



UNIVERSIDAD
PABLO DE OLAVIDE
SEVILLA



REVISTA DE MÉTODOS CUANTITATIVOS PARA
LA ECONOMÍA Y LA EMPRESA (14). Páginas 36–53.
Diciembre de 2012. ISSN: 1886-516X. D.L: SE-2927-06.
URL: <http://www.upo.es/RevMetCuant/art.php?id=61>

Un problema de consenso para problemas de toma de decisiones multicriterio en grupo mediante relaciones de preferencia intervalares difusas lingüísticas

TAPIA GARCÍA, JUAN MIGUEL

Departamento de Métodos Cuantitativos para la Economía y la Empresa
Universidad de Granada

Correo electrónico: jmtaga@ugr.es

DEL MORAL ÁVILA, MARÍA JOSÉ

Departamento de Estadística e Investigación Operativa
Universidad de Granada

Correo electrónico: delmoral@ugr.es

TAPIA GARCÍA, CRISTÓBAL

Departamento de Electrónica, Automática e Informática Industrial
Universidad Politécnica de Madrid

Correo electrónico: cristobal.tapia@upm.es

MARTÍNEZ, MARÍA DE LOS ÁNGELES

Grupo SECABA

Correo electrónico: mnglsmartinez@gmail.com

AMOR PULIDO, RAÚL

Departamento de Métodos Cuantitativos para la Economía y la Empresa
Universidad de Granada

Correo electrónico: ramor@ugr.es

RESUMEN

En el contexto de toma de decisiones multicriterio y bajo ciertas circunstancias, puede ocurrir que no se pueda expresar una cierta valoración mediante una única etiqueta lingüística, ya que puede haber duda en esa valoración. En este trabajo, presentamos un modelo de consenso para problemas de toma de decisiones en grupo con relaciones de preferencia intervalares lingüísticas. Este modelo está basado en dos criterios de consenso, una medida de consenso y una de proximidad, y en el concepto de coincidencia entre preferencias. Calcularemos ambos criterios en los tres niveles de representación de una relación de preferencia y diseñaremos un mecanismo de realimentación automático para guiar a los expertos en el proceso para alcanzar el consenso.

Palabras clave: toma de decisiones multicriterio; consenso; relaciones de preferencia intervalares lingüísticas.

Clasificación JEL: C19; C63; C69.

MSC2010: 62C86; 90B50.

Artículo recibido el 12 de octubre de 2011 y aceptado el 17 de septiembre de 2012.

A Consensus Model for Group Multicriteria Decision Making Problems with Interval Fuzzy Preference Relations

ABSTRACT

In some circumstances a decision maker, expert, in a group decision making problem cannot express his/her preferences with a unique linguistic fuzzy preference because he/she is dubious into some preferences. In this paper, we present a consensus model for group decision making problems with interval fuzzy preference relations. This model is based on two consensus criteria, a consensus measure and a proximity measure, and on the concept of coincidence among preferences. We compute both consensus criteria in the three representation levels of a preference relation and design an automatic feedback mechanism to guide experts in the consensus reaching process.

Keywords: group multicriteria decision making; consensus; linguistic interval fuzzy preference relations.

JEL classification: C19; C63; C69.

MSC2010: 62C86; 90B50.



1. INTRODUCCIÓN

Un problema de toma de decisiones multicriterio en grupo (GMCDM) podría definirse como un problema de decisión con varias alternativas y un conjunto de elementos –expertos– que toman decisiones intentando alcanzar una solución común, expresando sus opiniones –preferencias– sobre ellas según unos criterios. Este tipo de problemas pueden modelizarse de forma similar a los problemas de toma de decisión en grupo GDM (Herrera y Herrera-Viedma, 2000). En los problemas de toma de decisiones, las relaciones de preferencia son la representación más frecuente de las preferencias de los expertos ya que son útiles para expresar la información sobre las alternativas. Así, un problema GMCDM puede modelarse mediante una matriz para cada experto cuya entrada en la fila i y columna j refleja la preferencia del experto sobre la alternativa i según el criterio j (Bustince *et al.*, 2012).

Algunos problemas presentan aspectos cuantitativos que pueden ser valorados mediante valores numéricos precisos con la ayuda de teoría de lógica difusa (Chiclana *et al.*, 1998; Herrera-Viedma *et al.*, 2007a; 2007b; Kacprzyk *et al.*, 1997). Sin embargo, algunos otros problemas también presentan aspectos cualitativos que son complejos de valorar mediante estos valores. En estas ocasiones la aproximación mediante lógica difusa lingüística (Herrera *et al.*, 2000; Herrera-Viedma *et al.*, 2005; Xu, 2005; 2007; Zadeh, 1975a; 1975b; 1975c) puede emplearse para obtener una mejor solución. Nuestro interés se centra en problemas GMCDM en los que los expertos expresan sus preferencias con el uso de términos lingüísticos en lugar de valores numéricos. Muchos de estos problemas usan variables lingüísticas clasificadas en un conjunto de términos lingüísticos –etiquetas– y que cada experto emplea para expresar su opinión sobre el conjunto de alternativas como una relación de preferencia difusa lingüística. Sin embargo, puede darse el caso de que un experto tenga un conocimiento vago sobre el grado de preferencia lingüística o etiqueta que quiere asignar a la alternativa i según el criterio j y no pueda asignarle una etiqueta exacta. En estos casos el uso de relaciones de preferencia difusas lingüísticas intervalares puede resultar útil. Hasta la fecha ninguna investigación se ha dedicado a presentar un modelo de consenso en problemas GMCDM bajo relaciones de preferencia difusas lingüísticas intervalares. Este trabajo se centra en la definición de un nuevo modelo de consenso para problemas GMCDM mediante relaciones de preferencia difusas lingüísticas intervalares basado en 2 criterios –medidas– de consenso. Hemos propuesto un modelo de consenso (Tapia *et al.*, 2012) bajo relaciones de preferencia difusas intervalares para problemas de toma de decisiones en grupo (GDM) y este trabajo extiende esta propuesta al caso de toma de decisiones multicriterio en contexto lingüístico.

Además, las medidas de consenso propuestas permiten controlar el estado de consenso a distintos niveles y con ello dan la posibilidad de generar recomendaciones personalizadas en el proceso de alcanzar el consenso. También, las medidas de consenso propuestas permiten automatizar el proceso de consenso y sustituir la tradicional figura del moderador.

Un método de resolución habitual para problemas GDM/GMCDM consiste en la aplicación de dos procesos diferentes (Alonso *et al.*, 2010; Herrera *et al.*, 1995, 1996; Pérez *et al.*, 2011): un *proceso de consenso* –claramente, en cualquier proceso de decisión, es preferible que los expertos alcancen un alto grado de acuerdo sobre la solución en el conjunto alternativas. Además, este proceso indica cómo obtener el máximo grado de consenso o acuerdo entre los expertos sobre las diferentes alternativas–; y un *proceso de selección* –este proceso consiste en la obtención de la solución del conjunto de alternativas a partir de las preferencias expresadas por los expertos–.

En general, en cualquier problema GMCDM, el grupo de expertos tendrá inicialmente preferencias diferentes y es necesario desarrollar un proceso para alcanzar el consenso. Así, un proceso de consenso puede ser visto como un proceso dinámico donde un moderador, vía cambio de información y argumentos racionales, intenta que los expertos actualicen sus preferencias. En cada paso, se mide el grado de consenso presente y la distancia al consenso ideal se mide. Este proceso se repite hasta que la distancia entre el consenso presente y el ideal se considere suficientemente pequeña. De forma tradicional, la idea de consenso ideal significa un acuerdo completo y unánime de las preferencias de todos los expertos. Este tipo de consenso es ideal y muy difícil de alcanzar. Esto ha llevado a la definición y uso de un nuevo concepto denominado grado de consenso “soft” (Kacprzyk 1987; Kacprzyk y Fedrizzi, 1986; 1988), que permite usar el grado de consenso de una forma más flexible. El consenso soft permite medir la proximidad entre las opiniones de los expertos basándose en el concepto de coincidencia (Cabrerizo *et al.*, 2010b; 2010c; Herrera *et al.*, 1997b).

El objetivo de este artículo es presentar un modelo de consenso basado en la coincidencia soft entre preferencias para problemas GMCDM bajo relaciones de preferencia difusas intervalares lingüísticas. Como en el caso de problemas GDM (Herrera *et al.*, 1996; 1997a; 1997b), este nuevo modelo está basado en dos criterios de consenso para guiar el proceso para alcanzar el consenso: una *medida de consenso* y una *medida de proximidad*. La medida de consenso evaluará el acuerdo entre los expertos y se empleará para guiar el proceso de consenso hasta que la solución final se alcance. La medida de proximidad evaluará el acuerdo entre opiniones individuales de cada experto y la opinión colectiva del grupo y se utilizará para guiar la discusión del grupo en el proceso de consenso. Estas medidas se aplican sobre los tres niveles de representación de una relación de preferencia difusa intervalar lingüística: nivel de pares, nivel de alternativa y nivel de relación. También, presentamos un mecanismo automático de *feedback* para guiar a los expertos en el proceso de consenso y facilitar/sustituir la acción del moderador.

Este trabajo se desarrolla de la siguiente forma. El problema GMCDM basado en relaciones de preferencia difusas intervalares lingüísticas se describe en la Sección 2. La

Sección 3 presenta el nuevo modelo de consenso. Un ejemplo de aplicación práctica se da en la Sección 4. Finalmente, en la Sección 5, presentamos nuestras conclusiones.

2. EL PROBLEMA GMCDM CON ESTRUCTURAS DE PREFERENCIA DIFUSAS INTERVALARES LINGÜÍSTICAS

En esta sección describimos brevemente el problema GMCDM basado en relaciones de preferencia difusas intervalares lingüísticas y el proceso de resolución empleado para obtener la solución del conjunto de alternativas.

2.1 El contexto lingüístico

Como ya se ha mencionado, puede que un experto no pueda estimar su grado de preferencia con un valor numérico exacto. En este caso otra posibilidad es emplear etiquetas lingüísticas. De acuerdo con Herrera *et al.* (1997), como las afirmaciones lingüísticas son meras aproximaciones dadas por los expertos, podemos considerar funciones de pertenencia trapezoidales lineales como herramientas lo suficientemente buenas para capturar la vaguedad de estas afirmaciones lingüísticas. Una función de pertenencia trapezoidal lineal se puede representar mediante un conjunto de 4-uplas $(a_i, b_i, \alpha_i, \beta_i)$. En la Figura 1, los dos primeros parámetros de cada 4-upla indican el intervalo en el cual el valor de pertenencia es 1.0 –máximo– y los dos últimos indican las amplitudes a izquierda y derecha de la distribución respectivamente.

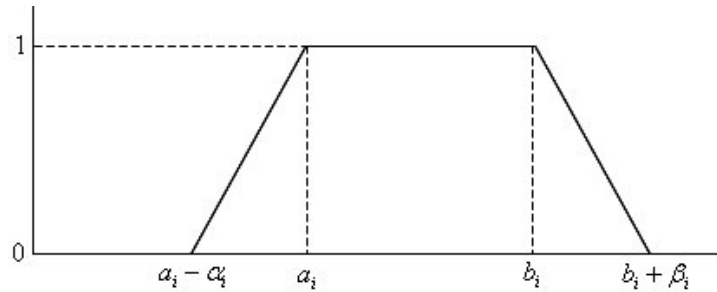


Figura 1. Representación de una función de pertenencia trapezoidal lineal.

Consideraremos un conjunto totalmente ordenado y finito $S = \{s_0, \dots, s_T\}$ con cardinalidad impar, donde cada etiqueta s_i representa un posible valor para una variable lingüística real.

Ejemplo 1. Consideraremos para el resto de los ejemplos el siguiente conjunto de etiquetas lingüísticas con su respectiva semántica asociada:

$s_8 = C$	<i>Completamente seguro/Pleno</i>	$(1.00, 1.00, 0.00, 0.00)$
$s_7 = MS$	<i>Muy seguro</i>	$(0.98, 0.99, 0.05, 0.01)$
$s_6 = BS$	<i>Bastante seguro</i>	$(0.78, 0.92, 0.06, 0.05)$
$s_5 = AS$	<i>Algo seguro</i>	$(0.63, 0.80, 0.05, 0.06)$
$s_4 = P$	<i>Puede</i>	$(0.41, 0.58, 0.09, 0.07)$

$s_3 = PS$	<i>Poco seguro</i>	$(0.22, 0.36, 0.05, 0.06)$
$s_2 = BPS$	<i>Bastante Poco seguro</i>	$(0.10, 0.18, 0.06, 0.05)$
$s_1 = MPS$	<i>Muy poco seguro</i>	$(0.01, 0.02, 0.01, 0.05)$
$s_0 = I$	<i>Nada seguro/Imposible</i>	$(0.00, 0.00, 0.00, 0.00)$

donde el término seguro puede alternarse por expresiones del tipo: alto-bajo, fiable... en donde proceda.

2.2 El problema GMCDM

Sea $X = \{x_1, \dots, x_n\}$ ($n \geq 2$) un conjunto finito de alternativas que serán evaluadas por un conjunto finito de expertos, $E = \{e^1, \dots, e^m\}$ ($m \geq 2$) según un conjunto finito de criterios $Y = \{y_1, \dots, y_t\}$ ($t \geq 2$). El proceso GMCDM consiste en encontrar la mejor alternativa de acuerdo a las preferencias de los expertos $\{P^1, \dots, P^m\}$.

En este trabajo asumimos que las preferencias de cada uno de los m expertos sobre X se describen mediante una relación de preferencias difusas intervalares para cada criterio, que notamos L^k , y que podemos representar mediante

$L^k \subset X \times Y$, cuya función de pertenencia L_k será:

$$L_k : X \times Y \rightarrow S, \text{ con } k = 1, \dots, m.$$

Para cada experto E^k , $L_k(x_i, y_j) = [p_{ij}^{k-}, p_{ij}^{k+}]$ denota el grado de preferencia difusa intervalar lingüístico de la alternativa x_i frente al criterio y_j , utilizando para ello el conjunto de etiquetas lingüísticas anteriormente mencionado con $s_0 \leq p_{ij}^{k-} \leq p_{ij}^{k+} \leq s_T$. Además, motivado por el orden inducido por el conjunto S , definimos una función natural:

$$s : S \rightarrow \mathbb{N}$$

mediante la asignación $s(p_{ij}^z) = a$ si $p_{ij}^z = s_a$. Por ejemplo, en el caso de que el tercer experto, E^3 , de un grupo exprese su preferencia sobre el grado de cumplimiento del segundo criterio para la primera alternativa –empleando el conjunto de etiquetas del ejemplo 1– entre Muy seguro y Pleno, se tendrá que, $L_3(x_1, y_2) = [p_{12}^{3-}, p_{12}^{3+}] = [MS, C]$. Además, si $p_{12}^{3-} = s_7 = MS$, se tiene que $s(p_{12}^{3-}) = 7$.

2.3 El operador LOWA

Se define en este apartado el operador de media ponderada ordenado lingüístico –linguistic ordered weighted averaging (LOWA) – que se aplica en la fase de agregación. Este operador se basa en el operador OWA definido por Yager (1988). Sea $\{a_1, \dots, a_m\}$ un conjunto de etiquetas que han de ser agregadas. El operador LOWA, Φ , puede definirse en la forma:

$$\Phi(a_1, \dots, a_m) = W \cdot B^T = C^m \{w_k, b_k, k = 1, \dots, m\} = w_1 \odot b_1 \oplus (1 - w_1) \odot C^{m-1} \{\beta_h, b_h, h = 2, \dots, m\}$$

donde $W = [w_1, \dots, w_m]$ es un vector de ponderaciones tal que:

$$w_i \in [0,1]; \sum_i w_i = 1; \beta_h = \frac{w_h}{\sum_2^m w_k}; h = 2, \dots, m$$

y B es vector ordenado de etiquetas asociado. Cada elemento $b_i \in B$ es la i -ésima mayor etiqueta en la colección $\{a_1, \dots, a_m\}$. C^m es el operador combinación convexa de m etiquetas y, si $m = 2$, entonces se define:

$$C^2 \{w_i, b_i, k = 1, 2\} = w_1 \odot s_j \oplus (1 - w_1) \odot s_i = s_k; s_i, s_j \in S; j \geq i,$$

donde $k = \text{mínimo}\{T, i + \text{round}(w_1 \cdot (j - i))\}$, siendo round la operación usual de redondeo y $b_1 = s_j, b_2 = s_i$.

Si $w_j = 1$ y $w_i = 0$, con $i \neq j$, entonces la combinación convexa se define:

$$C^m \{w_i, b_i, k = 1, \dots, m\} = b_j$$

Yager (1988) sugirió un método para el cálculo de las ponderaciones del operador de agregación OWA empleando cuantificadores lingüísticos, que en el caso de un cuantificador proporcional no-decreciente Q , viene dado por la expresión:

$$w_i = Q\left(\frac{i}{n}\right) - Q\left(\frac{i-1}{n}\right), \quad i = 1, \dots, n,$$

donde la función Q puede definirse mediante:

$$Q(r) = \begin{cases} 0 & \text{si } r < a, \\ \frac{r-a}{b-a} & \text{si } a \leq r \leq b, \\ 1 & \text{si } r > b, \end{cases}$$

con $a, b, r \in [0, 1]$. Cuando un cuantificador difuso lingüístico Q se emplea para calcular las ponderaciones del operador LOWA Φ , lo simbolizaremos por Φ_Q .

2.4 Proceso de resolución del problema GMCDM

De forma habitual, el proceso de resolución de un problema GMCDM consiste en la obtención de un conjunto de alternativas que son solución, partiendo de las preferencias dadas por los expertos. Este proceso de resolución consta de dos fases: fase de consenso y fase de selección. La fase de selección es la última de ellas y nos permite obtener el conjunto de soluciones. Se compone, a su vez de dos procedimientos (Herrera y Herrera-Viedma, 1997): (i) *agregación* y (ii) *explotación*. Estos dos procedimientos los veremos a continuación para centrarnos en el modelo de consenso que es el objetivo de la siguiente sección.

i) Procedimiento de agregación

Este procedimiento define una relación de preferencia difusa intervalar lingüística colectiva obtenida mediante la agregación de todas las relaciones de preferencias difusas intervalares lingüísticas. Esta relación colectiva, denominada U , indica la preferencia global entre cada par

ordenado de alternativas de acuerdo a la opinión mayoritaria de los expertos. En nuestro caso, de entre las agregaciones posibles utilizaremos la siguiente expresión:

$$U = (U_{ij}) \quad \text{para } i = 1, \dots, n \text{ y } j = 1, \dots, t \text{ con}$$

$$U_{ij} = U[p_{ij}^-, p_{ij}^+] = [\Phi_-(p_{ij}^{k-}), \Phi_+(p_{ij}^{k+})] = [\min_k(p_{ij}^{k-}), \max_k(p_{ij}^{k+})]$$

$$\text{con } w_- = \{0, \dots, 0, 1\} \text{ en } \Phi_- \text{ y } w_+ = \{1, 0, \dots, 0\} \text{ en } \Phi_+ \text{ y } k = 1, \dots, m$$

Ejemplo 2. Dos expertos expresan sus opiniones sobre el cumplimiento en dos empresas de tres criterios: rentabilidad, administración de materias primas y gestión medioambiental de los residuos. Para ello emplean el conjunto de etiquetas lingüísticas del Ejemplo 1. Sus preferencias vienen dadas mediante las siguientes relaciones de preferencia difusas intervalares lingüísticas:

$$e^1 = \begin{pmatrix} [BPS, P] & [AS, C] & [AS, MS] \\ [PS, P] & [MS, MS] & [BS, C] \end{pmatrix} \quad e^2 = \begin{pmatrix} [PS, AS] & [MS, MS] & [BS, MS] \\ [AS, BS] & [BS, MS] & [BS, BS] \end{pmatrix}$$

Así, por ejemplo, el experto 1 valora en la primera empresa un cumplimiento entre Poco Seguro y Puede para la rentabilidad, de Algo Seguro a Completamente Seguro en la administración de materias primas y de Algo Seguro a Muy Seguro en la gestión medioambiental de los residuos, etc. Empleando la herramienta de agregación previamente definida se obtiene la siguiente relación de preferencia colectiva U :

$$U = \begin{pmatrix} [BPS, AS] & [AS, C] & [AS, MS] \\ [PS, BS] & [BS, MS] & [BS, C] \end{pmatrix}$$

ii) Procedimiento de explotación

En este procedimiento se transforma la información colectiva sobre las alternativas en una clasificación global de ellas, pudiendo así elegir el conjunto solución. Para ello, una elección habitual es una función de elección de alternativas que se aplica sobre la relación de preferencia colectiva para obtener esta clasificación (Herrera y Herrera-Viedma, 2000). Por ejemplo, podemos definir una función de elección empleando el concepto de dominancia (Herrera y Herrera-Viedma, 2000). Así, para cada alternativa x_i podemos calcular su grado de dominancia en la relación de preferencia colectiva px_i como:

$$px_i = \sum_{j=1}^n (s(p_{ij}^-) + s(p_{ij}^+))$$

De esta forma, obtenemos una clasificación de las alternativas:

$$\text{Si } px_i > px_j \text{ entonces } x_i \text{ es preferido a } x_j.$$

Ejemplo 3. De la relación de preferencia difusa intervalar colectiva obtenida en el Ejemplo 2, se pueden caracterizar las dos posibles alternativas con los siguientes grados de dominancia, empleando las etiquetas que aparecen en U

$$U = \begin{pmatrix} [BPS = s_2, AS = s_5] & [AS = s_5, C = s_8] & [AS = s_5, MS = s_7] \\ [PS = s_3, BS = s_6] & [BS = s_6, MS = s_7] & [BS = s_6, C = s_8] \end{pmatrix},$$

de donde la expresión anterior se concreta en:

$$px_1 = (2 + 5) + (5 + 8) + (5 + 7) = 32;$$

$$px_2 = (3 + 6) + (6 + 7) + (6 + 8) = 36.$$

Así por ejemplo en U, para la primera empresa y en el primer criterio, aparecen la etiquetas $BPS = s_2$ y $AS = s_5$; entonces, mediante la función s , los dos primeros sumandos de px_1 son 2 y 5...

De esta forma se obtiene una clasificación de las empresas de mayor a menor preferencia:

$$x_2 > x_1$$

Por tanto, la segunda empresa es la mejor valorada por los expertos según los criterios empleados.

Evidentemente, no es igual que el conjunto de soluciones tenga un alto reconocimiento por parte de los expertos o que, por el contrario, sea discutido por ellos. Por tanto, es muy importante tener en cuenta el grado de consenso que lleva aparejado el conjunto de soluciones referido. Como se mencionó anteriormente, no existe un modelo de consenso para tratar problemas GMCDM bajo relaciones de preferencia difusas intervalares lingüísticas. En la siguiente sección presentamos un proceso de consenso para problemas GMCDM mediante relaciones de preferencia difusas intervalares en un contexto lingüístico.

3. MODELO DE CONSENSO

En esta sección presentamos un modelo consenso definido para problemas GMCDM asumiendo que los expertos expresan sus preferencias por medio de relaciones de preferencia difusas intervalares lingüísticas. Este modelo presenta las siguientes características principales:

- a) Se basa en dos criterios de consenso soft: una medida de consenso y una medida de proximidad.
- b) Ambos criterios de consenso están definidos usando la coincidencia entre relaciones de preferencia difusas intervalares dadas por los expertos (Cabrerizo *et al.*, 2010a).
- c) Incorpora un mecanismo de feedback que genera recomendaciones para los expertos sobre cómo cambiar sus relaciones de preferencia durante el proceso de consenso.

Inicialmente podemos considerar que, en cualquier problema GMCDM no trivial, las preferencias de los expertos son diferentes. Así, este consenso puede ser visto como un proceso iterativo, lo que significa que el acuerdo se obtiene tras varios turnos de consultas. En cada turno calculamos los dos criterios para el consenso (Herrera *et al.*, 1997). El supervisor evalúa el nivel de acuerdo entre todos los expertos y guía el proceso de consenso, midiendo posteriormente la distancia entre las preferencias individuales de los expertos y la preferencia colectiva, y soporta la fase de discusión del proceso de consenso. Para ello, calculamos la coincidencia entre relaciones de preferencia difusas intervalares lingüísticas.

El principal problema es cómo encontrar un modo de hacer que las posiciones individuales converjan. Para esto, se fija inicialmente un nivel de consenso (A) requerido en cada situación. Cuando la medida del consenso alcanza este nivel, la sesión de toma de decisiones finaliza y se obtiene la solución aplicando el proceso de selección. Si este no es el caso, las opiniones de los expertos deben ser modificadas. Esto se hace en una sesión de discusión en grupo en la que un mecanismo de feedback se usa para apoyar el cambio de opinión de los expertos. Este mecanismo de feedback se define usando las medidas de proximidad (Herrera-Viedma *et al.*, 2002; 2005; 2007a; Mata *et al.*, 2009). Para evitar que la solución colectiva no converja tras varias sesiones de discusión, es posible fijar un número máximo de turnos. El esquema de este modelo de consenso para GMCDM se presenta en la Figura 2. En las siguientes subsecciones presentamos los componentes de este modelo de consenso en detalle, es decir, los criterios de consenso y el mecanismo de feedback.

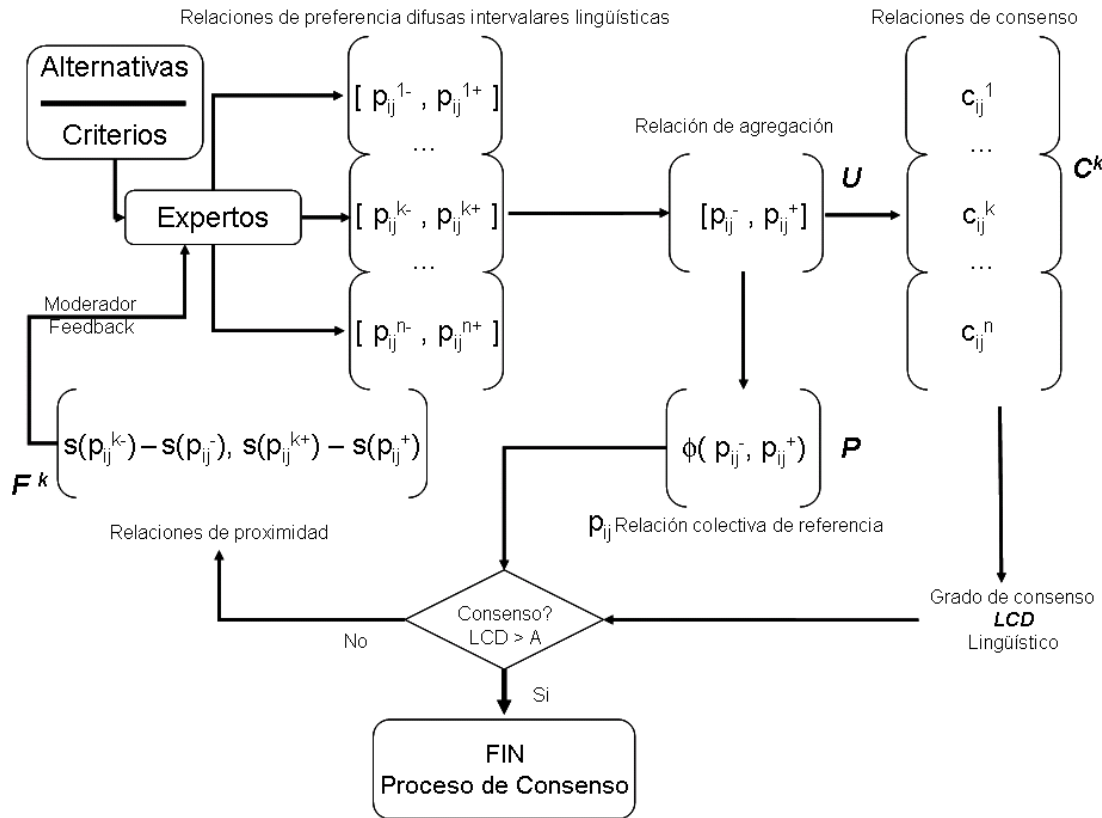


Figura 2. Modelo de consenso para GMCDM con relaciones de preferencias difusas intervalares lingüísticas.

3.1 Medidas de consenso y proximidad

Calculamos ambos indicadores del consenso mediante los pasos siguientes:

En primer lugar, calculamos las relaciones de consenso de cada experto e^k , llamadas C^k , con respecto a las relaciones de preferencia colectiva

$$C^k = (C_{ij}^k) \quad \text{con}$$

$$C_{ij}^k = \frac{\left(\left| s(p_{ij}^{k-}) - s(p_{ij}^-) \right| + \left| s(p_{ij}^{k+}) - s(p_{ij}^+) \right| \right)}{T} \quad \text{para } i = 1, \dots, n \text{ y } j = 1, \dots, t.$$

En esta relación de consenso, cada valor C_{ij}^k representa el grado de acuerdo del experto e^k con el grupo de expertos sobre la preferencia p_{ij} . Entonces, definimos el *grado de consenso lingüístico sobre una preferencia* p_{ij} :

$$LCD_{ij} = 1 - \sum_{k=1}^m \frac{C_{ij}^k}{m}$$

Tendremos un consenso total en la preferencia p_{ij} si $LCD_{ij} = 1$.

Definimos el *grado de consenso lingüístico en la alternativa* x_i :

$$LCD_i = 1 - \sum_{j=1}^t \sum_{k=1}^m \frac{C_{ij}^k}{t \cdot m}$$

Tendremos un consenso total en la x_i si $LCD_i = 1$.

Definimos el *grado de consenso lingüístico global* LCD :

$$LCD = 1 - \sum_{i=1}^n \sum_{j=1}^t \sum_{k=1}^m \frac{C_{ij}^k}{t \cdot n \cdot m}$$

En este caso, $0 \leq LCD \leq 1$ o, equivalentemente, $0\% \leq LCD \leq 100\%$. Tendremos un consenso total si $LCD = 1$ o $LCD = 100\%$.

Ejemplo 4. De la relación de preferencia colectiva lingüística obtenida en el Ejemplo 2, se obtienen las siguientes relaciones de consenso:

$$C^1 = \begin{pmatrix} 0.125 & 0.000 & 0.000 \\ 0.250 & 0.125 & 0.000 \end{pmatrix} \text{ y } C^2 = \begin{pmatrix} 0.125 & 0.375 & 0.125 \\ 0.250 & 0.000 & 0.250 \end{pmatrix}$$

Aplicando las anteriores expresiones, el grado de consenso lingüístico global es $LCD = 0.8645$ o $LCD = 86.45\%$ y, por ejemplo, el grado de consenso lingüístico en la alternativa x_1 es $LCD_1 = 0.875$ o $LCD_1 = 87.5\%$, y el grado de consenso lingüístico en la preferencia p_{13} es $LCD_{13} = 0.938$ o $LCD_{13} = 93.8\%$.

Ahora, continuamos el proceso con el cálculo de las medidas de proximidad. Primero calculamos la relaciones de proximidad de los expertos, que denominaremos F^k , con respecto a la relación de preferencia colectiva U :

$$F^k = (F_{ij}^k), \text{ con}$$

$$F_{ij}^k = (s(p_{ij}^{k-}) - s(p_{ij}), s(p_{ij}^{k+}) - s(p_{ij})) = (f_{ij}^{k-}, f_{ij}^{k+}) \quad \text{para } i = 1, \dots, n \text{ y } j = 1, \dots, t$$

$$\text{y } p_{ij} = \Phi_Q(p_{ij}^-, p_{ij}^+), \text{ con } s(p_{ij}) = n \text{ si } p_{ij} = s_n.$$

Entonces definimos la *medida de proximidad del experto* e^k sobre una preferencia p_{ij} :

$$PM_{ij}^k = \frac{(|f_{ij}^{k-}| + |f_{ij}^{k+}|)}{2T}$$

y la medida de proximidad del experto e^k en una alternativa x_i :

$$PM_i^k = \sum_{j=1}^t \frac{PM_{ij}^k}{t}$$

Finalmente definimos la medida de proximidad global del experto e^k :

$$PM^k = \sum_{i=1}^n \frac{PM_i^k}{n}$$

Ejemplo 5. Empleando los datos del Ejemplo 2, se obtienen las siguientes relaciones de proximidad para los dos expertos e^1 y e^2 , respectivamente, mediante la comparación de cada matriz de preferencias (e^i) con la matriz Φ_Q —construida mediante la agregación de los elementos de U —:

A partir de

$$U = \begin{pmatrix} [BPS = s_2, AS = s_5] & [AS = s_5, C = s_8] & [AS = s_5, MS = s_7] \\ [PS = s_3, BS = s_6] & [BS = s_6, MS = s_7] & [BS = s_6, C = s_8] \end{pmatrix}$$

mediante agregación —empleando un LOWA— resulta:

$$\Phi_Q(p_{ij}^-, p_{ij}^+) = \begin{pmatrix} s_4 & s_7 & s_6 \\ s_5 & s_7 & s_7 \end{pmatrix}$$

De la comparación término a término de esta nueva matriz y de la matriz de preferencias del experto primero

$$e^1 = \begin{pmatrix} [BPS = s_2, P = s_4] & [AS = s_5, C = s_8] & [AS = s_5, MS = s_7] \\ [PS = s_3, P = s_4] & [MS = s_7, MS = s_7] & [BS = s_6, C = s_8] \end{pmatrix},$$

se obtiene la matriz de relaciones de proximidad del primer experto con respecto a la relación de preferencia colectiva:

$$F^1 = \begin{pmatrix} (-2, +0) & (-2, +1) & (-1, +1) \\ (-2, -1) & (+0, +0) & (-1, +1) \end{pmatrix}$$

Así, por ejemplo, $(-2, +0)$ surge de la comparación de las posiciones 1,1 de ambas matrices: s_4 y $[BPS = s_2, P = s_4]$. De forma análoga se obtiene para el segundo experto:

$$F^2 = \begin{pmatrix} (-1, +1) & (+0, +0) & (-1, +1) \\ (+0, +1) & (-1, +0) & (-1, +1) \end{pmatrix}$$

También empleando las fórmulas propuestas anteriormente, se obtienen las medidas de proximidad de cada alternativa para ambos expertos:

$$PM_I^1 = 0.146 \quad PM_I^2 = 0.063$$

$$PM^l_2 = 0.083 \quad PM^2_2 = 0.083$$

y para el conjunto de preferencias:

$$PM^l = 0.115 \quad PM^2 = 0.073$$

3.2 Proceso de feedback/moderador

Podemos aplicar un mecanismo de feedback para guiar el cambio de las opiniones de los expertos mediante el uso de las matrices de proximidad F^k (Herrera-Viedma *et al.*, 2002; 2005; 2007b; Mata, 2009). Este mecanismo puede ayudar al moderador en sus tareas o incluso sustituir las acciones de éste en el proceso de consenso. De esta forma, el proceso de feedback ayuda a los expertos a cambiar sus preferencias hacia la obtención de un grado de acuerdo apropiado. Como se ha mencionado, el principal problema para el mecanismo de feedback es cómo encontrar un modo para que las posiciones individuales converjan apoyando a los expertos en la búsqueda de la solución (Herrera-Viedma *et al.*, 2002). Usualmente el proceso de feedback se lleva a cabo en dos fases: *fase de identificación* y *fase de recomendación*.

Fase de identificación. Previamente es necesario comparar el grado de consenso lingüístico global LCD y el valor límite fijado previamente A . Si $LCD > A$ o $LCD = A$, el proceso de consenso se detiene. Por otro lado, si $LCD < A$, tendrá lugar un nuevo turno de discusión para el consenso. Si el acuerdo entre los expertos es bajo, existirán muchas preferencias de los expertos en desacuerdo. En esta situación, para mejorar el acuerdo, el número de cambios en las preferencias debería ser alto. Sin embargo, si el acuerdo es elevado, la mayoría de las preferencias estarán próximas y solo un pequeño número de preferencias estarán en desacuerdo; parece entonces razonable cambiar solo estas referencias en particular. El procedimiento sugiere modificar las valoraciones de preferencia sobre todos los pares de alternativas dónde el acuerdo no es suficientemente alto. Para identificar el conjunto de las preferencias que deberían modificarse se propone:

- a) Identificar los pares de alternativas con un grado de consenso menor que el valor límite A , $LCD_{ij} < A$.
- b) Identificar, en segundo lugar, a los expertos que deberían modificar los pares de alternativas identificadas en a). Para esto, usaremos las medidas de proximidad PM^k y PM^k_i , y también fijaremos un valor límite B , de tal forma que los expertos que son requeridos para modificar sus preferencias en dichos pares de alternativas son aquellos donde $PM^k > B$.

Fase de recomendación. En esta fase se recomienda a los expertos que cambien sus preferencias de acuerdo a ciertas reglas. Así, aplicaremos unas reglas de recomendación que informarán a los expertos de la dirección apropiada de los cambios para mejorar el acuerdo. Encontramos la dirección de cambio que se debe aplicar sobre las preferencias p_{ij}^{k+} o p_{ij}^{k-} de cada experto sobre una preferencia concreta, aplicando las siguientes reglas:

- 1) Si $s(p_{ij}^{k-}) - s(p_{ij}) = f_{ij}^{k-} > 0$, entonces el experto e^k debería decrementar su valoración asociada al par de alternativas (x_i, y_j) .
- 2) Si $s(p_{ij}^{k+}) - s(p_{ij}) = f_{ij}^{k+} < 0$, entonces el experto e^k debería incrementar su valoración asociada al par de alternativas (x_i, y_j) .
- 3) Si $f_{ij}^{k-} < 0 < f_{ij}^{k+}$, entonces el experto e^k debería decrementar su valoración p_{ij}^{k+} y debería incrementar su valoración p_{ij}^{k-} asociadas al par de alternativas (x_i, y_j) .

4. EJEMPLO

Tres analistas valoran cuatro fondos de inversión bajo tres criterios para una entidad que desea invertir en ellos. Las opiniones de los expertos se reflejan en las siguientes relaciones de preferencia difusas intervalares lingüísticas:

$$e^1 = \begin{pmatrix} [PS, AS] & [P, BS] & [AS, AS] \\ [BS, C] & [BS, BS] & [MS, C] \\ [MS, MS] & [C, C] & [MS, C] \\ [AS, AS] & [AS, MS] & [BS, MS] \end{pmatrix} \quad e^2 = \begin{pmatrix} [P, P] & [AS, BS] & [BS, BS] \\ [BS, MS] & [AS, AS] & [P, BS] \\ [MS, MS] & [BS, C] & [BS, MS] \\ [AS, BS] & [BS, BS] & [BS, MS] \end{pmatrix}$$

$$e^3 = \begin{pmatrix} [MPS, BPS] & [BPS, P] & [PS, PS] \\ [PS, AS] & [P, P] & [BPS, PS] \\ [P, P] & [AS, BS] & [P, MS] \\ [BPS, BPS] & [AS, AS] & [PS, P] \end{pmatrix}$$

Aplicando agregación, se obtiene la matriz de las relaciones de preferencia difusas intervalares colectiva siguiente:

$$U = \begin{pmatrix} [MPS, AS] & [BPS, BS] & [PS, BS] \\ [PS, C] & [P, BS] & [BPS, C] \\ [P, MS] & [AS, C] & [P, C] \\ [BPS, AS] & [AS, MS] & [PS, MS] \end{pmatrix}$$

Se calculan las relaciones de consenso de cada experto: C^1 , C^2 y C^3 , obteniendo los grados de consenso sobre las preferencias $[p_{ij}]$ como:

$$\begin{pmatrix} 0.6250 & 0.7083 & 0.6250 \\ 0.5833 & 0.7500 & 0.4167 \\ 0.6250 & 0.7500 & 0.7083 \\ 0.5417 & 0.8333 & 0.6250 \end{pmatrix}$$

y el grado de consenso global es $LCD = 0.6493$ o $LCD = 64.93\%$. Supuesto fijado un límite mínimo para el consenso de $3/4 = 0.75$ o 75% , se tiene en este caso que, como $LCD < 0.75$, se debe seguir con el turno de consultas a los expertos.

Para orientar a los expertos en sus reflexiones, se calcula F^k para cada experto:

$$F^1 = \begin{pmatrix} (+0, +2) & (+0, +2) & (+0, +0) \\ (+0, +2) & (+1, +1) & (+2, +3) \\ (+1, +1) & (+1, +1) & (+1, +2) \\ (+1, +1) & (-1, +1) & (+1, +2) \end{pmatrix} \quad F^2 = \begin{pmatrix} (+1, +1) & (+1, +2) & (+1, +1) \\ (+0, +1) & (+0, +0) & (-1, +1) \\ (+1, +1) & (-1, +1) & (+0, +1) \\ (+1, +2) & (+0, +0) & (+1, +2) \end{pmatrix}$$

$$F^3 = \begin{pmatrix} (-2, -1) & (-2, +0) & (-2, -2) \\ (-3, -1) & (-1, -1) & (-3, -2) \\ (-2, -2) & (-2, -1) & (-2, +1) \\ (-2, -2) & (-1, -1) & (-2, -1) \end{pmatrix}$$

Las medidas de proximidad para los expertos son:

$$\begin{array}{lll} PM^1_1 = 0.083 & PM^2_1 = 0.146 & PM^3_1 = 0.188 \\ PM^1_2 = 0.188 & PM^2_2 = 0.063 & PM^3_2 = 0.229 \\ PM^1_3 = 0.146 & PM^2_3 = 0.104 & PM^3_3 = 0.208 \\ PM^1_4 = 0.146 & PM^2_4 = 0.125 & PM^3_4 = 0.188 \end{array}$$

y

$$PM^1 = 0.141 \quad PM^2 = 0.109 \quad PM^3 = 0.203$$

Aplicando el mecanismo de feedback se tiene:

- Si se observa la matriz de los grados de consenso sobre las preferencias $[p_{ij}]$ todas aquellas preferencias en las que el valor no sea 0.75 o superior son candidatas a ser mejoradas, por ejemplo: p_{11} , p_{12} , p_{13} , p_{21} , p_{23} , ...
- Si se fija un valor límite experimental de 0.15 para identificar aquellos expertos que preferentemente deberían cambiar sus valoraciones, se encuentra que el experto 3 es aquel que más se aleja de la solución colectiva -0.203- (aunque el experto 1 está próximo) y sobre quién mayor incidencia se debería tener en aquellas preferencias anteriormente seleccionadas para que valore estas, con esta información.
- Por ejemplo, una recomendación podría ser del tipo: el experto 3 podría incrementar su preferencia en la alternativa p_{11} .

Tras varios turnos, las preferencias de los expertos son:

$$e^1 = \begin{pmatrix} [P, AS] & [AS, AS] & [AS, AS] \\ [BS, MS] & [AS, BS] & [BS, MS] \\ [MS, MS] & [MS, C] & [MS, MS] \\ [P, AS] & [AS, BS] & [BS, BS] \end{pmatrix} \quad e^2 = \begin{pmatrix} [P, P] & [AS, AS] & [AS, BS] \\ [BS, MS] & [AS, AS] & [P, AS] \\ [BS, MS] & [BS, MS] & [BS, MS] \\ [AS, AS] & [BS, BS] & [BS, BS] \end{pmatrix}$$

$$e^3 = \begin{pmatrix} [PS, P] & [PS, P] & [P, AS] \\ [AS, AS] & [P, AS] & [AS, AS] \\ [P, BS] & [AS, BS] & [AS, BS] \\ [P, P] & [AS, AS] & [P, P] \end{pmatrix}$$

dando lugar a la siguiente relación colectiva:

$$U = \begin{pmatrix} [PS, AS] & [PS, AS] & [P, BS] \\ [AS, MS] & [P, BS] & [P, MS] \\ [P, MS] & [AS, C] & [AS, MS] \\ [P, AS] & [AS, BS] & [P, BS] \end{pmatrix}$$

Las relaciones de consenso para cada experto son:

$$C^1 = \begin{pmatrix} 1 & 2 & 2 \\ 1 & 1 & 2 \\ 3 & 2 & 2 \\ 0 & 0 & 2 \end{pmatrix} \quad C^2 = \begin{pmatrix} 2 & 2 & 1 \\ 1 & 2 & 2 \\ 2 & 2 & 1 \\ 1 & 1 & 2 \end{pmatrix} \quad C^3 = \begin{pmatrix} 1 & 1 & 1 \\ 2 & 1 & 3 \\ 1 & 2 & 1 \\ 1 & 1 & 2 \end{pmatrix}$$

En este caso, se obtiene un grado de consenso global $LCD = 0.8125$ o $LCD = 81.25\%$, que supera el mínimo fijado para finalizar el proceso. A partir de la matriz colectiva final de preferencias U se obtienen los siguientes grados de dominancia:

$$px_1 = 26 \quad px_2 = 33 \quad px_3 = 36 \quad px_4 = 30$$

Por tanto, los fondos de inversión pueden ordenarse de mayor a menor valoración como:

$$x_3 > x_2 > x_4 > x_1$$

De esta forma, el fondo de inversión recomendado por los analistas sería el tercero, seguido del segundo, del cuarto y finalmente del primero y ello con un consenso de más del 81%.

5. CONCLUSIONES

En este trabajo hemos presentado un nuevo modelo de consenso para problemas de toma de decisiones multicriterio en grupo mediante relaciones de preferencia difusas intervalares lingüísticas. Este modelo de consenso está basado en dos criterios de consenso, una medida de consenso y una medida de proximidad. Las medidas de consenso propuestas permiten controlar el estado de consenso a distintos niveles y con ello dan la posibilidad de generar recomendaciones personalizadas en el proceso de alcanzar el consenso y, también, los criterios de consenso propuestos permiten automatizar el proceso de consenso y sustituir la tradicional figura del moderador. Finalmente, presentamos ejemplos para mostrar su posible aplicación.

AGRADECIMIENTOS

Los autores desean dar las gracias a los editores y a los evaluadores anónimos, por sus comentarios y sugerencias, que han permitido la publicación del presente trabajo. Este estudio ha sido desarrollado con la financiación del proyecto MTM2009-08886.

REFERENCIAS

- Alonso, S.; Herrera-Viedma, E.; Chiclana, F.; Herrera, F. (2010). A Web Based Consensus Support System for Group Decision Making Problems and Incomplete Preferences. *Information Sciences*, 180:23 (2010), 4477–4495.
- Bustince, H.; Fernandez, J.; Sanz, J.; Galar, M.; Mesiar, R. (2012). Multicriteria decision making by means of interval-valued Choquet integrals. *Advances in Intelligent and Soft Computing*, 107, 269–278.
- Cabrerizo, F.J.; López-Gijón, J.; Ruíz-Rodríguez, A.A.; Herrera-Viedma, E. (2010a). A Model Based on Fuzzy Linguistic Information to Evaluate the Quality of Digital Libraries. *Int. J. Inform. Technol. Decis. Making* 9:3, 455–472.
- Cabrerizo, F.J.; Moreno, J.M.; Pérez, I.J.; Herrera-Viedma, E. (2010b). Analyzing consensus approaches in fuzzy group decision making: advantages and drawbacks. *Soft Computing* 14:5, 451–463.
- Cabrerizo, F.J.; Pérez, I.J.; Herrera-Viedma, E. (2010c). Managing the Consensus in Group Decision Making in an Unbalanced Fuzzy Linguistic Context with Incomplete Information. *Knowledge-Based Systems* 23:2, 169–181.
- Chiclana, F.; Herrera, F.; Herrera-Viedma, E. (1998). Integrating three representation models in fuzzy multipurpose decision making based on fuzzy preference relations. *Fuzzy Sets and Systems*, 97(1), 33–48.
- Herrera, F.; Herrera-Viedma, E. (1997). Aggregation operators for linguistic weighted information. *IEEE Trans. System Man Cybernet.* 27, 646–656.
- Herrera, F.; Herrera-Viedma, E. (2000). Linguistic decision analysis: steps for solving decision problems under linguistic information. *Fuzzy Sets and Systems* 115, 67–82.
- Herrera, F.; Herrera-Viedma, E.; Verdegay, J.L. (1995). A sequential selection process in group decision making with linguistic assessment. *Inform. Sci.* 85, 223–239.
- Herrera, F.; Herrera-Viedma, E.; Verdegay, J.L. (1996). A model of consensus in group decision making under linguistic assessments. *Fuzzy Sets and Systems* 78, 73–87.
- Herrera, F.; Herrera-Viedma, E.; Verdegay, J.L. (1997a). A rational consensus model in group decision making under linguistic assessment. *Fuzzy Sets and Systems* 88, 31–49.
- Herrera, F.; Herrera-Viedma, E.; Verdegay, J.L. (1997b). Linguistic Measures Based on Fuzzy Coincidence for Reaching Consensus in Group Decision Making. *International Journal of Approximate Reasoning* 16, 309–334.
- Herrera-Viedma, E.; Alonso, S.; Chiclana, F.; Herrera, F. (2007a). A consensus model for group decision making with incomplete fuzzy preference relations. *IEEE Trans. Fuzzy Syst.* 15, 863–877.

- Herrera-Viedma, E.; Chiclana, F.; Herrera, F.; Alonso, S. (2007b). Group Decision-Making Model with Incomplete Fuzzy Preference Relations Based on Additive Consistency. *IEEE Transactions on Systems, Man and Cybernetics, Part B, Cybernetics*, 37:1, 176–189.
- Herrera-Viedma, E.; Herrera, F.; Chiclana, F. (2002). A consensus model for multiperson decision making with different preference structures. *IEEE Transactions on Systems Man and Cybernetics Part A - Systems and Humans* 32 394–402.
- Herrera-Viedma, E.; Martínez, L.; Mata, F.; Chiclana, F. (2005). A consensus support system model for group decision-making problems with multi-granular linguistic preference relations. *IEEE Trans. Fuzzy Syst.* 13, 644–645.
- Jiang Y. (2007). An approach to group decision making based on interval fuzzy preference relations. *Journal of Systems Science and Systems Engineering*, 16:1, 113–120.
- Kacprzyk, J. (1987). On some fuzzy cores and “soft” consensus measures in group decision making, en *The Analysis of Fuzzy Information* (J. Bezdek Ed.). *CRC Press*, 119–130.
- Kacprzyk, J.; Fedrizzi, M. (1986). “Soft” consensus measure for monitoring real consensus reaching processes under fuzzy preferences. *Control Cybernet.* 15, 309–323.
- Kacprzyk, J.; Fedrizzi, M. (1988). A “soft” measure of consensus in the setting of partial (fuzzy) preferences. *European J. Oper. Res.* 34, 316–323.
- Mata, F.; Martínez, L.; Herrera-Viedma, E. (2009). An adaptive consensus support model for group decision making problems in a multi-granular fuzzy linguistic context. *IEEE Transactions on Fuzzy Systems*, 17(2), 279–290.
- Pérez, I.J.; Cabrerizo, F.J.; Herrera-Viedma, E. (2011). Group Decision Making Problems in a Linguistic and Dynamic Context. *Expert Systems with Applications* 38:3, 1675–1688.
- Xu, Z.S. (2005). An approach to group decision making based on incomplete linguistic preference relations. *International Journal of Information Technology and Decision Making* 4(1), 153–160.
- Xu, Z.S. (2007). A survey of preference relations. *International Journal of General Systems* 36, 179–203.
- Xu, Z.S. (2010). Interactive group decision making procedure based on uncertain multiplicative linguistic preference relations. *Journal of Systems Engineering and Electronics* 21, 408–415.
- Yager, R.R. (1988). On ordered weighted averaging aggregation operators in multicriteria decision making. *IEEE Trans. Systems Man Cybernet.* 18, 183–190.
- Zadeh, L.A. (1975a). The concept of a linguistic variable and its applications to approximate reasoning. Part I. *Information Sciences* 8(3), 199–249.
- Zadeh, L.A. (1975b). The concept of a linguistic variable and its applications to approximate reasoning. Part II. *Information Sciences* 8(4), 301–357.
- Zadeh, L.A. (1975c). The concept of a linguistic variable and its applications to approximate reasoning. Part III. *Information Sciences* 9(1), 43–80.