New York, USA, 1980.

- [128] M. Salles. *Handbook of Utility Theory*, chapter Fuzzy Utility. Kluwer Academic Publishers, 1998.
- [129] P.J. Sánchez, L.G. Pérez, F. Mata, and A.G. López-Herrera. A multi-granular linguistic model to evaluate the suitability of installing an ERP system. *Mathware*, XII(2-3):217–233, 2005.
- [130] F. Seo and M. Sakawa. Fuzzy multiattribute utility analysis for collective choice. *IEEE Transactions on Systems, Man and Cybernetics*, 15:45–53, 1985.
- [131] M.G. Shields. *E-Business and ERP. Rapid Implementation and Project Planning*. John Wiley & Sons Inc., 2001.
- [132] Y. Tang and J. Zheng. Linguistic modelling based on semantic similarity relation among linguistic labels. *Fuzzy Sets and Systems*, 157(12):1662–1673, 2006.
- [133] T. Tanino. Fuzzy preference orderings in group decision making. *Fuzzy Sets and Systems*, 12:117–131, 1984.
- [134] T. Tanino. *On Group Decision Making Under Fuzzy Preferences*, pages 172–185. in: J. Kacprzyk and M. Fedrizzi, Eds., *Multiperson Decision Making Using Fuzzy Sets and Possibility Theory*. Kluwer Academic Publishers, 1990.
- [135] J.F. Le Teno and B. Mareschal. An interval version of PROMETHEE for the comparison of building products’ design with ill-defined data on environmental quality. *European Journal of Operational Research*, 109:522–529, 1998.

- [136] M. Tong and P.P. Bonissone. A linguistic approach to decision making with fuzzy sets. *IEEE Transactions on Systems, Man and Cybernetics*, 10:716–723, 1980.
- [137] M. Tong and P.P. Bonissone. Linguistic solution to fuzzy decision problems. *Studies in the Management Science*, 20:323–334, 1984.
- [138] V. Torra. Negation function based semantics for ordered linguistic labels. *International Journal of Intelligent Systems*, 11:975–988, 1996.
- [139] V. Torra. Linguistic aggregation in non-unified domains. In *EUROFUSE-SIC 99*, pages 188–193, Budapest, 1999.
- [140] V. Torra. Aggregation of linguistic labels when semantics is based on antonyms. *International Journal of Intelligent Systems*, 16:513–524, 2001.
- [141] E. Triantaphyllou. *Multi-Criteria DM Methods: A comparative Study*. Kluwer Academic Publishers, 2000.
- [142] I.B. Turksen. Meta-linguistic axioms as a foundation for computing with words. *Information Sciences*, 177(2):332–359, 2007.
- [143] J.H. Wang and J. Hao. An approach to computing with words based on canonical characteristic values of linguistic labels. *IEEE Transactions on Fuzzy Systems*, 15(4):593–604, 2007.
- [144] R.C. Wang and S.J. Chuu. Group decision-making using a fuzzy linguistic approach for evaluating the flexibility in a manufacturing system. *European Journal of Operational Research*, 153(3):563–572, 2004.
- [145] P.P. Wang(Ed.). *Computing with Words*. Wiley Series on Intelligent Systems. John Wiley and Sons, 2001.

- [146] Z.S. Xu. A method based on linguistic aggregation operators for group decision making with linguistic preference relations. *Information Science*, 166:19–30, 2004.
- [147] Z.S. Xu. Uncertain linguistic aggregation operators based approach to multiple attribute group decision making under uncertain linguistic environment. *Information Sciences*, 168:171–184, 2004.
- [148] Z.S. Xu. An approach to group decision making based on incomplete linguistic preference relations. *International Journal of Information Technology and Decision Making*, 4(1):153–160, 2005.
- [149] Z.S. Xu. Deviation measures of linguistic preference relations in group decision making. *Omega*, 33(3):249–254, 2005.
- [150] Z.S. Xu. Deviation measures of linguistic preference relations in group decision making. *Omega*, 33(3):249–254, 2005.
- [151] Z.S. Xu. A direct approach to group decision making with uncertain additive linguistic preference relations. *Fuzzy Optimization and Decision Making*, 5(1):21–32, 2006.
- [152] Z.S. Xu. An interactive procedure for linguistic multiple attribute decision making with incomplete weight information. *Fuzzy Optimization and Decision Making*, 6(1):17–27, 2007.
- [153] R.R. Yager. On ordered weighted averaging aggregation operators in multicriteria decision making. *IEEE Transaction on Systems, Man, and Cybernetics*, 18:183–190, 1988.
- [154] R.R. Yager. Aggregation operators and fuzzy system modelling. *Fuzzy Sets and Systems*, 67:129–145, 1993.

- [155] R.R. Yager. Non-numeric multi-criteria multi-person decision making. *Group Decision and Negotiation*, 2:81–93, 1993.
- [156] R.R. Yager. An approach to ordinal decision making. *International Journal of Approximate Reasoning*, 12:237–261, 1995.
- [157] R.R. Yager and V. Kreinovich. Decision making under interval probabilities. *International Journal of Approximate Reasoning*, 22(3):195–215, 1999.
- [158] A. Yazici and R. George. *Fuzzy Database Modeling*. Physica-Verlag, 1999.
- [159] C.Y. Yue, S.B. Yao, and P. Zhang. Rough approximation of a preference relation for stochastic multi-attribute decision problems. *Lecture Notes in Artificial Intelligence*, 3613:1242–1245, 2005.
- [160] L.A. Zadeh. Fuzzy sets. *Information and Control*, 8:338–353, 1965.
- [161] L.A. Zadeh. The concept of a linguistic variable and its applications to approximate reasoning. *Information Sciences, Part I, II, III*, 8,8,9:199–249,301–357,43–80, 1975.
- [162] L.A. Zadeh. A computational approach to fuzzy quantifiers in natural languages. *Computers and Mathematics with Applications*, (9):149–184, 1983.
- [163] L.A. Zadeh. Fuzzy logic = computing with words. *IEEE Transactions on Fuzzy Systems*, 4(2):103–111, May 1996.
- [164] L.A. Zadeh and J. Kacprzyk. *Computing with Words in Information/Intelligent Systems - Part 1: Foundations; Part 2: Applications*. Physica-Verlag, Heidelberg, Germany, 1999.

-
- [165] S. Zadrozny and J. Kacprzyk. Computing with words for text processing: An approach to the text categorization. *Information Sciences*, 176(4):415–437, 2006.
- [166] Q. Zhang, J.C.H. Chen, and P.P. Chong. Decision consolidation: Criteria weight determination using multiple preference formats. *Decision Support Systems*, 38(2):247–258, 2004.
- [167] H.J. Zimmermann. *Fuzzy sets: Theory and its applications*. Kluwer Academic, 1996.

