

Tecnologías y Lógica Fuzzy

Actas del Cuarto Congreso
de la Asociación Española de Tecnologías y Lógica Fuzzy
FUZZY'94

14, 15, y 16 de Septiembre de 1994
Institut d'Investigació en Intel·ligència Artificial
Blanes

Editores:

Francesc Esteva
Pere García

Publicado por:

Institut d'Investigació en Intel·ligència Artificial – CSIC.
Camí de Sta. Bàrbara, s/n. 17300 Blanes, España.

ISBN: 84-00-07440-8

Depósito legal: GI-1346-94

Indice

Prefacio	ix
----------------	----

Comités de Programa y de Organización	xi
---	----

CONFERENCIAS INVITADAS

Fuzzy sets in automated reasoning : Problems and methods	3
<i>D. Dubois & H. Prade</i>	

La Incertidumbre en la Economía y Gestión de Empresas	9
<i>J. Gil Aluja</i>	

Regulación basada en reglas y regulación P.I.D.: Diferencias y complementariedades	15
<i>J. Aguilar & J. Quevedo</i>	

Neural Distributed Processing Systems	23
<i>A.F. Rocha</i>	

COMUNICACIONES

Lógica Fuzzy

Razonamiento Modus Tollens en bases de conocimiento borroso con encadenamiento	31
<i>Bugarín, P. Felix & S. Barro</i>	

El uso de conectivos difusos en el diseño de algoritmos genéticos	37
<i>F. Herrera, M. Lozano & J.L. Verdegay</i>	

A temporal reasoning system based on fuzzy temporal constraints	43
<i>Ll. Godo & Ll. Vila</i>	

Relaciones de indistinguibilidad descomponibles	49
<i>J. Jacas, J. Recasens</i>	

Una aproximación al razonamiento temporal borroso en medicina	53
<i>R. Marín, R. Ruiz, F. Martín , S. Barro & F. Palacios</i>	

Análisis de los rasgos característicos de las funciones de implicación según su comportamiento en un sistema borroso	59
<i>J. Villar & M.A. Sanz</i>	

Query answering in fuzzy temporal constraints	65
<i>Ll. Vila & Ll. Godo</i>	

Fuzzy logic as families of bivaluated logics	71
<i>J.L. Castro</i>	

Sobre la identidad de condicionales materiales	75
<i>E. Trillas, S. Cubillo & A. Rodríguez</i>	

Setup semantics of systems of fuzzy deontic logic	79
<i>F.J. Díez & L. Peña</i>	

Generadores y similitudes duales	85
<i>F. Boixader & J. Jacas</i>	

Teoría

Valuaciones y conjuntos convexos de probabilidad	91
<i>A. Cano & M.T. Lamata</i>	
Números intuicionistas fuzzy	97
<i>P. Burillo, H. Bustince & V. Mohedano</i>	
Negaciones fractales	103
<i>G. Mayor & T. Calvo</i>	
Estudio del retículo de conceptos L-fuzzy	109
<i>A. Burusco & R. Fuentes</i>	
Acerca de la calculabilidad de las funciones difusas	115
<i>R. Morales, J.L. Pérez, B. Clarcs & R. Conejo</i>	
Propiedades de transformación en conjuntos fuzzy	121
<i>M.L. Eraso & M.I. Portilla</i>	
Aplicación de algoritmos evolutivos para el problema de la triangulación en redes causales	127
<i>L.D. Hernández & M.J. Bolaños</i>	
Una estructura elemental de conocimiento fuzzy	133
<i>A. Núñez & J. Delegido</i>	
Semioraciones y semisartas (Estudio de algunos aspectos)	139
<i>A. Sobrino & J.A. Olivas</i>	
Independencia en la teoría de la posibilidad	145
<i>L.M. Campos & J.F. Huete</i>	
An approach to first-order multiple-valued logic programing	151
<i>G. Escalada & F. Manyà</i>	
Síntesis de funciones de pertenencia en la determinación de radios de influencia	157
<i>V. Torra</i>	

Hardware

Realizaciones hardware de sistemas basados en lógica difusa	165
<i>S. Sánchez, C.J. Jiménez, A. Barriga & J.L. Huertas</i>	
Arquitecturas sistólicas de procesamiento difuso y su simulación	171
<i>L. de Salvador, M. García & J. Gutierrez</i>	
Controlador difuso tipo "singleton" en modo de corriente	177
<i>I. Baturone, S. Sánchez, A. Barriga & J.L. Huertas</i>	
A design approach for neuro/fuzzy chips	183
<i>F. Vidal & A. Rodríguez</i>	

Decisión y Optimización

Decisión interactiva para un problema general fuzzy multiobjetivo	191
<i>J.M. Cadenas & F. Jiménez</i>	

Un modelo general de decisión	197
<i>M.T. Lamata</i>	
Ordenación de números borrosos mediante la comparación de sus intervalos esperados	203
<i>M. Jiménez</i>	
El problema del transporte sólido multiobjetivo no lineal con metas fuzzy: un método de solución basado en algoritmos genéticos	209
<i>J.M. Cadenas & F. Jiménez</i>	
Un proceso de decisión en grupo con preferencias lingüísticas y selección secuencial	215
<i>F. Herrera, E. Herrera & J.L. Verdegay</i>	
Un modelo de Goal programming lineal difuso con números difusos triangulares en coeficientes y niveles de aspiración ..	221
<i>M.M. Arenas & M.V. Rodríguez</i>	

Control Fuzzy

Un algoritmo iterativo para la generación de reglas difusas	229
<i>A. González & R. Pérez</i>	
Supervisión de reguladores adaptativos clásicos mediante técnicas difusas	235
<i>R. Ferreiro & B. Rodríguez</i>	
Aprendizaje mediante la aplicación de técnicas fuzzy a espacios discretos	241
<i>J. Riquelme & M. Toro</i>	
Aplicación de algoritmos borrosos al guiado de un robot móvil con visión artificial	247
<i>E. Santiso, M. Mazo, C. Serra, J.J. García & F.J. Rodríguez</i>	
Depuración y generación de reglas de forma automática. Aplicación sobre el control de una planta piloto utilizando autómatas programables con módulo fuzzy	253
<i>J.C. Hernández, F. Jiménez, J. Quevedo & J. Aguilar</i>	
Influencia de las T-normas en controladores difusos	259
<i>E. Cárdenas, J.C. Castillo, O. Cerdón & A. Peregrín</i>	

Aprendizaje

Una metodología de redes neuronales para resolver ecuaciones relacionales	267
<i>A. Blanco, M. Delgado & I. Requena</i>	
Modelado borroso de procesos mediante un algoritmo de simplificación	273
<i>L. Sánchez</i>	
Un método para validar procesos de clustering difuso	279
<i>M. Delgado, M. A. Vila & A.F. Gómez</i>	
Modelo para codificación de información difusa sobre memorias asociativas discretas	285
<i>A. Blanco, W. Fajardo & I. Requena</i>	
Un procedimiento empírico para obtener reglas difusas usando redes neuronales	291
<i>J. M. Benítez, A. Blanco & I. Requena</i>	

Aplicaciones

Análisis de las relaciones entre los atributos de una base de conocimiento usando medidas de información ..	297
<i>M.D. Villatoro, C. Hervás & J.A. Gavilán</i>	

Aplicación de la lógica fuzzy al control de inventarios	303
<i>L.M. Marín & F. Mesa</i>	
Estudio de funciones de composición de evidencias en un problema de reconocimiento de textos	309
<i>J. Etxanobe, J.R. González & J.R. Garitagoitia</i>	
Un procesador elemental de fuzzy SQL	315
<i>J.M. Medina, O. Pons & A. Vila</i>	
Detección de contornos escalón mediante algoritmos difusos	321
<i>E. Montseny, P. Sobrevilla & C. García</i>	
Fuzzy similarities in the field of image analysis	329
<i>R. Albersmann, R. Felix & T. Kretzberg</i>	
Aplicación de una metodología fuzzy a la evaluación de la calidad universitaria	333
<i>F. Mesa & L.M. Marín</i>	
Análisis de texturas naturales mediante técnicas difusas	339
<i>S. Alvarez, D. Puig, P. Sobrevilla & E. Montseny</i>	
 DEMOSTRACIONES DE SOFTWARE	
FuzzyControlDesK	347
<i>R. Albersmann & R. Felix</i>	
FuzzyDecisionDesK	351
<i>R. Albersmann & R. Felix</i>	
MILORD II	353
<i>J. Puyol & C. Sierra</i>	
Índice de autores	355

Prefacio

En vuestras manos tenéis el libro con las conferencias y comunicaciones presentadas en el IV Congreso Español sobre Tecnologías y Lógica Fuzzy auspiciado por la asociación del mismo nombre y cuyo objetivo ha sido, como en ediciones anteriores, promover la investigación en estos campos poniendo un especial énfasis tanto en aquellos que se inician en la investigación como en la interacción entre las Tecnologías Fuzzy y el mundo empresarial. No en vano las aplicaciones de las Tecnologías Fuzzy a nivel mundial han sufrido un enorme desarrollo duplicándose anualmente el número de aplicaciones comerciales durante los últimos cuatro años y llegando a alcanzar en 1993 las 1500 aplicaciones*.

Japón es el país que ha desarrollado más iniciativas en este campo y se encuentra a la cabeza en cuanto a número de aplicaciones desarrolladas, habiendo logrado liderar el mercado mundial. EEUU y Europa han reaccionado más tardíamente ante el empuje y el impacto de los sistemas Fuzzy. Así, por ejemplo, el IV programa marco enfatiza los *"Fuzzy Logic-based Systems"* como tema a desarrollar, lo que demuestra la importancia que la CEE concede al tema.

En nuestro país la investigación en este campo cuenta con una tradición y una calidad nada desdeñables. Se empezó a trabajar en este campo desde casi el mismo inicio de la teoría de los *Fuzzy Sets* en 1965 y hoy en día podemos decir que estamos en el quinto lugar del mundo en cuanto a la producción de artículos científicos si atendemos al número de trabajos presentados en los congresos y publicados en las revistas especializadas. No obstante, se ha de reconocer que las aplicaciones comerciales todavía son pocas y pensamos que debemos hacer un esfuerzo en este sentido si no queremos que una vez más nos encontremos que aquí creamos teorías interesantes y otros países son los que las aplican y obtienen los beneficios comerciales de ello.

Buena muestra de la cantidad y calidad de los grupos de investigación que trabajan en España en este campo son los cincuenta y dos trabajos que se han presentado en este congreso y que se recogen en estas actas agrupados temáticamente en cinco apartados: Lógica Fuzzy, Teoría, Hardware, Decisión y Optimización, Control Fuzzy, Aprendizaje y Aplicaciones. También este año hemos aumentado el número de conferenciantes invitados a cuatro y creemos que la calidad de los temas que desarrollaran puede ayudarnos a tener una idea global del estado de la cuestión en la investigación en Tecnologías y Lógica Fuzzy y en algunas de sus aplicaciones. Por último, hemos intentado que el apartado de las demostraciones de programas permitiera ver aplicaciones a diferentes campos aunque estuvieran en este momento a nivel de prototipo. Así, al cierre de estas actas, podemos asegurar que el congreso contará con demostraciones de aplicaciones de tecnologías Fuzzy al control automático, a las bases de datos deductivas o a los sistemas expertos.

Desde aquí queremos agradecer a los conferenciantes su esfuerzo e interés por colaborar en nuestro congreso y su contribución a la calidad del mismo.

Queremos asimismo expresar nuestro agradecimiento al comité de programa por haber realizado las revisiones de los trabajos con total prontitud y eficacia contribuyendo con ello a la calidad y

Fuzzy Systems - An Overview*. T. Munkata & Y. Jani. *Communications of the ACM*, marzo 1994, **37(3).

al buen desarrollo del congreso, a todas aquellas personas y entidades que de alguna forma han colaborado en este congreso y muy especialmente al Comité Organizador por su abnegada labor, al IIA por su acogida y al Excelentísimo Ayuntamiento de Blanes, a la CIRIT y al CSIC por su colaboración.

Francesc Esteva
Pere García

Presidente

Francesc Esteva *IIIA-CSIC*

Comité de Programa

Josep Aguilar *LAAS-CNRS*
Claudi Alsina *Univ. Politècnica de Catalunya*
Senen Barro *Univ. de Santiago de Compostela*
Miguel Delgado *Univ. de Granada*
Pere García *IIIA-CSIC*
Lluís Godó *IIIA-CSIC*
Julio Gutiérrez *INTA*
Francisco Herrera *Univ. de Granada*
Joan Jacas *Univ. Politècnica de Catalunya*
Ramón López de Mántaras *IIIA-CSIC*
Gaspar Mayor *Univ. de les Illes Balears*
Joseba Quevedo *Univ. Politècnica de Catalunya*
Alejandro Sobrino *Univ. de Santiago de Compostela*
Enric Trillas *Univ. Politècnica de Madrid*
Llorenç Valverde *Univ. de les Illes Balears*
José Luis Verdegay *Univ. de Granada*
Amparo Vila *Univ. de Granada*

Comité Organizador

María del Mar Cuñado
Gonzalo Escalada
Eva López
Felip Manyà
Josep Puyol
M. Angels Valls
Lluís Vila

Entidades Colaboradoras

Comissió Interdepartamental per a la Recerca i Innovació Tecnològica (CIRIT)
Consejo Superior de Investigaciones Científicas (CSIC)
Ajuntament de Blanes
Caixa d'Estalvis i Pensions de Barcelona "La Caixa"

CONFERENCIAS INVITADAS

Fuzzy Sets in Automated Reasoning : Problems and Methods

Didier Dubois and Henri Prade

Institut de Recherche en Informatique de Toulouse (I.R.I.T.) – C.N.R.S.
Université Paul Sabatier, 118 route de Narbonne
31062 Toulouse Cedex – France
Email: {dubois, prade}@irit.irit.fr – Fax: (+33) 61.55.62.39

Abstract

This short paper about fuzzy set-based reasoning first emphasizes the three main semantics for fuzzy sets : similarity, preference and uncertainty. The difference between truth-functional many-valued logics of vague or gradual propositions and possibilistic logic (which handles uncertainty in a non-compositional way) is stressed. Two fuzzy extensions of classical logic entailment, pertaining respectively to nonmonotonic reasoning and to similarity-based approximate reasoning are contrasted. Various kinds of fuzzy set-based reasoning are briefly outlined : fuzzy deductive reasoning about flexible constraints, reasoning under uncertainty and inconsistency, exception-tolerant plausible reasoning using generic knowledge, hypothetical reasoning in possibilistic logic, interpolative reasoning and abductive reasoning with uncertain relationship between disorders and manifestations and with uncertain observations. Open problems are listed in the conclusion.

Key Words

similarity, preference, uncertainty, nonmonotonic reasoning, approximate reasoning, possibilistic logic

1 The semantics of fuzzy sets

Fuzzy sets seem to be relevant in three classes of applications : classification and data analysis, reasoning under uncertainty, and decision-making problems. These three directions, that have been investigated by many researchers, actually correspond and/or exploit three semantics of the membership grades, respectively in terms of similarity, uncertainty and preference. Indeed, considering the degree of membership $\mu_F(u)$ of an element u in a fuzzy set F , defined on a referential U , one can find in the literature, three interpretations of this degree

- degree of similarity : $\mu_F(u)$ is the degree of proximity of u from prototype elements of F . Historically, this is the oldest semantics of membership grades (Bellman, Kalaba and Zadeh, 1966). This view is particularly suitable in classification, clustering, regression analysis and the

like, where the problem is that of abstraction from a set of data ;

- degree of preference : F represents a set of more or less preferred objects (or values of a decision variable x) and $\mu_F(u)$ represents an intensity of preference in favor of this object, or the feasibility of selecting u as a value of x . This view is the one later put forward by Bellman and Zadeh (1970) ; it has given birth to an abundant literature on fuzzy optimisation and decision analysis. Applications pertain to design and planning problems ;

- degree of uncertainty : this interpretation was proposed by Zadeh (1978) when he introduced possibility theory. $\mu_F(u)$ is then the degree of possibility that a parameter x has value u , given that all that is known about it is that " x is F ". Possibility can convey an epistemic meaning (F then describes the more or less plausible values of x) or a physical meaning ($\mu_F(u)$ being the degree of ease of having $x=u$). Viewing $\mu_F(u)$ as a degree of uncertainty only refers to the epistemic interpretation. The physical meaning of possibility has more to do with preference and feasibility.

Note that some semantics are in some sense more basic than others in the sense that they can be operationally measured : similarity is evaluated via distance indices, and feasibility may be measured by costs (as in toll sets, Dubois and Prade, 1993). Preference and plausibility are derived notions, among others because of their epistemic nature. Preference may be a function of cost ; it may also reflect similarity with respect to an ideal object. Plausibility may also be viewed as similarity with respect to a maximally plausible situation. However it can be as well construed as an upper bound on a frequency (and this is how probability and possibility get connected) or a likelihood function.

Another point is that although fuzzy sets have been introduced with a numerical flavor, a membership function is not necessarily mapped on a set of numbers, but an ordered set such as a complete lattice is enough. A membership function can even be construed as an ordering relation $\geq_F$, attached to a

predicate F , where $u \geq_F u'$ means that u is more F than u' . A fuzzy relation R on $U \times U$ can also be viewed as a ternary relation (on U^3), i.e., a collection $\{\geq_u, u \in U\}$ of binary relations that are complete preorderings. Then $\mu_R(u, u') \geq \mu_R(u, u'')$ can represent a situation where $u \geq_u u''$, which reads u' is closer to u than u'' . These structures are common in conditional logics of counterfactuals (Lewis, 1973).

However, reducing a fuzzy set to an ordering relation between elements prevents one from directly defining fuzzy set-theoretic aggregations. It has been well known from a long time that there is no simple way of aggregating ordering relations (due to Arrow's impossibility theorem). These difficulties are by-passed by assuming that degrees of membership to fuzzy sets pertaining to unrelated concepts are commensurate. This is done by resorting to a common membership scale (that need not be numerical). This commensurability assumption is often taken for granted and never emphasized in the literature. Yet, it is better to state it clearly in the face of newcomers to the theory, as it sheds light on potential limitations of the fuzzy approach and helps locating it within alternative settings such as measurement theory and multiple-criteria decision analysis.

2 Fuzziness versus uncertainty

From now on, let us consider the relevance of fuzzy set theory for knowledge representation and automated reasoning, and more generally Artificial Intelligence. First there always remains some confusion in the fuzzy literature on the potential of fuzzy sets for handling uncertainty. This state of confusion can be exemplified in some texts by fuzzy set proponents claiming that probability theory models randomness while fuzzy set models subjective uncertainty (hence ignoring subjective probability). It is also present in the expert systems literature where certainty factors have been confused with membership grades. It also pervades the antifuzzy literature where the truth-functionality of conjunction, disjunction and negation in fuzzy logic is considered as mathematically inconsistent [see (Elkan, 1993) for a recent restatement of this fallacy].

In fact, insofar as fuzzy sets model vague predicates, membership grades model degrees of truth of vague propositions, not degrees of uncertainty. In the scope of knowledge representation, truth is a matter of convention, while uncertainty reflects incomplete or contradictory knowledge. In classical logic the convention is that truth is binary. Fuzzy set theory (and before, multivalued logics) has modified this convention. This shift in convention does not entitle degrees of

truth to be interpreted as degrees of uncertainty.

Degrees of uncertainty can be attached to a crisp proposition in order to model the fact that it is not known whether this proposition is true or false. Uncertainty is at the meta-level with respect to truth. In classical logic truth is binary (true or false) while uncertainty is ternary (surely true, surely false, or unknown). In probability theory, truth is usually binary (crisp propositions) while uncertainty takes on all values in the unit interval. In fuzzy set theory truth is many-valued but there is no uncertainty insofar as the element, the membership grade of which is computed, is precisely located. Fuzzy-truth values (which Zadeh has claimed to be typical of fuzzy logic) are possibility distributions that describe partially unknown truth-values.

As pointed out above, it happens that a degree of membership $\mu_F(u)$ is interpreted as a degree of uncertainty instead of a degree of truth. But this degree of uncertainty is not attached to the fuzzy proposition ' X is F ' when it is known that $X=u$, but to the non-fuzzy proposition $X=u$, when all that is known is that the value of X is somewhere in the support of F .

This discussion (see also Dubois and Prade, 1994) leads to two entirely different extensions of classical logic that exploits the notion of a fuzzy set :

- *many-valued logics* that were proposed before fuzzy set theory came to light. They are exclusively devoted to the handling of vague propositions $\tilde{p}$. The underlying algebraic structure is weaker than a Boolean algebra, and can be consistent with truth-values $t(\tilde{p})$ that lie in the unit interval and remain truth-functional. The truth-functionality is not compulsory however. For instance suppose that the vagueness of predicates stems from an indistinguishability relation (Ruspini, 1991) that blurs the set of interpretations of the language. Then a Boolean proposition p is actually understood as a fuzzy proposition $\tilde{p}$ whose models are *close* to models of p . Let $[p]$ be the set of models of p and $R(p)=[p] \circ R$ the fuzzy set of models close (in the sense of fuzzy relation R) to models of p . Then generally, $R(p \wedge q) \subseteq R(p) \cap R(q)$, without equality, so that truth-values $t(p)$ are not truth-functional with respect to conjunction.

- *possibilistic logic* that is built on top of classical logic, and where each crisp proposition is attached a degree of certainty $N(p)$ such that $N(p) = 1$ iff p is surely true and $N(p)=0$ expresses the complete lack of certainty that p is true (either p is false when $N(\neg p) = 1$, or it is unknown if p is true or false and then $N(\neg p)=0$). The degree $N(p)$ is compositional *for conjunction only* ($N(p \wedge q) = \min(N(p), N(q))$), and

$N(p \vee q) \geq \max(N(p), N(q))$ generally. For instance, if $q = \neg p$, $p \vee q$ is tautological, hence surely true ($N(p \vee q) = 1$), but p may be unknown ($N(p) = N(\neg p) = 0$). Moreover $N(\neg p) = 1 - \prod(p)$ where $\prod(p)$ is the degree of possibility of proposition p . Functions N and $\prod$ stem from the existence of a fuzzy set of more or less possible worlds, one of which is the actual one. It is described by means of a possibility distribution π on the interpretations of the language, and $N(p) = 1 - \sup\{\pi(\omega), \omega \models \neg p\} = 1 - \prod(\neg p)$, i.e., $N(p)$ is computed as the degree of impossibility of the proposition $\neg p$.

In possibilistic logic a fuzzy set describes incomplete knowledge about where the actual world is, while in many-valued logics fuzzy sets describe the extensions of vague predicates.

3 Two fuzzy extensions of classical entailment

In classical logic, a proposition p entails another proposition q whenever each situation where p is true is a situation where q is true. Entailment is denoted $\models$, and $p \models q$ means $[p] \subseteq [q]$ where $[p]$ is the set of models of p .

In possibilistic logic, the type of inference which is at work is plausible inference. The possibility distribution π on the set of interpretations encodes an ordering relation that ranks possible worlds ω in terms of plausibility. Then p plausibly entails q if and only if q is true in the most plausible situations where p is true, i.e., $\max[p] \subseteq [q]$, where $\max[p] = \{\omega \in [p], \pi(\omega) \text{ maximal}\}$. This type of inference is also called "preferential inference" in nonmonotonic reasoning. In contrast, inference in similarity logics (introduced by Ruspini (1991)) is dual to preferential inference. The set of interpretations of the language is equipped with a similarity relation R . Then p entails approximately q in similarity logic if all the situations where p is true are close to situations where q is true, i.e., $[p] \subseteq \text{support}[R(q)]$. More generally, a degree of strength of the entailment can be computed as

$I(q \mid p) = \inf_{\omega \in [p]} \sup_{\omega' \in [q]} \mu_R(\omega, \omega')$. It plays in similarity logic the same role as a degree of confirmation in inductive logic. In possibilistic logic the counterpart of $I(q \mid p)$ is the conditional necessity $N(q \mid p)$ computed from π , i.e., $N(q \mid p) = N(\neg p \vee q) > 0$ if $\prod(q \wedge p) > \prod(\neg q \wedge p)$, and $N(q \mid p) = 0$ otherwise.

These newly emerged notions of fuzzy set-based inference certainly deserve further developments, and lead to very different types of logic. Of interest is the

study of their links to the usual entailment principle of Zadeh, and various extensions of consequence relations in multiple-valued logic, as studied by Chakraborty (1988), Castro et al. (1994). Comparison and combinations of similarity logic and possibilistic logic are also worth-investigating (e.g., Godo et al., 1994 ; Dubois and Prade, 1992).

4 Various kinds of approximate reasoning

Fuzzy sets offer a powerful tool for the modelling of various kinds of commonsense reasoning, where the three semantics of membership functions interfere. Let us review some of them.

Fuzzy deductive inference

This type of approximate reasoning has been advocated by Zadeh in the mid-seventies, as a calculus of fuzzy restrictions. The principles of this approach are as follows : consider a set of statements (in natural language, for instance). Each statement is translated into a possibility distribution that relates concerned variables. Then all possibility distributions are conjunctively combined (using minimum operation) into an overall possibility distribution π . Reasoning consists in projecting π over various subsets of variables of interest. This type of inference is a generalization of constraint propagation to flexible constraints, provided that one interprets each statement as a requirement that some controllable variable must satisfy. Then the possibility distributions model preference, and inference come down to consistency analysis (such as arc-consistency, path-consistency, etc) in the terminology of constraint-directed reasoning. The advantage of fuzzy deductive inference is to directly account for flexible constraints and prioritized constraints, where the priorities are modelled by means of necessity functions (see Dubois, Fargier and Prade, 1993).

Reasoning under uncertainty and inconsistency

A possibilistic knowledge base K is a set of pairs (p, s) where p is a classical logic formula and s is a lower bound of a degree of necessity ($N(p) \geq s$). It can be viewed as a stratified deductive data base where the higher s , the safer the piece of knowledge p . Reasoning from K means using the safest part of K to make inference, whenever possible. Denoting $K_\alpha = \{p, (p, s) \in K, s \geq \alpha\}$, the entailment $K \vdash (p, \alpha)$ means that $K_\alpha \vdash p$. K can be inconsistent and its inconsistency degree is $\text{inc}(K) = \sup\{\alpha, K \vdash (\perp, \alpha)\}$ where $\perp$ denotes the contradiction. In contrast with classical logic, inference in the presence of

inconsistency becomes non trivial. This is the case when $K \vdash (p, \alpha)$ where $\alpha > \text{inc}(K)$. Then it means that p follows from a consistent and safe part of K (at least at level α). This kind of syntactic non-trivial inference is sound and complete with respect to the above defined preferential entailment. Moreover adding p to K and nontrivially entailing q from $K \cup \{p\}$ corresponds to revising K upon learning p , and having q as a consequence of the revised knowledge base. This notion of revision is exactly the one studied by Gardenfors (1988) at the axiomatic level.

Nonmonotonic plausible inference using generic knowledge

Possibilistic logic does not allow for a direct encoding of pieces of generic knowledge such as "birds fly". However, it provides a target language in which plausible inference from generic knowledge can be achieved in the face of incomplete evidence. In possibility theory "p generally entails q" is understood as " $p \wedge q$ is a more plausible situation than $p \wedge \neg q$ ". It defines a constraint of the form $\prod(p \wedge q) > \prod(p \wedge \neg q)$ that restricts a set of possibility distributions. Given a set S of generic knowledge statements of the form " p_i generally entails q_i ", a possibilistic base can be computed as follows. For each interpretation ω of the language, the maximal possibility degree $\pi(\omega)$ is computed, that obeys the set of constraints in S . This is done by virtue of the principle of minimal specificity (or commitment) that assumes each situation as a possible one insofar as it has not been ruled out. Then each generic statement is turned into a material implication $\neg p_i \vee q_i$, to which $N(\neg p_i \vee q_i)$ is attached. It comes down, as shown in Benferhat et al. (1992) to rank-order the generic rules giving priority to the most specific ones, as done in Pearl (1990)'s system Z . A very important property of this approach is that it is exception-tolerant. It offers a convenient framework for implementing a basic form of nonmonotonic system called "rational closure" (Lehmann and Magidor, 1992), and addresses a basic problem in the expert system literature, that is, handling exceptions in uncertain rules.

Hypothetical reasoning

The idea is to cope with incomplete information by explicitly handling assumptions under which conclusions can be derived. To this end some literals in the language are distinguished as being assumptions. Possibilistic logic offers a tool for reasoning with assumptions. It is based on the fact

that in possibilistic logic a clause $(\neg h \vee q, \alpha)$ is semantically equivalent to the formula with a symbolic weight $(q, \min(\alpha, t(h)))$ where $t(h)$ is the (possibly unknown) truth value of h . The set of environments in which a proposition p is true can thus be calculated by putting all assumptions in the weight slots, carrying out possibilistic inference so as to derive p . The subsets of assumptions under which p is true with more or less certainty can be retrieved from the weight attached to p . This technique can be used to detect minimal inconsistent subsets of a propositional knowledge base (see Benferhat et al., 1994).

Interpolative Reasoning

This type of reasoning is at work in fuzzy control applications, albeit without clear logical foundations. Klawonn and Kruse (1993) have shown that a set of fuzzy rules can be viewed as a set of crisp rules along with a set of similarity relations. Moreover an interpolation-dedicated fuzzy rule 'if is A then Y is B ' can be understood as "the more x is A the more Y is B " and the corresponding inference means that if $X=x$ and $\alpha = \mu_A(x)$ then Y lies in the level cut B_α . When two rules are at work, such that $\alpha_1 = \mu_{A_1}(x)$, $\alpha_2 = \mu_{A_2}(x)$, then the conclusion $Y \in (B_1)_{\alpha_1} \cap (B_2)_{\alpha_2}$ lies between the cores of B_1 and B_2 , i.e., on ordered universes, an interpolation effect is obtained. It can be proved that Sugeno's fuzzy reasoning method for control can be cast in this framework (Dubois, Grabisch, Prade, 1994). More generally interpolation is clearly a kind of reasoning based on similarity (rather than uncertainty) and it should be related to current research on similarity logics. More generally similarity relations and fuzzy interpolation methods should impact on current research in case-based reasoning.

Abductive reasoning

Abductive reasoning is viewed as the task of retrieving plausible explanations of available observations on the basis of causal knowledge. In fuzzy set theory causal knowledge has often been represented by means of fuzzy relations relating a set of causes C to a set of observations S . However the problem of the semantics of this relation has often been overlooked. $\mu_R(c, s)$ may be viewed either as a degree of intensity or a degree of uncertainty. Namely, when observations are not binary, $\mu_R(c, s)$ can be understood as the intensity of presence of observed symptom s when the cause c is present. This is the traditional view in fuzzy set theory. It leads to Sanchez (1977) approach to abduction, based

on fuzzy relational equations. Another view has been recently proposed by the authors, where by $\mu_R(c,s)$ is understood as the degree of certainty that a binary symptom s is present when c is present. A dual causal matrix R' must be used where $\mu_{R'}(c,s)$ is the degree of certainty that a binary symptom s is absent when c is present. On such a basis the theory of parsimonious covering for causal diagnosis by Peng and Reggia (1990) can be extended to the case of uncertain causal knowledge and incomplete observations. This method is currently applied to satellite failure diagnosis (Cayrac et al., 1994).

5 Conclusion and open problems

Fuzzy sets open a lot of new areas of research, or at least a new way of investigating old questions in a unified way. They are useful to model vague predicates in reasoning, but also for coping with incomplete information. Fuzzy sets are also a natural framework for reasoning with similarity and preference. A number of open problems remains, some of which are listed hereafter :

A syntax for fuzzy logic

A lot of works have been done in fuzzy logic at the semantic level, but few papers address the question of finding a proper syntax. Most of these papers (e.g., Pavelka, 1979) put intermediary truth-values in the object language. The only well-known axiomatic system is the one for Lukasiewicz logic, that relies on MV-algebras. But Zadeh's fuzzy logic is not based on MV-algebra. Offering a syntax to Zadeh's fuzzy logic may cancel the main objection raised by logicians against fuzzy set theory.

Refining the min-based ordering

Fuzzy constraint satisfaction is based on the idea of finding a solution that best satisfies the most violated constraint. A lot of solutions may remain indiscriminated in such a process. The question is how to refine this ordering while remaining faithful to that egalitarian principle. Several proposals are being studied that derive from set-inclusion and lexicographic methods borrowed from nonmonotonic reasoning and social sciences (e.g., Fargier et al., 1993).

Implementing nonmonotonic reasoning systems in possibilistic logic.

So far only rational and preferential inference have been captured. Some other forms of nonmonotonic reasoning approaches such as circumscription (Benferhat et al., 1993) and default theory (Yager, 1987) should be investigated with this purpose in view.

Expressing and handling independence statements in possibility theory

Commonsense exception-tolerant reasoning relies not only on generic rules, but also on independence assumptions. In logic the treatment of independence is done systematically due to the monotonicity of the inference. Independence is well-mastered in probability theory but deserves a lot of investigations in possibility theory. Preliminary results can be found in (Dubois et al., 1994)

Defining a full-fledged logic of similarity

The purpose is to handle approximate reasoning problems and account for interpolation. Especially it might be useful to cast both similarity logic and possibilistic logic in the same setting so as to come up with a system of approximate inference under incomplete information.

Fuzzy abduction with prior knowledge

All fuzzy relational abductive reasoning methods neglect the fact that the retrieval of causes is partially based on prior knowledge especially when observations are scarce. Bayesian methods of diagnosis exploit prior knowledge. How to account for prior knowledge using fuzzy methods ?

References

- [1] BELLMAN R., KALABA L. AND ZADEH L.A. (1966) Abstraction and pattern classification. J. Math. Anal. Appl., 13,1-7
- [2] BELLMAN R., ZADEH L.A. (1970) Decision making in a fuzzy environment. Management Science, 17, B141-B164
- [3] BENFERHAT S., DUBOIS D. AND PRADE H. (1992) Representing default rules in possibilistic logic, Proc. of the 3rd Inter. Conf. on Principles of Knowledge Representation and Reasoning, KR'92 Cambridge, Ma, 673-684
- [4] BENFERHAT S., DUBOIS D., LANG J. AND PRADE H. (1994) Hypothetical reasoning in possibilistic logic : basic notions, applications and implementation issues. In : Between Mind and Computer. Fuzzy Science and Engineering (Wang P.Z., Loe K.F. eds.), 1-29
- [5] CASTRO J.L., TRILLAS S. AND CUBILLA S. (1994) On consequence in approximate reasoning. J. Applied Non-Classical Logics, 4, 91-103
- [6] CAYRAC D., DUBOIS D., HAZIZA M. AND PRADE H. Possibility theory in "fault mode effect analyses" -A satellite fault diagnosis application Proc IEEE World Cong. on Computational Intelligence Orlando, Fl, June 26-July 2, 1994, 1176-1181

- [7] CHAKRABORTY M. (1988) Use of fuzzy set in introducing graded consequence in multiple-valued logic. In : *Fuzzy Logic in Knowledge Based Systems-Decision and Control* (Gupta M.M., Yamakawa T., eds) North-Holland, 247-257
- [8] DUBOIS D., FARGIER H. AND PRADE H. (1993) Flexible constraint satisfaction problems with application to scheduling problems. Report IRIT/93-30R, Toulouse, France
- [9] DUBOIS D., FARINAS DEL CERRO L., HERZIG A. AND PRADE H. (1994) An ordinal independence with application to plausible reasoning. Proc. 10th Uncertainty in AI Conference, Seattle, Wa, July, 28-30, 1994
- [10] DUBOIS D., GRABISCH M. AND PRADE H. (1994) Gradual rules and the approximation of control laws. In : *Theoretical Aspects of Fuzzy Control* (H.T. Nguyen, M. Sugeno, R. Tong, R. Yager, eds.), Wiley, to appear
- [11] DUBOIS D., LANG J. AND PRADE H. (1994) Possibilistic logic. In : *Handbook of Logic in Artificial Intelligence and Logic Programming* (Gabbay D. et al., eds.) Oxford Univ. Press, Vol. 3, 439-513
- [12] DUBOIS D., PRADE H. (1992) Putting rough sets and fuzzy sets together. In *Intelligent Decision Support* : (Slowinski R., ed.) Kluwer Acad. Pub., 203-232
- [13] DUBOIS D., PRADE H. (1993) Toll sets and toll logic. In : *Fuzzy logic- State of the Art*. (Lowen R., Roubens M., eds) Kluwer Academic Pub., 169-180
- [14] DUBOIS D., PRADE H. (1994) Can we enforce compositionality in uncertainty calculi ? Proc. of the 12th American Association for AI Conference, Seattle, Wa, July 31-Aug. 4, 1994
- [15] ELKAN C. (1993) The paradoxical success of fuzzy logic. Proc. of the 11th. American Association for AI Conf., Washinston DC., 698-703
- [16] ESTEVA F., GARCIA-CALVES P. AND GODO L. Relating and extending semantical approaches to possibilistic reasoning. *Int. J. of Approximate Reasoning*, 10,311-344
- [17] FARGIER H., LANG J. AND SCHIEX T. Selecting preferred solutions in fuzzy constraint satisfaction problems. Proc. 1st Europ. Cong. on Fuzzy and Intellig. Technologies, EUFIT'93, Aachen, 1128-1134
- [18] GARDENFORS P. (1988) *Knowledge in Flux*. MIT Press
- [19] KLAWONN F., KRUSE R. (1993) Equality relations as a basis for fuzzy control. *Fuzzy Sets & Syst.*, 54, 147-156
- [20] LEHMANN D., MAGIDOR M. (1992) What does a conditionnal knowledge base entail ? *Artificial Intelligence*, 55, 1-60
- [21] LEWIS D. (1973) *Counterfactuals*. Basil Blackwell
- [22] PAVELKA J. (1979) On fuzzy logic. I-II-III. *Zeitschr. f. Math. Logik und Grundlagen d. Math.*, 25, 45-52, 119-134, 447-464
- [23] PEARL J. (1990) System Z : a natural ordering of defaults with tractable applications to default reasoning. Proc. 3rd Conf. on the Theoretical Aspects of Reasonig About Knowledge TARK'90, Morgan & Kaufmann, 1990, 121-135
- [24] PENG Y., REGGIA (1990) *Abductive Inference Models for Diagnostic Problem-Solving*. New York : Springer Verlag
- [25] RUSPINI E. (1991) On the semantics of fuzzy logic. *Int. J. of Approximate Reasoning*, 5, 45-88
- [26] SANCHEZ E. (1977) Solutions in conposite fuzzy relations equations. Application to medical diagnosis in Brouwerian logic. In : *Fuzzy Automated and Decision Processes* (M.M. Gupta et al. ed.) North-Holland, 221-234
- [27] YAGER R.R. (1987) Using approximate reasoning to represent default knowledge. *Artificial Intelligence*, 31, 99-112
- [28] ZADEH L. (1978) Fuzzy sets as a basis for a theory of possibility. *Fuzzy Sets and Systems*, 1,3-28

La Incertidumbre en la Economía y Gestión de Empresas

Prof. Dr. Jaime Gil Aluja

Universidad de Barcelona

Avda. Diagonal 612

Fax: + 34.3.201 80 47

08021 Barcelona, España

Resumen

El largo letargo en el que estaban sumidas las técnicas operativas para el estudio de la economía y gestión de empresas, parece haber llegado a su fin con la incorporación de unos nuevos elementos, con base en la teoría de los subconjuntos borrosos.

Los textos vigentes hasta hace poco en las Universidades de los países con mayor tradición y prestigio, han utilizado como instrumentos para la descripción y tratamiento de la realidad económica de los países y las empresas, unos modelos contruidos con la utilización de las matemáticas de la certeza y del azar. Pero las transformaciones operadas en la sociedad, los cambios que se han ido produciendo en los sistemas empresariales, han obligado a una reflexión en torno a la idoneidad de tales modelos.

Tras unos años de trabajos y después de no pocos ensayos tendentes a insertar una nueva manera de enfocar la solución a los problemas económicos-financieros, se ha llegado a un momento en el que parece reducirse la resistencia, por otro lado comprensible, de quienes, profesores e investigadores, se hallan asentados en unos conocimientos que la tradición les ha legado. y así, junto a unas matemáticas de la certeza y del azar, se están empleando, ahora, unas matemáticas de la incertidumbre.

Palabras clave: economía, expertos, gestión, incertidumbre, subconjunto borroso.

1 Introducción

El saber científico no debe describir y tratar el universo que nos gustaría vivir sino el que

realmente vivimos. Para ello será necesario poner fin a los conocimientos "sagrados" de las leyes ciertas que describen un mundo estable y crear una nueva racionalidad basada en la inestabilidad que conduce a la incertidumbre, aunque transgredir estas leyes exija volver a pensar las ciencias e incluso la filosofía. Asimismo será necesario hallar un lenguaje que permita responder a las preguntas profundas que se están planteando en estos últimos años del siglo.

No es posible olvidar que la ciencia occidental, consagrada durante muchos siglos, se basa en una idea que es, en sí misma, original: la idea de las leyes de la naturaleza. Según ella, la naturaleza está obligada a seguir ciertas reglas que conducen a estructuras basadas en la certeza, en las cuales no existe diferencia entre pasado y futuro: se trata, por lo tanto, de una visión intemporal del universo. Esto queda ya patente en las leyes de Newton, y curiosamente, aquellas que se han considerado como las dos grandes revoluciones del siglo XX, la "mecánica cuántica" y la "relatividad", no han hecho más que abonar esta visión. A pesar de ello, cada vez resulta más patente que esos planteamientos se hallan en contradicción con el aspecto evolutivo del universo, y, por ello, con el aspecto evolutivo del ser humano.

Creemos que hay que llegar a una idea darwiniana de la evolución del universo, y dejar aparte una idea geométrica. Es necesario describir el universo y como este universo es inestable se produce la incertidumbre. Pero incluso de la incertidumbre se pueden extraer ciertos comportamientos expresables la mayor parte de ellos mediante posibilidades, algunos a través de probabilidades y muy pocos por la certeza.

Pero no resulta honesto cerrar los ojos y decir que se cree en las certezas, aunque éstas no residen en nuestro mundo. La incertidumbre puede ser una forma de representar las inestabilidades y a través de ella intentar comprender el papel que juegan en unas reglas de la naturaleza. En el fondo, intentamos encontrar una puerta angosta, en un universo determinista que resulta alienante ya que en él todo se halla predeterminado e inscrito en el Big-Bang. Será un universo incierto, pero que responde a unas determinadas normas de conducta. La novedad se halla en la conversión de normas de conducta de los sistemas inciertos en leyes que pueden ser formalizadas de una manera cierta.

2 Notas acerca de la evolución de las ciencias económicas

En este marco, la fenomenología económica no constituye una excepción. A diferencia de cuanto sucedía en épocas pretéritas, en las cuales los acontecimientos tenían lugar de manera lenta y la evolución se producía a través de extensos períodos en los cuales la capacidad de reacción frente a los cambios era total, actualmente la actividad social se halla en permanente ebullición.

Ha sido sobre todo a lo largo de los últimos decenios cuando más importantes han sido los cambios, no sólo en cuanto a los hechos y los fenómenos, sino también en los comportamientos y en las ideas. Así, valores tradicionales como la laboriosidad, perseverancia, paciencia que durante tantos siglos han constituido un todo monolítico y una guía para los seres en sociedad, han estallado hecho añicos para ser sustituidos por otros valores tales como la audacia, el espíritu competitivo, ...apareciendo lo que ya se ha denominado "reino de la imagen". Ante este contexto, parece lícito preguntarse cómo se puede concebir una actividad científica cuando el pensamiento humano, cargado de un alto grado de subjetividad, intenta encontrar, entre tanto cambio, lo objetivo. Quizás sea en la búsqueda y no en el hallazgo, cuando se desarrolla el conocimiento.

Pero como afirmaba, con tanta frecuencia, François Perroux: la ciencia se desarrolla a través de las necesidades de la realidad de cada

momento y como consecuencia de ello, continuaba, las estructuras sociales actúan y condicionan el pensamiento económico. No es de extrañar, por ello, que la estabilidad en el progreso de los conocimientos económicos ha dado paso a la eclosión de nuevas ideas que, en direcciones muy diferenciadas, pretenden dar respuesta a los numerosos problemas que la sociedad actual tiene planteados.

La economía, quizás la más joven entre la ciencias sociales, aparece de manera tardía y el pensamiento que gira a su entorno, se estructura inicialmente en base a una matemática mecanicista entre 1880 y 1914 con el Equilibrio General (Walras, Pareto, Cournot, Edgeworth, ...). Se observa que la incorporación de conceptos, métodos y técnicas, no siempre ha sido fruto de una elaboración interna.

En efecto, frente a una realidad cuya característica fundamental es la estabilidad en la vida social y en las relaciones económicas, grupos de estudiosos intentan formalizar los procesos que en ella acontecen. En este intento se aproximan a otros campos del conocimiento y observan los elementos que en ellos se utilizan. En el ámbito de la física y de la química, entre otros, la experimentación tiene lugar en objetos o sustancias, cuyo comportamiento se cree mecánico, por lo que la matemática de la certeza constituye un instrumento idóneo para describir la naturaleza de los fenómenos que se producen.

La transposición al campo de la economía da lugar al nacimiento del marginalismo, que necesita crear un ente artificial como sujeto de la actividad económica: el "homo oeconomicus". Este ser, desprovisto de toda sensación, toma decisiones únicamente a través de estímulos racionales, puramente mecánicos. Se trata, en definitiva de una primera versión del "hombre-robot", un hombre carente de alma y de imaginación. En este contexto se utiliza la mecánica clásica de Lagrange, lo que da una impresión de rigor, en lo que Perroux llama al "laxismo del discurso económico". El hombre se halla sujeto a unas leyes, que le llevan indefectiblemente a un futuro predeterminado.

Esta deformación y a la vez resunción de los acontecimientos, ha hecho que la matemática del determinismo haya tenido tanto predicamento, y haya imperado y continúe imperando aún hoy, en muchos ámbitos de la actividad científica en la economía y gestión de empresas. Pero con el

devenir de los tiempos, se inician en la sociedad importantes cambios que tienen cada vez mayor presencia en el campo económico. Se levanten ciertas voces clamando por una nueva manera de enfocar los problemas y subrayando la insuficiencia de la matemática mecanicista para describir la nueva sociedad que estaba emergiendo.

Este cambio radical, se produce a partir de la segunda guerra mundial. Se renuncia a la figura del "hombre robot", se considera el tiempo irreversible, se evita formalizar al fatalismo de la predestinación, dando al hombre oportunidad de elegir libremente su futuro, un futuro del que es protagonista activo y no mero engranaje de una cadena predeterminada.

De nuevo una segunda generación de estudiosos de la economía buscan, en otros ámbitos del conocimiento, los instrumentos que permitan describir primero y tratar después, una realidad que va conociendo sucesivas mutaciones y que da lugar a una fenomenología de muy difícil descripción con los conocimientos oficiales existentes. Una vez más se gira la cabeza hacia otras ciencias y se comprueba que, por ejemplo, en el ámbito de la física o de la química, un experimento es capaz de ser repetido prácticamente en las mismas o similares condiciones un elevado número de veces. De estas repeticiones los científicos son capaces de extraer una ley de comportamiento. Se incorporan, así, las matemáticas del azar en los estadios económicos y de gestión. Se trata de un nuevo mecanicismo pero que resulta mucho más adecuado para tratar un mundo en el que los cambios constituyen, cada vez más, uno de los elementos centrales. A partir de ley empírica del azar primero, y de los axiomas de Borel-Kolmogorov después, el concepto de probabilidad constituye el elemento que permite un cambio fundamental en la investigación científica de la economía y gestión de empresas.

Pero la situación actual, caracterizada por unos cambios brusco e inesperados en direcciones muchas veces contrapuestas, nos ha llevado, en los últimos años, a replantear de nuevo el empleo de las técnicas normalmente utilizadas para el tratamiento de una realidad, que de tan cambiante se ha convertido en incierta. Porque ¿es posible pensar en este final de siglo que los fenómenos sociales, económicos y de gestión, que se producen en el gobierno de las

naciones y en la dirección de las empresas son susceptibles de repetición en iguales o parecidas circunstancias? Desde otra perspectiva ¿se puede decir que en la actualidad los fenómenos sociales, económicos o de gestión cumplen la axiomática de Borel-Kolmogorov? Difícilmente nuestra realidad permite una respuesta afirmativa.

3 La incertidumbre en la economía y la gestión de empresas

En el ámbito de las ciencias económicas, el concepto de decisión constituye uno de los términos más utilizados. Tanto es así que, para muchos, la economía es la ciencia de la decisión. Pues bien, por el hecho de que en los sistemas económicos se están produciendo procesos de aceleración y desaceleración que no siempre van en la misma dirección, tienen lugar en un seno tensiones, de naturaleza diversa, que provocan importantes problemas derivados de la ausencia de una plataforma de futuro con la suficiente estabilidad para tomar decisiones en base a la previsión de magnitudes que permita acotar convenientemente el devenir de los acontecimientos.

En este ambiente, los ejecutivos deben adoptar unas decisiones con una repercusión económica y financiera que no se limita al momento en que son tomadas sino que se prolongan, en muchos casos, a lo largo de varios años. Las dificultades de previsión y estimación, consustanciales en todo ejecutivo, van aumentando cada vez más como consecuencia de un reciente clima de incertidumbre. Resulta evidente que "los hechos de la naturaleza son inciertos; el entorno económico, social, financiero de las empresas cambia incesantemente; los actos del hombre - porque es libre y dotado de imaginación- como las relaciones entre los hombres porque éstos no son robots, son las causas profundas de la incertidumbre" [1].

Nuestra preocupación, nuestros trabajos y reflexiones [2] van encaminadas a poner de manifiesto que incluso sin poder medir de manera formal o mediante la probabilidad, también se puede aspirar a un comportamiento racional. Los hechos susceptibles de verdadera repetición, pertenecen quizás al ámbito de la naturaleza, de la física, de la química, de la astronomía, e incluso de la biología, pero el

hombre introduce, además de los hechos inciertos de la naturaleza, los que provienen de su libertad, de su poder de imaginación.

En el ámbito de la economía y la gestión de empresas, se han realizado intentos, creemos que de manera parcial logrados, de crear unos elementos capaces de realizar el adecuado tratamiento de los fenómenos que tienen lugar en el seno de los estados y de las empresas, cuando su conocimiento tiene lugar de manera poco precisa. Para ello se han utilizado la teoría de los errores, la de los intervalos de confianza, la de los números borrosos, de los subconjuntos borrosos y todas las generalizaciones propuestas que ya hemos empleado[3]. Pero insistimos en un punto: si se puede medir, hay que hacerlo. Y si no se puede medir utilizar en su defecto lo perceptible.

En este sentido, se puede comprobar que, en los esquemas tradicionales, dada la imposibilidad de recoger con precisión la compleja e incierta realidad económica, se recurre a una simplificación inicial para realizar los desarrollos posteriores en base a estos elementos simplificados. Las posibles desviaciones iniciales se van acumulando y ampliando a medida que el proceso operativo avanza. Se pierde además una información desde el principio que ya no se va a recuperar.

Por nuestra parte preferimos recoger los fenómenos económicos y de gestión con su incertidumbre, para realizar los pertinentes desarrollos conservando la imprecisión (y también toda la información) para hacerla "caer" lo más tarde posible, dado que siempre es posible (perdiendo información) reducir la incertidumbre.

Las posibilidades de utilización de este camino han sido numerosas y van desde las previsiones a largo y corto plazo, pasando por la selección de personal, círculos de calidad y un largo etcétera.

4 La subjetividad en el expertizaje

Se ha repetido hasta la saciedad que fundamentar los estudios económicos y de gestión en la opinión de un experto, por muy cualificado que sea comporta, en el mejor de los casos, la asunción de un alto riesgo que difícilmente podrá ser asumido por quienes tienen la responsabilidad de la gestión en empresas e instituciones. A pesar de ello, en una

primera fase nuestros esfuerzos se encaminaron a recoger, y en lo posible ampliar, los elementos lógicos y matemáticos que pudieran ser utilizados, previamente instrumentalizados, para el tratamiento de los problemas en un mundo de incertidumbre prácticamente imposible de probabilizar.

A pesar de que los acontecimientos se suceden, se superponen, nacen y mueren y nos traen junto con el temor la alegría de algunos éxitos, alegrías más raras que las inquietudes, es necesario realizar previsiones de la mejor manera posible, y ello, con los conocimientos y las informaciones que casi siempre son subjetivas, deficientemente transmitidas o deformadas de manera más o menos voluntaria. Existen, evidentemente, verdades que son absolutamente objetivas y que negarlo rozaría la insensatez, pero hay muchísimas que sólo son subjetivas, sobre todo las que se refieren al futuro próximo, a medio plazo, y, con mayor razón, a largo plazo.

Las empresas y las instituciones, se hallan rodeadas de consejeros, de expertos, de especialistas, de personas que se consideran con sentido común. Sin embargo, las opiniones individuales no son suficientes y, aún resultando cada vez más complicado es conveniente realizar deducciones hacia uno o varios objetivos elegidos y adaptados.

Y como se han hallado algunos buenos modelos matemáticos, han sido empleados para escapar de un exceso de desorden conceptual. Es necesario que exista un cierto desorden, ya que sin él no poseeríamos el don preciado de la imaginación. También es necesario el orden, pero dentro de unos ciertos límites. Demasiado orden convertiría al hombre en un robot. La playa de entropía, en donde se sitúa el pensamiento y la decisión entre el exceso de desorden y el exceso de orden, constituye una concepción muy lógica y portadora de los más adecuados frutos para una existencia equilibrada y, en definitiva, feliz.

Ahora bien, una opinión de experto no es una decisión, ni siquiera la preparación de decisiones. Esta opinión (u opiniones) se halla en un estrato superior y puede servir de base a las elecciones que deberán ser tomadas en el futuro. Y, como consecuencia de que en economía cada vez más se debe recurrir a la opinión de los expertos, al ser muy difícil y a

veces imposible realizar medidas, se escogen como informaciones básicas las opiniones de personas cualificadas, a veces especialistas, a veces personas que poseen unos conocimientos en numerosos ámbitos.

La subjetividad hace referencia a un individuo o a unos pocos. Este tipo de conocimiento constituye también un modelo (el cerebro humano está construido para modelizar), pero comparado con el de los otros puede resultar diferente e incluso muy diferente. En esto consiste nuestra libertad de pensamiento, esta valiosísima libertad de pensamiento.

5 La agregación de opiniones

En un segundo estadio hemos manejado, de la mejor manera posible, la opinión de los expertos, subjetiva y en algunas ocasiones objetiva, a través de ciertos métodos que han permitido un nuevo camino para el tratamiento de las informaciones necesarias con objeto de preparar una decisión. Manipular adecuadamente la intuición, nuestra intuición y la de los demás, en esquemas correctos ha sido una de las bases para el desarrollo de una economía de la incertidumbre.

Se acostumbra a tomar como punto de partida de todo expertizaje la comparación, a través del conocimiento abstracto o concreto. Para casos importantes, se recurre a varios expertos, lo que obliga a realizar comparaciones, agregaciones, a saber si se aceptan o rechazan las opiniones, a establecer contraexpertizajes, y de manera más general, a buscar una verdad que en el ámbito de la economía y de la gestión de empresas no se puede, o raramente se puede, medir. Para ello las matemáticas pueden aportar, en este campo en el que reina la subjetividad, una ayuda muy interesante, pudiéndose esbozar una teoría matemática del expertizaje. Esta teoría, que quizás hoy es sólo un esbozo [4], se halla asentada principalmente en conocimientos situados en ciertos ámbitos de las matemáticas formales (teoría de conjuntos) y de las matemáticas borrosas (subconjuntos borrosos, intervalos de confianza, expertones, etc.).

Hemos intentado agregar subjetividades [5] de tal manera que puedan llegar a ser aceptadas por un número suficiente de personas y así alcanzar el mayor nivel posible de objetividad.

Saber sacar fruto de manera adecuada a la opinión de los expertos no resulta tan fácil como

podría suponerse a primera vista. El abuso de las medias y otras características de la estadística puede conducirnos a decisiones incorrectas, a sofismos. La manera de agregar las opiniones próximas o alejadas, la manera de tratar esta incertidumbre subjetiva para que aparezca una cuasicerteza aceptada, constituye uno de los planteamientos más interesantes para la economía y gestión de empresas.

Los estudios realizados no pretenden ser exhaustivos (¿cómo sería esto posible!). Sin embargo, ciertas "recetas" para la agregación de opiniones de expertos están siendo, así lo esperamos, de mucha utilidad para los economistas.

6 Conclusiones

De los trabajos realizados han surgido nuevos elementos que se sitúan en cuatro esferas distintas del conocimiento: lógica, matemática, investigación operativa y economía y gestión de empresas.

En cada uno de estos niveles será necesario mejorar y tratar con más amplitud y exhaustividad los temas abordados, dentro de algunos años, cuando sean aceptadas nuevas condiciones y existan unas costumbres distintas de las actuales. La informática y la comunicación progresan a pasos agigantados y transforman la sociedad industrial y comercial, e incluso van a modificar, no lo dudemos, el funcionamiento de nuestras democracias.

Cada vez más se adquiere consciencia, a través de las informaciones vehiculizadas por los distintos medios, del papel más importante que juegan los expertos. Se busca la objetividad, lo cual es una actitud digna de elogio; la prueba científica es la que mejor se acepta como deseable por la mayoría, en una sociedad libre y desarrollada. Pero muchas de las situaciones que exigen una decisión en el ámbito de la economía no son susceptibles de ser conocidas con la suficiente objetividad y son objeto de controversia. El papel de los expertos consiste en aportar sus conocimientos y colaborar en la búsqueda de la verdad. Si existe coincidencia en todos los expertos nos acercamos a la objetividad, en caso contrario será necesario hallar conclusiones, sea a través del compromiso, sea por dominio evidente.

No hemos intentado aislar y controlar los principales problemas económicos, sino

únicamente suministrar algunos instrumentos útiles, algunos de los cuales poco conocidos. Se han planteado nuevas perspectivas sobre la manera de cuantificar la fenomenología de las instituciones y empresas cuando ello ha sido posible.

Quedan aún muchas cuestiones a tratar, más aún si tenemos en cuenta que el expertizaje y sobre todo el multiexpertizaje se convertirá en una práctica económica y empresarial cada vez más generalizada. es bien cierto que una persona puede tener razón frente a la opinión de muchos otros, sin embargo será necesario que se intente comprender la causa de este fenómeno y buscar el consenso que se asemeje más a la verdad que se resiste a ser aceptada.

En el ámbito de la economía no existen panaceas, pero, en cambio, es posible a través de instrumentos, conseguir una mayor eficacia en los conocimientos, sobre todo en las ciencias humanas, mucho más inciertas que las ciencias que permiten la prueba a través de experiencias repetidas y extensivas.

El expertizaje realizado con la ayuda de las nuevas técnicas operativas de gestión debe completar lo que puede ser efectivamente probado.

El pensamiento humano es muy sutil e incluso versátil, tenemos hoy oportunidad de hacer uso de la libertad de opinión. Esta libertad es para nosotros lo más querido..., y lo más difícil de conservar.

Referencias

- [1] BARRE, Raymond: Prólogo a la obra de Kaufmann, A. y Gil Aluja, J.: Técnicas operativas de gestión para el tratamiento de la incertidumbre. Ed. Hispano Europea, Barcelona, 1987.
- [2] KAUFMANN, A. y GIL ALUJA, J.: Introducción de la teoría de los subconjuntos borrosos a la gestión de las empresas. Ed. Milladoiro. Santiago de Compostela, 1986.
- [3] KAUFMANN, A. y GIL ALUJA, J.: Técnicas operativas de gestión para el tratamiento de la incertidumbre. Ed. Hispano-Europea. Barcelona, 1987.
- [4] KAUFMANN, A. y GIL ALUJA, J.: Técnicas especiales para la gestión de expertos. Ed. Milladoiro. Santiago de Compostela, 1993.
- [5] KAUFMANN, A. y GIL ALUJA, J.: Técnicas de gestión de empresa. Previsiones, decisiones y estrategias. Ed. Pirámide. Madrid, 1992.

Regulación basada en reglas y regulación P.I.D. Diferencias y complementaridades

Joseph Aguilar Martín^{1,2}

(¹)LAAS-CNRS, Toulouse, FRANCE

Joseba Quevedo Casín²

(²)Dept ESAII-ETSEIT-UPC, Terrassa
(Catalunya) SPAIN

Resumen

El conocimiento sobre un proceso industrial es relativo a su estructura y a su funcionamiento y se puede considerar como la anticipación de los efectos que producen las acciones de control.

La regulación es el campo de la Automática que pretende imponer una respuesta deseada a un sistema físico, a pesar de las perturbaciones e incertezas a que esté sometido. Tanto la técnica como el sentido común han impuesto las soluciones mas clásicas a este problema, entre las cuales la mas corriente es el llamado regulador P.I.D.

Recientemente el desarrollo de la Inteligencia Artificial ha permitido replantearse el control como un problema de razonamiento sobre el comportamiento de un sistema en funcion de los conocimientos globales que se tienen de él, i no ya únicamente de un modelo establecido en forma de ecuaciones diferenciales o recurrentes. La regulación Fuzzy es la solución mas utilizada de este nuevo tipo de control basado en reglas.

Nos proponemos aqui de analizar las diferencias y similitudes que hacen de la regulación basada en reglas una herramienta que puede remplazar los reguladores clásicos con variadas prestaciones. Mostraremos las componentes básicas para el diseño de reguladores fuzzy, y finalmente evocaremos la importancia del planteo borroso en la supervision de procesos industriales.

Palabras clave : Fuzzy regulation, PID, Knowledge Based control.

1 - Principio de la regulación

Desde el regulador de Watt hasta los más complejos sistemas de control actuales, la idea directriz es siempre la misma: contrarestar los efectos y comportamientos indeseables de un sistema dinámico por medio de acciones adecuadas. Dicha afirmación presupone que se conocen los efectos de un conjunto de acciones posibles, es decir la respuesta del sistema, y que, por otra parte se saben cuales son los comportamientos deseados, y en cierto modo los limites que no deben ser franqueados.

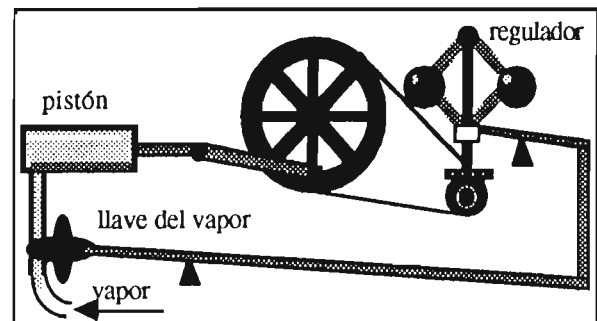


Figura 1

En la figura 1 se puede ver que abriendo el paso del vapor se acelera la máquina de Watt, y cerrandolo se disminuye su velocidad; basado en este conocimiento, James Watt ideó un mecanismo que, en el caso en que la velocidad de la máquina aumentase por cualquier razon externa, cierra parcialmente el paso del vapor, e inversamente lo abre en el caso de una reducción intempestiva de la velocidad [Winton 1885].

La construcción mecánica de este dispositivo está basada en dos mecanismos : el primero es la "medida" de la velocidad, lo cual se obtiene gracias a la fuerza centrífuga que separa mas o menos las bolas del famoso regulador, el segundo es la transmisión de esta información y la acción en la llave de paso del vapor.

La acción resultante ha de hacerse en el mismo sentido que el razonamiento anterior, es decir cerrar si se acelera y abrir si se decelera; además la obertura ha de hacerse progresivamente y ser adecuada al efecto que se quiere producir, es decir pequeña si la variación de la velocidad ha sido pequeña, y grande si esta última ha sido grande. Hay que notar que la obertura del vapor tiene dos límites bien claros : no se puede abrir mas del diametro de la tubería, ni se puede cerrar más que totalmente. La medida de la velocidad tambien tiene sus límites naturales, el ángulo que hacen los brazos que sostienen las bolas con el eje vertical va unicamente de nulo a ángulo recto : en el caso en que la velocidad está por debajo de cierto límite, las bolas no se separan del eje vertical y el ángulo es nulo, y a partir de cierta velocidad el peso de las bolas es insuficiente para influir de manera apreciable en el ángulo de los brazos.

2 - Síntesis del regulador

A partir de las condiciones claramente expuestas aquí en lenguaje natural, hay que confesar que el regulador no puede construirse, ya que quedan muchas cosas por decidir, todas ellas referentes a la relación entre la abertura y la velocidad dentro de los límites.

James Watt estaba reducido a mecanismos mecánicos y ello le limitaba la elección de esta relación. Por medio de palancas y poleas o engranajes sólo podía realizar funciones proporcionales al desplazamiento del cojinete inferior móvil que sostiene las bolas. La constante de proporcionalidad viene dada por la geometría de las articulaciones y de la transmisión. No existía posibilidad de imaginar que la manera en que la velocidad varía tambien podría ser utilizada en la regulación, si bien se concibió que a pesar de este regulador, podía quedar una diferencia entre la velocidad deseada y la velocidad obtenida, si se observaba una persistencia de dicho error se podía compensar mediante una acción manual que accionaba sobre las palancas de transmisión.

Esta descripción puramente lingüística de la regulación demuestra que los conceptos elementales de la automática no son fórmulas esotéricas de matemáticos, sino mas bien, la formalización del sentido común. El regulador de Watt efectúa un **control proporcional**, que mas tarde se convirtió en integral cuando se introdujo la acción correctiva del operador que constataba la existencia de un error remanente.

Hoy en día el mecanismo de la medida de la velocidad puede hacerse de un sinfín de maneras ópticas, mecánicas, electicas, etc que dan una señal numérica mas o menos fiable segun el precio del aparato, también la transmisión de esta información al actuador ha ganado en grados de libertad ya que prácticamemnte cualquier relación instantánea o dinámica y cumulativa puede realizarse mediante la micro-electrónica adecuada. Sin embargo no sabemos analizar ni prever el funcionamiento que se podría obtener con cualquier función; sólo en el caso lineal, y algunos otros casos elementales, se tienen principios básicos para la síntesis de funciones debidamente adecuadas y que proporcionen funcionamientos deseados.

3 - La regulación P. I. D.

Una vez mas el sentido común nos ayuda a pensar que si tuviéramos una idea de la manera en que varía la velocidad de nuestra máquina de vapor, podríamos prever si la acción de compensación ha de ser enérgica o leve : **control derivado**.

También podemos automatizar la constatación de un error demasiado permanente por medio de la acumulación del error instantáneo : **control integral**.

La acción de control efectiva será siempre del mismo tipo, es decir abrir más o menos la llave de paso del vapor, sin embargo, gracias a las facilidades de la técnica moderna podremos hacerla, no solo en funcion de la velocidad actual, sino que también dependerá de la derivada y de su integral : **control P.I.D. (proporcional + derivado + integral)**.

El control PID está, como hemos visto basado en el sentido común, y presupone un conocimiento cualitativo del sistema a regular. Sin embargo quedan tres parámetros a ajustar : las ponderaciones relativas de cada uno de sus componentes. La mayor parte de la teoría sobre este tipo de regulación se basa en la evaluación de

esos parámetros a partir de modelos matemáticos ideales del sistema a controlar. Cuando estos modelos son conocidos, y además lineales, el ajuste es fácil, y se pueden demostrar y anticipar las prestaciones del sistema controlado; cuando estos modelos no son conocidos, se trata de construir modelos lineales cuyo funcionamiento tenga cierta similitud con el del sistema, ya sea mediante estimación de parámetros, o por análisis de respuestas a experimentos, tras lo cual se aplican los ajustes como en el caso anterior.

4 - Regulación basada en límites.

En ninguno de los casos citados de regulación PID no se utiliza la información sobre los límites en las acciones ni en las medidas, se trata de una regulación atractiva, que está basada únicamente en la distancia hacia el punto descrito por los valores que se desean (error nulo, error permanente nulo, variación de la velocidad nula)

Por analogía con el caminante que presta mas atención en alejarse del precipicio que en seguir una trayectoria especial, podemos imaginar una regulación repulsiva en la cual, en vez de indicar una trayectoria a seguir, se indicasen trayectorias o conjuntos de valores de los cuales es preciso alejarse.

La teoría del control repulsivo no existe todavía, y la regulación tal como la conocemos toma sus bases en el análisis de la estabilidad, o sea la atracción causada por un punto o trayectoria de equilibrio.

Los límites de la acción de control son imperativos y corresponden a imposibilidades físicas. En cambio los límites de la variable controlada son contingentes a la precisión de su medida, o sea al funcionamiento deseado. En el caso de la máquina de vapor, si la velocidad está, por ejemplo, entre 80 y 100 rpm, nos daremos por satisfechos; en tal situación, contrariamente a la regulación proporcional clásica, no se reclama ninguna acción correctiva.

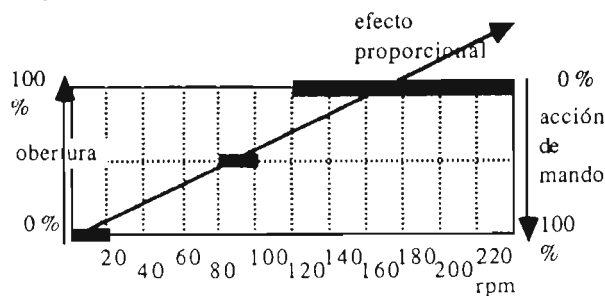


Figura 2

El diseño del coeficiente de transmisión entre el eje principal y el eje del regulador determina la precisión y la zona de validez de la medida. Supongamos que ésta va de 10 rpm a 220 rpm, como se vé en la figura 2, podremos entonces construir reglas como por ejemplo :

Cada vez que la velocidad se acerque peligrosamente a la velocidad máxima admisible, cerrar al máximo la admisión del vapor.

La acción inducida por esta regla nos alejará lo mas deprisa posible de los límites inaceptables. Esta regla, dictada por el sentido común y el conocimiento de la máquina no corresponde a ninguna regulación de tipo atractivo, y por lo tanto no tiene su equivalente en la regulación PID. Tampoco se considera en los reguladores tradicionales la zona de inacción debida a la satisfacción del objetivo, como se observa en la figura 2.

5 - Planteos de la regulación.

Mostraremos diversos planteos del problema de regulación, en el cual aparecerá la dificultad matemática del planteo de la regulación repulsiva mediante los medios clásicos del control y la optimización, comparada con otros conceptos de regulación.

5 -1- Regulación atractiva LQR

Sistema descrito por una función de transferencia :

$$y = \frac{N(s)}{D(s)} (u + \epsilon) + \eta$$

donde "u" es el control, "y" la respuesta, se han añadido dos variables aleatorias : "ε" la perturbación en la acción y "η" el error de observación, ambas supuestos de media nula.

Los índices siguientes caracterizan el objetivo de la regulación :

- calidad de la regulación: $V_y = \int_0^\infty y^2 dt$

- esfuerzo de control: $V_u = \int_0^\infty u^2 dt$

- importancia relativa del esfuerzo respecto a la calidad $\rho = \frac{V_y}{V_u}$

El planteo del problema consiste entonces en minimizar el criterio $V = V_y + V_u$

La solución óptima clásica conocida por el nombre de LQG (Lineal Cuadrática Gaussiana) es la siguiente :

Representación de estado equivalente

$$\frac{dx}{dt} = A.x + B.u + B.e$$

$$y = C.x + \eta$$

Si la ecuación algebraica de Riccati

$$A^T P + P A - \frac{1}{\rho} P B B^T P + C^T C = 0$$

tiene P^* por solución definida positiva, entonces el regulador óptimo viene dado por

$$u = -\frac{1}{\rho} B^T P^* x, \text{ lo cual supone que el estado } x \text{ es}$$

accesible o reconstructible. Si suponemos que las perturbaciones aleatorias son gaussianas y sus potencias respectivas sobre la acción y la medida son σ y ψ , podemos reconstruir el estimador de x mediante un filtro de Kalman-Bucy.

Si la ecuación algebraica de Riccati

$$F A^T + A F + \frac{1}{\psi} F C^T C F + \sigma B B^T = 0.$$

tiene F^* por solución definida positiva, entonces el estimador óptimo del estado x viene dado por :

$$\frac{d\hat{x}}{dt} = A.\hat{x} + B.u + \frac{1}{\psi} F^* C^T (y - C \hat{x}) \text{ y el}$$

$$\text{regulador es entonces } u = -\frac{1}{\rho} B^T P^* \hat{x},$$

5 -2- Regulación repulsiva:

Consideremos el mismo sistema, en el cual se precisan los límites admisibles de la respuesta y de la acción, o sean los intervalos admisibles para estas, como se ilustra para "y" en la figura 3 :

$$y \in [\underline{y}, \bar{y}] , u \in [\underline{u}, \bar{u}] .$$

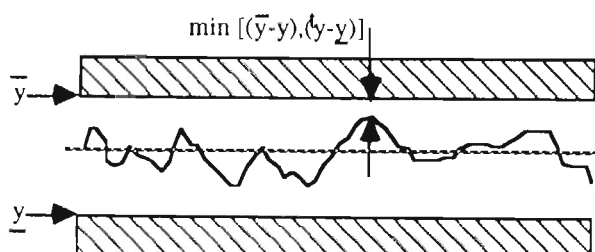


figura 3

La calidad de la regulación viene ahora dada en función de la distancia a los límites :

$$V_y = \min [(y - \underline{y}), (\bar{y} - y)]$$

Igualmente la calidad del esfuerzo de control es

$$V_u = \min [(u - \underline{u}), (\bar{u} - u)]$$

Y la importancia relativa del esfuerzo respecto al

$$\text{resultado es } \rho = \frac{V_y}{V_u} .$$

5 - 3- Control PID

Si se quiere evitar que un error permanente se instale, sabemos que basta con introducir una acción integral.

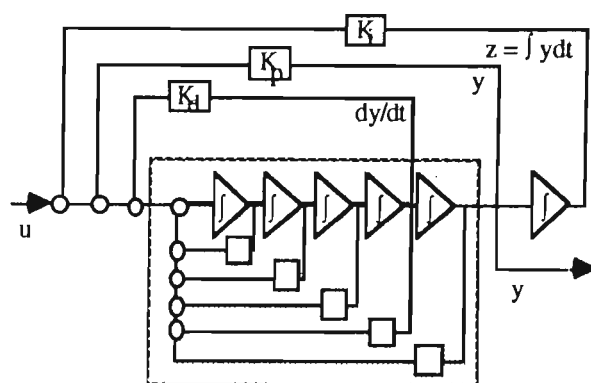


Figura 4

La figura 4 muestra como el regulador PID, en un sistema descrito en forma canónica, equivale a añadir un integrador, y por lo tanto a aumentar la dimensión del espacio de estado. Se puede plantear entonces el problema como anteriormente bajo la forma LQG, o sea como una optimización basada en el sistema aumentado, con el criterio :

$$V = \int_{t_0}^{\infty} (\alpha.y^2 + \rho.u^2 + \lambda.z^2).dt$$

$$\text{con } z = \int_{t_0}^t y . dt$$

Con este planteo se obtiene una familia de reguladores PID parametrizada por α , ρ y λ , con lo cual se puede dar más o menos importancia relativa a la calidad del control, al error permanente, o al esfuerzo de acción. Para obtener funcionamientos más o menos alejados de los límites admisibles se pueden ajustar estos parámetros, así se trata de resolver indirectamente el control repulsivo.

6 - La regulación Fuzzy.

Un controlador basado en lógica borrosa [Quevedo 1993] tiene 4 componentes básicas :

- la interfase de borrosificación o fuzzificador que deberá traducir el error numérico de regulación, y sus derivadas, y/o integrales, en valores lingüísticos o conjuntos borrosos.

- la base de conocimiento, formada por una base de datos que recoge la definición de las funciones de pertenencia de las entradas y salidas del controlador y la base de reglas, IF...THEN..., donde el diseñador del controlador describe las diferentes acciones a realizar en función del estado del proceso.

- el motor de inferencia, que evalúa todas las reglas para inferir las acciones de control borroso.

- la interface de desborrosificación o de-fuzzificador que deberá traducir las acciones borrosas en acciones susceptibles de actuar sobre el proceso.

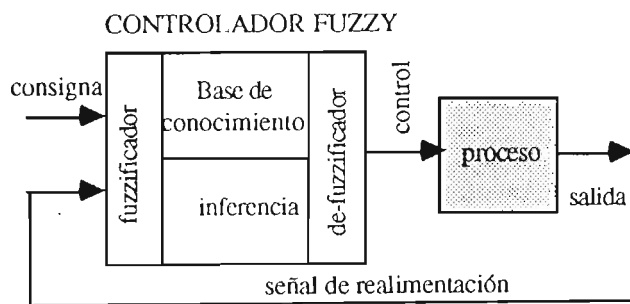


Figura 5

En la figura 5 se muestra el esquema del anillo cerrado por un controlador borroso o fuzzy.

El diseñador de un controlador borroso necesita definir un conjunto de parámetros para su correcto funcionamiento, los cuales pueden agruparse [Martínez 1994] en :

Parámetros estructurales :

- número de variables de salida
- número de variables de entrada
- número de entradas derivadas, acumuladas, etc.
- variables lingüísticas
- estructura de la base de reglas
- operaciones de lógica borrosa.
- procedimiento de inferencia borrosa
- método de de-fuzzificación.

Sintonía de parámetros :

- rango de los valores de las variables
- parámetros de las funciones de pertenencia

- ganancias estáticas de las entradas y salidas
- construcción de la base de reglas
- introducir nuevas medidas y refinar la base de reglas, etc.

Muchos de los interrogantes presentados quedan pre-definidos a priori por el tipo de herramienta que se utiliza, por ejemplo: máximo 7 variables lingüísticas por cada entrada o salida en el autómata programable SYS200H de OMRON, o simplemente por la experiencia del diseñador en otras aplicaciones similares.

Uno de los aspectos más importantes del diseño de controladores borrosos es cómo construir la base de reglas. Una alternativa es ajustar por prueba-error y comparar sus resultados con los obtenidos con controladores convencionales [Tang 1997] otra alternativa es construir un comportamiento similar al obtenido con un controlador PID lineal, ampliamente utilizado en el control de procesos, y después mejorarlo en línea con el proceso mediante prueba-error [Albertos 1992] y otra alternativa es la adquisición automática del comportamiento humano en el control manual de procesos industriales en el trabajo presentado en estas jornadas [Hernández 1994].

En este último caso, existe una metodología propuesta en [Wang 1992], para generar la base de reglas mediante aprendizaje con ejemplos. Este método puede resumirse en 5 pasos :

- 1) División de los espacios de las variables de entrada y salida en un número impar ($2n+1$) de conjuntos borrosos, con funciones de pertenencias triangulares solapadas.

- 2) Generación de reglas borrosas IF...THEN... a partir de la máxima pertenencia de los datos de entrada-salida adquiridos en el control manual del proceso.

- 3) Asignar un grado de confianza a las reglas borrosas generadas, en caso de conflicto con igual antecedente y consecuente diferente, teniendo en cuenta el producto de la pertenencia de los datos de entrada-salida a los conjuntos borrosos y la información a priori del experto.

- 4) Creación de una base de reglas borrosas, teniendo en cuenta el conocimiento obtenido en el segundo paso y la depuración del tercero.

5) Utilización de una estrategia de defuzzificación utilizando la operación producto de la función de pertenencia de los antecedentes de una regla para obtener el de la salida de control, y posteriormente, se aplica la fórmula del centroide para todas las reglas activas simultáneamente.

Una metodología semejante a esta [Hernandez 1994] ha sido utilizada en nuestro Laboratorio docente para controlar con un autómata programable comercial el nivel de 2 depósitos interconectados y el posicionamiento de un móvil en una guía con resultados muy satisfactorios.

7 - La supervisión de procesos.

El problema del control y de la supervisión de procesos industriales complejos abarca los aspectos funcionales importantes que se articulan según la figura 4 :

- función regulación : estabilización de variables como temperatura, presión, caudal, velocidad,... en unos valores de referencia óptimos de funcionamiento
- función control : seguimiento de referencias dinámicas como trayectorias de referencia, rampas, .. por parte de ciertas variables, posición, temperatura,... de los procesos.
- función supervisión : sistemas de alarma y de protección automática del correcto funcionamiento de los procesos, con posible toma de decisiones sobre la reconfiguración del sistema de control, y/o acciones redundantes.
- función gestión : La optimización, económica, energética, de seguridad,... de los diferentes sub-procesos sujetos a ciertas restricciones físicas y la coordinación de estos, así como su coordinación y la gestión integral de todo el proceso.

En la figura 4 aparecen todos estos aspectos funcionales reflejados en forma de niveles interrelacionados. Los niveles bajos de regulación y control reaccionan rápidamente, actúan de forma local y el conocimiento analítico del proceso a este nivel es, muchas veces, más que suficiente para diseñar controladores convencionales, PID,... que reaclicen adecuadamente estas funciones. Los niveles altos de supervisión y gestión óptima reaccionan más lentamente y actúan de forma global, siendo más difícil obtener modelos globales de comportamiento y las soluciones analíticas son prácticamente inabordables. En estos niveles el razonamiento heurístico o aproximado de los

expertos permite obtener soluciones eficaces a los problemas planteados.

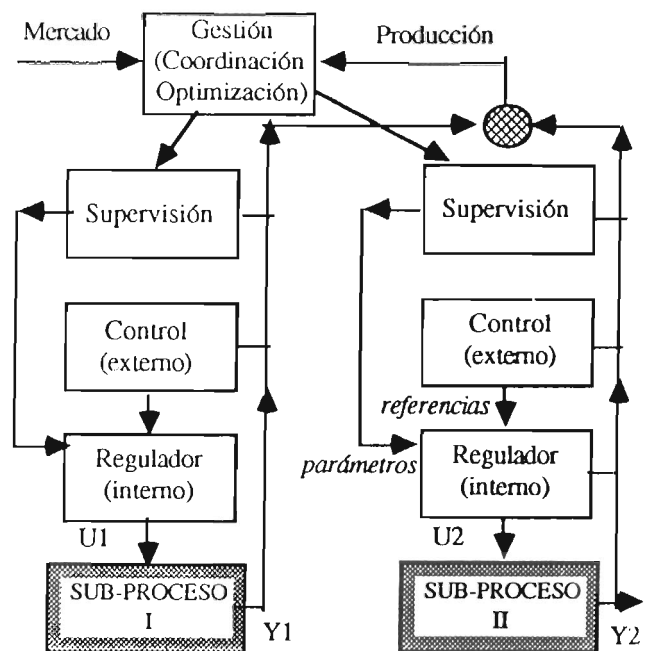


figura 4

Así, en estos últimos niveles, e incluso en los niveles bajos, cuando el comportamiento del proceso es muy no-lineal o impreciso, la teoría de la lógica borrosa proporciona una opción seria para diseñar el sistema de control.

8 - Conclusión

En este artículo hemos querido situar la regulación fuzzy con respecto a los métodos tradicionales fruto de la teoría del control. Visto de un modo abstracto, un regulador fuzzy es una de las posibles realizaciones de reguladores basados en reglas, sin embargo hay que reconocer que actualmente es el único que disfruta de un amplio conjunto de herramientas y métodos para diseñarlo.

Además la lógica borrosa, de la cual es una aplicación, posee hoy un gran número de resultados gracias a los cuales se pueden defender teóricamente la mayoría de los esquemas de regulación fuzzy, aunque muchos de ellos hayan sido propuestos al origen, de manera heurística y justificados sólo experimentalmente.

Referencias.

J. G. Winton, *Modern steam practice and engineering*, libro, ed Blackie & Son, Old Bailey E. C., 1885

L.X. Wang & J.M. Mendel, *Generating fuzzy rules by learning from examples*, IEEE Trans. SMC, vol. 22 n. 6, pp. 1414-1427, 1992

P. Albertos, *Fuzzy controllers. Applications of Artificial Intelligence in Process Control*, Ed. Pergamon Press ISBN 0-08-042016-8, 1992.

M. Martinez, P. Albertos & J.L. Navarro, *Fuzzy Control Methodology in the Process Control*, Preprints of SICICA symposium, pp. 37-47, Budapest 1994.

K. Tang & R. Mulholland, *Comparing Fuzzy Logic with classical controller designs*, IEEE Trans on SMC, vol. 17 n. 6, pp. 1085-1087, 1987

J.C. Hernandez, F. Jimenez & J. Quevedo, *Experiencias con un autómata programable fuzzy sobre una planta piloto*. CONAI 1994, La Habana, Cuba 1994.

J. Quevedo & A. González, *Confidence degree on the advanced stand alone controllers*, Preprints of SICICA symposium, pp. 422-426, Budapest 1994.

J. Quevedo, *Aplicaciones de la lógica borrosa en un laboratorio de Automática*, revista ARBOR, CXLVI, n. 573-574, pp. 221-229, 1993

Neural Distributed Processing Systems

A. F. Rocha

Dep. Physiology and Biophysics
Dep. Computer Engineering and Automation
Dep. Petroelum Engineering
UNICAMP
e-mail: eina@bruc.bitnet

1 Introduction

Distributed Intelligent Problem Solving Systems (henceforth, called DIPS) are defined as a collection of agents (intelligent or not) which interact to solve a problem in a defined knowledge domain (Chandrasekaran, 1981; Durfee and Montgomery, 1991; Hewitt and Inman, 1981; Lesser, 1991). According to this approach, the solution of the problem may be decomposed in a set of solutions of sub-tasks (problems), and each DIPS agent is assigned to solve one of these sub-tasks. The sub-tasks may be of any sort, e.g., data base activities, microtheory proving, efferent device control, data acquisition, etc. Since sub-tasks are supposed not necessarily require intelligent solutions, this property arises mainly by means of the way agents interact to solve the whole problem. Basically, agents ask and furnish information from and to other agents. Information flow may be specified by means of either mail systems when at least one of the agents knows precisely the address of the other; or by means of message broadcasting, when agents select messages they are interested based on their contents and these messages are distributed in the system by means of unspecific systems like blackboards, buses, etc. Implicit on both strategies is that communication is supported by a set of specific or general relations. In both cases, some sort of connectionism is assumed. One popular way of specifying DIPS information flow is contract nets. However, agents are not prohibited to use logic to take care of its obligations concerning the whole DIPS.

The knowledge accumulated about the biochemistry of the synapsis in the last two decades, because it completely changes the notion of brain processing from an exclusive electrical mechanism toward a complex message exchange occurring both locally at the synaptic site as well as at distances depending on the solubility of chemical substances in the extracellular environment. Synaptic transactions cannot be assumed anymore as simple process of propagating numbers powered by a factor measuring the pre-synaptic capacity of influencing the post-synaptic electrical activity. Briefly, the electrical activity arriving at the pre-synaptic terminal releases specific molecules called transmitters (t) into the extracellular space in amounts proportional to the degree of this incoming

activity. These transmitters have to binding to specific post-synaptic molecules called receptors (r) in order to modify the post-synaptic electrical and/or chemical activity. These actions are exerted by other special molecules called generically controllers (c) or second messengers. In other words (Rocha, 1992):

$$t + r \longrightarrow c \simeq \text{action} \quad (1)$$

the concatenation of the chemical strings t and r produces another string c whose semantics is defined by its action over the neuronal economy. A special kind of controller is the ionic gate, whose action is to modify the electrical properties of the post-synaptic membrane. The amount $a(c)$ of the activated controller is related to 1) the degree of pre-synaptic activity expressed by the amount $a(t)$ released transmitter; 2) the amount $a(r)$ of available post-synaptic receptor, and 3) the degree $\mu(t, r)$ of the t/r matching, that is:

$$a(c) = f(a(t), a(r), \mu(t, r)) \quad (2)$$

Because of this, it may be said the synaptic transactions involve both numeric and symbolic calculations; numbers measuring the degree of symbol truth.

The neuron defined by (1) and (2) may be considered as an computational agent able to exchange messages with other neurons. In this line of reasoning, the brain may be understood as distributed intelligent processing system (DIPS), reasoning by means of message exchange among their elements. The purpose of the present paper is to introduce and discuss these ideas.

2 The Formal Neuron

A new formal model of the neuron incorporating this new knowledge about the synaptic biochemistry may be obtained by combining the McCulloch-Pitts neurons and some notions borrowed from Petri Nets Theory (Rocha, 1992, Rocha et al., 1993). A Petri Net PN is as structure:

$$PN = (P, T, \theta, A, M) \quad (3)$$

where:

P is a set of places;

T is a family of tokens;

θ is a set of transitions;

A is a set of arcs joining places to transitions and transitions to places, and

M is a map assigning tokens to places.

Whenever a transition (e.g., the synapsis) is activated (e.g., t/r binding), it uses tokens from its input (pre-synaptic) places to produce tokens at its output (post-synaptic) places according to the rules specified at the transition (the t/r binding).

The amount $a(i, c)$ of controllers activated by the i th synapsis of a neuron n_j , is dependent on the:

a) amount $a(i, r)$ of available receptor r at the post-synaptic site;

b) amount $a(i, t)$ of the transmitter t released at the pre-synaptic terminal n_i , and

c) chemical affinity $\mu(t, r)$ among the tokens t and r .

Now:

d) if $A(i, t)$ is the total amount of transmitter t stored at the pre-synaptic site, and

e) if $A(i, r)$ is the total amount of the receptor r stored by the post-synaptic cell, then

f) the synaptic weight w_i is calculated as

$$w_i = A(i, t)\rho A(i, r)\lambda\mu(t, r) \quad (4)$$

where ρ and λ are, in general, t-norms.

In this condition, the actual amount $a(i, t)$ of transmitter t_i bound to the receptor r_i at the i -th synapsis may be calculated as

$$a(i, t) = s_i * (A(i, t)\rho A(i, r)\lambda\mu(t, r)) \quad (5)$$

where $*$ is, in general, a t-norm, and s_i is the pre-synaptic activity. The amount $a(i, c)$ of the controller activated by the t/r binding may be assumed to be:

$$a(i, c) = g(a(i, t)) \quad (6)$$

In the case, g is the identity function:

$$a(i, c) = s_i * w_i \quad (7)$$

But, if the controller c_i may activate the axon then the different post-synaptic places are associated to a post-synaptic axonic place by means of special transitions whose function is to simulate the McCulloch-Pitts Neuron, such that the amount of tokens s_j produced at the post-axonic place is obtained as

$$a = \sum_{i=1}^n a(i, c) \quad (8)$$

$$s_j = f(a) \quad (9)$$

In other words, s_j is a function f of the powered averaging a of the inputs s_i . If f is taken as a boolean function then:

$$s_j = \begin{cases} 1 & \text{if } a > \alpha \\ 0 & \text{otherwise} \end{cases} \quad (10)$$

Different transmitters t_i may be released by the different pre-synaptic neurons n_i to bind different post-synaptic receptors r_i and to produce different controllers c_i :

$$t_i \oplus r_i \longrightarrow c_i \quad (11)$$

meaning:

$$\text{If } t_i \text{ matches } r_i \text{ then } c_i \text{ is available} \quad (12)$$

which supports the following fuzzy production rule:

$$\text{If } t_i \text{ is } r_i \text{ then } a_j \text{ is } c_i \quad (13)$$

where a_j is the activation of the post-synaptic neuron n_j induced by the pre-synaptic neuron n_i . Because

$$s_j = f\left(\sum_{i=1}^n s_i * (A(i, t)\rho A(i, r)\lambda\mu(t, r))\right) \quad (14)$$

then n_j may be assumed to be able to handle a knowledge base of the type (Yager, 1990; Rocha, 1992):

$$\text{if } Q \text{ (relevant } X_i \text{ is } A_i \text{ are true) then } Y \text{ is } B \quad (15)$$

where Q is a linguistic quantifier like *Most of*, *If*, *X of N*, etc, the degree of relevance of X_i is A_i is given by w_i and the confidence X_i is A_i is given by s_i . Because of this, the artificial neuron for (4) to (7) may be considered a processing device supporting the Generalized Modus Ponens.

3 Sensory Neural Agents

Let the fuzzy variable X to be defined as the restriction R over the universe of discourse U . In other words:

$$R : X \times U \longrightarrow [0, 1] \quad (16)$$

$$x = R(u), \quad u \in U \quad (17)$$

Given a set T of linguistic terms and a set M defining the semantics of these terms, a linguistic variable X is defined by a family of n restrictions R_i over the universe of discourse, each one defining the semantics $m_i \in M$ of $t_i \in T$. In other words:

$$R_i : X \times T \times U \longrightarrow [0, 1] \quad (18)$$

$$x = \{t_i = R_i(u) \mid t_i \in T, u \in U \text{ and } i = 1, \dots, n\} \quad (19)$$

A sensory neuron n_j is defined if a in (1) is assumed to be a measurement in the universe of discourse U and the filtering function f is taken as the restriction R . In other words:

$$a = u, \quad s_j = f(a) \quad (20)$$

In this condition, n_j is said to be in charge of reading the fuzzy variable X in the universe of discourse U . Since different filtering functions may be assigned to the different axonic branches of the same neuron, n_j may also be programmed to read a linguistic variable X in U . In this case, each term $t_i \in T$ is associated to a different pre-synaptic token at the different axonic terminals. In other words:

$$s_j = \{t_i = R_i(u) \mid t_i \in T, u \in U \text{ and } i = 1, \dots, n\} \quad (21)$$

The neuron defined in (20) will be called here **fuzzy sensory neuron**, and that defined by (21) will be called **linguistic sensory neuron**. As a matter of fact, the fuzzy sensory neuron is a special case of the linguistic sensory neuron, when just one term (token) t_i is defined in T . Both the set of terms T and the semantic set M may be assumed to be specified by controller tokens released by other neurons, what means that sensory neurons may be specified to make specific readings of U according to the reasoning involved in the decision making.

Let P be a set of measures x_i^j provided by a family S of m sensory neurons at different instants $i = 1$ to n . In other words:

$$P = [x_i^1, \dots, x_i^m \mid i = 1, n] \quad (22)$$

is the trajectory (or pathway) in the space X^m related to the observation of U provided by S .

Different neural agents a_p may now be designated to analyze different types of relations shared by different points $p_j, p_k \in P$. For example, these agents may be in charge of classifying directions between consecutive points of P according to some preferential directions; of to detect changes of direction in P , etc. Let us call these agents P feature detecting agents.

The role of the agent a_p is to recode P into a string p of fuzzy symbols ($c_1, \dots, c_j$) according to the relations shared by points $p_j, p_k \in P$. Whenever a string p containing a symbolic description of events in P is produced, a vector v may be created containing the coordinates $x^1, \dots, x^m$ associated to each element of p . Different symbolic and vectorial descriptions p, v of P will be generated by the distinct agents a_p .

Hence, given a family A of feature detecting agents, a picture D of P is provided as family of partial descriptions p, v , whose cardinality n is determined by the number of its specialized agents:

$$D(P) = \{(p_1, v_1), (p_2, v_2), \dots, (p_n, v_n)\} \quad (23)$$

The role of any primitive $a_p \in A$ is to recode P into a string p of symbols ($c_1, \dots, c_j \in C$) according to the relation R shared by points $x_j^1, \dots, x_j^m$ and $x_k^1, \dots, x_k^m$ of P in X^m . Both R and C specialize a_p .

4 Reasoning Neural Agents

Sowa, 1984 called **percepts** the elementary features extracted by low level neural circuits (or primitive agents) and proposed **concrete concepts** to be

supported by sets of learned relations among percepts. **Abstract concepts** are, then, built as a set of learned relations among concepts. Here, the *word* concept will be used in a general sense meaning either percept, or concrete or abstract concept, unless stated otherwise.

Let C to be a set of concepts; R a set of possible relations to be established among these concepts and F a set of facts about the universe U . A theory or model T about U is:

$$T \subset C \times R \Leftarrow F \quad (24)$$

where $\subset$ means contained and $\Leftarrow$ denotes supported. Thus, T is a set R of relations among concepts C learned or proved to hold in U because supported by a set F of facts extracted from U . A high order concept is a theory in a preceeding low level of analysis.

Given a sensory system S providing a picture $D(P)$ of the fact F , a set C of concrete concepts is defined as:

$$T \subset D(P) \times R \Leftarrow F \quad (25)$$

where R depends on both the specializations of each primary element of a_p as well as on the commitments established among the different agents of A . Specialization is supported by selection of tokens and commitment among neural agents is dependent on the amount of tokens they have to support message exchange (see, eqs. for (4) to (7)). In this way, R may be learned by both selecting adequate tokens used by agents and by strengthening successful relations (commitments) among these agents. Specialization and commitment may be also inherited (programmed).

Higher order institutional agents H_i^l may learn (or be crafted to handle) complex models or theories T_i^l about U supported by those facts $F \subset \{T_i^{l-1}\}_{i=1 \text{ to } n}$, that is:

$$T_i^l \subset C \times R \Leftarrow F \subset \{T_i^{l-1}\}_{i=1 \text{ to } n} \quad (26)$$

This means, that a theory T_i^l at the hierarchical level l is supported by facts provided by theories T_i^{l-1} constructed at level $l-1$. H_i^l is called reasoning agent.

The bottom-up process described by (25) and (26) may now be reverted if T_i^l is supposed to be used to control a set of effector agents E_e in charge of acting upon U . In other words, models of action M_j^{l-1} may now be supported by a set of high order theories $\{T_i^{l-1}\}_{i=1 \text{ to } n}$ such that:

$$M_j^{l-1} \subset C \times R \Leftarrow F \subset \{T_i^{l-1}\}_{i=1 \text{ to } n} \quad (27)$$

$$A \subset C \times R \Leftarrow F \subset \{M_j^{l-1}\}_{j=1 \text{ to } n} \quad (28)$$

where A is the desired action to be exercised upon U . In general, this action A is supposed to modify some variables $\{X^k\}$ defined over U . This is because (28) may in some conditions to be rewritten as:

$$P \subset C \times R \Leftarrow F \subset \{M_j^{l-1} \mid j=1 \text{ to } p\} \quad (29)$$

A high order theory L^l may now be learned by some specialized institutional agents S_l , A_l , H_l^i and E_l to act upon P_l such that:

$$P_l \simeq L^l \quad (30)$$

In other words, P_l becomes a projection of L^l in U . In this way, abstract concepts may be introduced in U by the complex agent L^l :

$$L^l = \{S_l, A_l, H_l^i, E_l\} \quad (31)$$

In this condition, L^l may introduce in U at least some of the T_l^i learned by the DIPS it belongs to. Hence it may speak about the DIPS' concepts. If different DIPSs share the same theory L^l they may also learn to share common theories T_l^i . This high order theory T_l^i is called the language spoken by these DIPSs and its complexity is defined by the actual value of l .

5 Conclusion

The concepts discussed in the previous sections were used to implement the following systems:

a) **Jargon**(c) (Rocha et al, 1992; Rocha, 1992) is a hierarchical system composed by three neural nets devoted, respectively, to learn the most frequent words, phrases and sets of phrases occurring in a set of texts. It has been successfully used to extract knowledge from different data base written in natural language (e.g.; Arruda et al, 1994; Morooka, et al 1993; Rocha and Rocha, 1994)

b) **Sensor** (Rocha and Serapiao, 1994) is a distributed sensory system aimed to use both numeric and symbolic processing in pattern recognition. Primitive sensory agents S_r specialize in detecting elementary features (percepts) on the incoming sensory data, such as: relative directions between neighbours sensory point in the sensory space S ; trends on the sensory trajectory P in this space; periodicities in P , etc. Associative sensory agents A_c reason about sensory patterns using both the symbolic description of P provided by a family $\{S_r\}_{r=1 \text{ to } p}$ of primitive sensory agents and the actual numeric values of the sensory points associated to these percepts.

c) **SICAD** (Alegre, 1993) is a distributed intelligent system devoted to the management of oil production. It uses some of the facilities of **Sensor** and **fuzzy reasoning** to diagnose problems in rod pumping and to optimize oil production.

The results obtained so far with this different systems supports the conclusion that Fuzzy Neural Distributed Intelligent Systems constitute a very powerful AI tool.

6 References

Alegre, L., C.K. Morooka and A. F. Rocha - Intelligent Diagnosis of Rod Pumping Problems. SPE 26516, 1994

Arruda, H.; L. H. Arruda; A. F. Rocha and M. Theoto - Using Neural Networks to analyse Medical Concepts in Different Populations. Proceeding of IPMU, Paris, 1994

Chandrasekaran, B. - Natural and social system metaphors for distributed problem solving: introduction to the issue. IEEE Trans. Sys. Man and Cybernetics, 1981/11:1-5.

Durfee, E. H and T. A. Montgomery - Coordination as distributed search in a hierarchical behavior space. IEEE Trans. Sys. Man and Cybernetics, 1991/21:1363-1378.

Hewitt, C. and J. Inman - DAI betwixt and between: From "Intelligent Agents" to Open System Science. IEEE Trans. Sys. Man and Cybernetics, 1991/21:1409-1419.

Lesser, V. R. - A retrospective view of FA/C distributed problem solving. IEEE Trans. Sys. Man and Cybernetics, 1991/21:1347-1362.

Morooka, C. K., A. R. Rocha; K. Miura and L. Alegre - Offshore Well Completion Operational Knowledge Acquisition and Structuring.. OMAE, Vol. 1, Offshore Technology, ASME, 1993

Rocha, A.F. - Neural Nets: A theory for brains and machines. In Lecture Notes in Artificial Intelligence, Springer Verlag, 1992/638

Rocha, A. F. , M. Theoto, A. M. K Myiadahira, M. S. Koizumi and I. R. Guilherme - A neural net for extracting knowledge from natural language data base. IEEE Trans. Neural Nets. 1992, 2/5,115-123

Rocha, A. F., F. Gomide and R. Machado - Updating the biology of the artificial neuron. Fuzzy Logic: The State of Art, Lowens and Rouben (ed.), Kluwer Academic Publishers, 1993

Rocha, M. T. and A. F. Rocha - Using neural tools in environmental data analysis. To appear in the Proceedings of 3rd. Int. Conf. On Fuzzy Logic, Neural Nets, Chaos and Soft Computing, Iizuka, Japan, August 1994

Rocha, A. F., M. T. Rocha and I. R. Guilherme - Neural Nets and Concept Reasoning. Proceeding of EUFIT'94, Aachen, Germany 1994

Rocha, A. F. and A. Serapiao - Neuro-fuzzy Systems and Pattern Recognition. 1994. To appear.

Yager, R. R. On a semantics for neural networks based on linguistic quantifiers. Tech. report MII-1103 Machine Intelligence Institute, Iona College, N. Rochelle, N. York.

COMUNICACIONES

Lógica Fuzzy

Razonamiento *Modus Tollens* en Bases de Conocimiento Borroso con Encadenamiento

A. Bugarín, P. Félix y S. Barro

Departamento de Electrónica e Computación.

Facultade de Física. Universidade de Santiago de Compostela.

15706 Santiago de Compostela.

e-mail: senen@gaes.usc.es

Resumen

La ejecución de Bases de Conocimiento Borroso (BCB) con encadenamiento entre las reglas que las forman resulta frecuentemente un proceso que conlleva un elevado coste computacional. Por este motivo se han planteado mecanismos de ejecución dirigidos a la reducción del número de operaciones a realizar, ya sea mediante la descripción paramétrica de los valores lingüísticos que intervienen en el proceso [7] o mediante la realización en tiempo diferido (momento de definición de la BCB) de parte de las operaciones asociadas a la ejecución de la misma [10]. En este trabajo se analiza el proceso de razonamiento *Modus Tollens* [6, 8] para BCB con encadenamiento entre reglas. De las expresiones obtenidas se deduce la posibilidad de realizar una compactación a priori de la BCB que permite la realización en tiempo diferido de parte de las operaciones asociadas a dicho proceso de razonamiento. El mecanismo de encadenamiento que se describe es más flexible que el utilizado habitualmente [9, 4] en este tipo de BCB. Consideramos que una regla está encadenada con otra por una variable intermedia si existen dos proposiciones, una de ellas en la parte consecuente de la primera regla y la otra en la parte antecedente de la segunda regla, que se establecen sobre dicha variable. El mecanismo de inferencia *Modus Tollens* en BCB con reglas encadenadas a nivel de variables se plantea a partir del principio de combinación/proyección [11] y en dos situaciones diferentes, que dependen de la relación de implicación utilizada: en la primera situación se consideran relaciones de implicación que verifican la propiedad conmutativa y en la segunda relaciones de implicación no conmutativas. En el primero de los casos, el proceso de inferencia *Modus Tollens* resulta idéntico a la realización de un mecanismo de razonamiento abductivo. Puede aplicarse directamente, por tanto, el proceso de compactación descrito para el razonamiento *Modus Ponens* [1, 2, 3], sin más que plantear la regla correspondiente invirtiendo el sentido de la implicación. El grado de compactación en esta situación resulta óptimo, permitiendo realizar el cálculo del Grado de Verificación (GDV) numérico de cualquier proposición en que intervenga una variable intermedia mediante operaciones simples (cálculo del valor

máximo) entre pares de números reales (los GDV de las proposiciones que se encuentran encadenadas con ella y los valores numéricos de compactación, calculados a priori). El proceso de ejecución de la BCB puede entenderse, en este contexto, como la "propagación" a través de la misma de dichos GDV numéricos. Para el caso de relaciones de implicación no conmutativas se analiza el mecanismo de inferencia *Modus Tollens* en el dominio lingüístico de la verdad, de forma análoga a la descrita en la bibliografía [10, 7] para el mecanismo de inferencia *Modus Ponens*. El paso al dominio de la verdad permite, por un lado, el tratamiento homogéneo de cualquier variable, independientemente del universo de discurso en que se haya definido. Por otra parte, permite interpretar el proceso de inferencia como la composición de valores lingüísticos de verdad, algunos de los cuales son calculados a priori. En concreto, el GDV lingüístico de cualquier proposición con variable intermedia se obtiene mediante la composición del valor lingüístico de compactación, que se calcula en el momento de definición de la BCB, y la función que implementa el razonamiento *Modus Tollens* en el dominio de la verdad (equivalente en *Modus Tollens* de la Función de Razonamiento hacia Adelante [7] descrita para el razonamiento *Modus Ponens*). De esta manera, el proceso de inferencia puede entenderse como la propagación de los GDV lingüísticos a lo largo de la BCB. La carga computacional correspondiente a las operaciones relativas al encadenamiento entre reglas se traslada en gran medida al momento de definición de la BCB a través de los valores de compactación lingüísticos, que resumen la relación existente entre las reglas encadenadas, reduciendo notablemente el número de operaciones a realizar durante la ejecución de la BCB.

1 Introducción

El ámbito de los sistemas basados en reglas se ha mostrado como una de las áreas de mayor éxito en la aplicación de la lógica borrosa. Un elemento fundamental a la hora de describir y tratar este tipo de sistemas es la Base de Conocimiento Borroso (BCB), donde se recoge el conocimiento disponible acerca del sistema. Este conocimiento se estructura habitualmente en forma de reglas condicionales SI-ENTONCES, en

cada una de las cuales se ligán las variables de la parte antecedente y las de la parte consecuente a través de una relación de implicación. La creciente aplicación de la lógica borrosa al control de sistemas de gran complejidad, junto con su uso en sistemas de soporte y ayuda a la decisión y, más propiamente, en Sistemas Expertos Borrosos, exigen una estructuración altamente flexible de la BCB, que facilite la elicitación y representación del conocimiento. De esta forma resulta conveniente alterar la estructura homogénea que presentan las BCB en sistemas de control para incluir variables intermedias que permitan reflejar estados más complejos del sistema. Estas variables, presentes simultáneamente en la parte antecedente de algunas reglas y en la parte consecuente de otras, dan lugar a la existencia de encadenamiento entre reglas, imponiendo, por tanto, una ordenación implícita a la hora de realizarse el proceso de ejecución. Todo esto, añadido a la propia utilización de la teoría de conjuntos borrosos, trae aparejado un coste computacional notablemente alto para el proceso de ejecución de la BCB. Por ello, en ciertas aplicaciones resulta interesante adoptar estrategias de ejecución que optimicen dicho proceso, simplificando el cálculo en función de los objetivos perseguidos en cada momento.

En la bibliografía se describen al menos dos métodos que abordan esta simplificación del cálculo: la parametrización de los valores lingüísticos que intervienen en el proceso [7] y la utilización de valores de compactación [10, 1]. Ambos abordan el problema de ejecución de reglas y BCB según el esquema de razonamiento *Modus Ponens*. En este trabajo abordaremos el tratamiento del esquema *Modus Tollens* en BCB con encadenamiento, siguiendo la estrategia de compactación que permite realizar a priori parte de las operaciones involucradas.

2 Razonamiento Modus Tollens

El interés principal de abordar esta estrategia de razonamiento [6, 8] radica en la necesidad de inferir información acerca de las variables de la parte antecedente de una regla a partir de la información disponible sobre las variables de la parte consecuente. Para ello se realiza, al igual que en el esquema de razonamiento *Modus Ponens*, la composición "Sup-*" [11] de la relación de implicación con la distribución de posibilidad obtenida a partir del conocimiento disponible de las variables de la parte consecuente. Se verá cómo este proceso resulta análogo al correspondiente para el esquema *Modus Ponens*. Posteriormente lo encuadraremos en el contexto de BCB con encadenamiento entre reglas, y analizaremos el mecanismo de compactación de las reglas a través de la utilización de valores lingüísticos de compactación. Estos valores pueden ser calculados a priori y resumen en una distribución de posibilidad las operaciones relacionadas con las variables intermedias. A partir de la descripción analítica del proceso, que se muestra seguidamente, discutiremos sus distintos ámbitos de aplicación, en función del tipo de relación de implicación considerada.

2.1 Razonamiento Modus Tollens en una regla

Describiremos la regla genérica r -ésima R^r de una BCB como

$$\begin{aligned} SI \ X_1^r \ ES \ A_1^r \ Y \ ... \ Y \ X_{M_r}^r \ ES \ A_{M_r}^r \\ \text{ENTONCES} \\ X_{M_r+1}^r \ ES \ B_1^r \ Y \ ... \ Y \ X_{M_r+N_r}^r \ ES \ B_{N_r}^r \end{aligned} \quad (1)$$

$X_{m_r}^r, m_r=1, \dots, M_r$ y $X_{n_r}^r, n_r=1, \dots, N_r$, son variables lingüísticas, M_r, N_r el número de proposiciones, respectivamente, en la parte antecedente y consecuente de la regla, $A_{m_r}^r, m_r=1, \dots, M_r, B_{n_r}^r, n_r=1, \dots, N_r$, son valores lingüísticos asociados a las distribuciones de posibilidad de igual nombre

$$A_{m_r}^r = \{a_{m_r,i}^r\} \text{ y } B_{n_r}^r = \{b_{n_r,j}^r\}, \ i, j = 1, \dots, I, \quad (2)$$

donde I es el grado de discretización de los universos de discurso (uniforme para todos ellos). Además, denotaremos por

$$\tilde{B}_{n_r}^r = \{\tilde{b}_{n_r,j}^r\}, \ j = 1, \dots, I, \ n_r = 1, \dots, N_r \quad (3)$$

las distribuciones de posibilidad conocidas en un momento dado para las variables de la parte consecuente de la regla, y por

$$\tilde{A}_{m_r}^r = \{\tilde{a}_{m_r,i}^r\}, \ i = 1, \dots, I, \ m_r = 1, \dots, M_r \quad (4)$$

las distribuciones de posibilidad inferidas para cada una de las variables de la parte antecedente, tras aplicar el esquema de razonamiento *Modus Tollens*. Este esquema puede realizarse aplicando la Regla Composicional de Inferencia "Sup-*", de modo que a cada variable de la parte antecedente de la regla $R^r, X_k^r, k \in \{1, \dots, M_r\}$, el proceso de inferencia *Modus Tollens* le asigna una distribución de posibilidad $\tilde{A}_k^r = \{\tilde{a}_{k,i}^r\}, i = 1, \dots, I$ definida como

$$\begin{aligned} \tilde{a}_{k,i}^r = \bigvee_{\substack{i_{m_r} \\ m_r \neq k}} \bigvee_{\substack{j_{n_r} \\ n_r = 1, \dots, N_r}} \sigma \left(\tilde{b}_{1,j_1}^r, \dots, \tilde{b}_{N_r,j_{N_r}}^r \right) * \\ * \rho \left(\sigma(a_{1,i_1}^r, \dots, a_{M_r,i_{M_r}}^r), \sigma(b_{1,j_1}^r, \dots, b_{N_r,j_{N_r}}^r) \right) \end{aligned} \quad (5)$$

donde los operadores " σ " y " ρ " expresan, respectivamente, las relaciones de conjunción e implicación escogidas en el proceso de inferencia.

A partir del análisis de la expresión 5 puede observarse que, si " ρ " es un operador conmutativo, la distribución de posibilidad obtenida resulta idéntica a la que se obtendría tras plantear la ejecución de la regla 1 invertida (intercambiando parte antecedente y consecuente) en el esquema *Modus Ponens*. De forma

análoga, si " ρ " verifica la propiedad contrapositiva¹ el planteamiento según el esquema *Modus Tollens* que acabamos de discutir resulta equivalente a desarrollar un esquema *Modus Ponens* sobre la regla 1 modificada negando todas las proposiciones descritas en ella.

No entraremos en la discusión de estos dos casos, puesto que su tratamiento se reduce a realizar la ejecución de una nueva regla siguiendo el esquema *Modus Ponens*, en la cual resultan aplicables mecanismos previos (parametrización de valores lingüísticos, uso de valores lingüísticos de compactación) descritos en la bibliografía [7, 10, 2]. Centraremos nuestro análisis en el tratamiento de operadores de implicación que no verifiquen ninguna de las dos propiedades mencionadas anteriormente. Con objeto de facilitar dicho análisis reescribiremos la expresión 5 de forma que nos permita tratar homogéneamente todas las variables y escoger dentro de un gran abanico los operadores más adecuados. Para ello, basta tener en cuenta que 5 puede reescribirse como

$$\begin{aligned} \bar{a}_{k,i_k} = & \bigvee_{\substack{a_{m_r} \in [0,1] \\ m_r \neq k}} \bigvee_{i_{m_r}/a_{m_r,i_{m_r}}=a_{m_r}} \bigvee_{\substack{b_{n_r} \in [0,1] \\ n_r=1,\dots,N_r}} \\ & \bigvee_{j_{n_r}/b_{n_r,j_{n_r}}=b_{n_r}} \sigma(\bar{b}_{1,j_1}, \dots, \bar{b}_{N_r,j_{N_r}}) * \\ & * \rho(\sigma(a_{1,i_1}, \dots, a_{M_r,i_{M_r}}), \sigma(b_{1,j_1}, \dots, b_{N_r,j_{N_r}})) \end{aligned} \quad (6)$$

Si exigimos que los operadores binarios "*" y " σ " sean monótonos crecientes frente a su primer argumento y la relación de implicación " ρ " sea monótona decreciente frente a su primer argumento², se puede reescribir 6 como

$$\bar{a}_{k,i_k} = \tau_{A^r}(a_{k,i_k}) \quad (7)$$

definiendo la **Función de Razonamiento *Modus Tollens*** τ_{A^r} como

$$\begin{aligned} \tau_{A^r} = & \bigvee_{\substack{b_{n_r} \in [0,1] \\ n_r=1,\dots,N_r}} \sigma(\tau_{B_1^r}(b_1), \dots, \tau_{B_{N_r}^r}(b_{N_r})) * \\ & * \rho(\sigma(0, 0, \dots, a), \sigma(b_1, \dots, b_{N_r})) \end{aligned} \quad (8)$$

y los **Grados de Verificación Lingüísticos**, $\tau_{B_{n_r}^r}(b)$ $n_r=1,\dots,N_r$, como

$$\tau_{B_{n_r}^r}(b) = \bigvee_{j_{n_r}/b_{n_r,j_{n_r}}=b} \bar{b}_{n_r,j_{n_r}}, \quad n_r=1, \dots, N_r, b \in [0,1] \quad (9)$$

La distribución de posibilidad que describe esta última ecuación puede interpretarse como una medida

¹Es decir, se cumple $\rho(a, b) = \rho(\bar{b}, \bar{a}) \forall a, b \in [0, 1]$.

²La primera exigencia se cumple trivialmente en el caso habitual de que "*" y " σ " sean t-normas. La segunda es una exigencia axiomática que verifican las relaciones de implicación procedentes de la lógica clásica y las que reflejan una ordenación parcial de las proposiciones.

de la compatibilidad existente entre las distribuciones de posibilidad asociadas a las variables de la parte consecuente de la regla ($B_{n_r}^r$, $n_r=1,\dots,N_r$) y las correspondientes distribuciones que expresan el conocimiento disponible (por observación o por inferencias realizadas previamente) sobre dichas variables ($\bar{B}_{n_r}^r$, $n_r=1,\dots,N_r$).

La expresión 8 constituye la versión para el esquema *Modus Tollens* de la **Función de Razonamiento Hacia Adelante** [7] del *Modus Ponens*. La aplicación de esta función sobre las distribuciones de posibilidad descritas en la parte antecedente de la regla da lugar al proceso de inferencia, tal y como muestra la expresión 7. El proceso puede describirse, por tanto, en tres fases:

- Obtención de los Grados de Verificación (GDV) lingüísticos correspondientes a las proposiciones de la Parte Consecuente de la regla en la que se realiza el paso de los correspondientes universos de discurso a un universo de discurso común para todas las variables: el intervalo $[0,1]$ o dominio de la verdad.
- El cálculo de la Función de Razonamiento *Modus Tollens* (igualmente en el dominio de la verdad).
- La inferencia de las distribuciones de posibilidad asociadas a las variables de la parte antecedente de la regla, donde se realiza el paso del dominio de la verdad a los universos de discurso correspondientes a cada variable.

La Figura 1 muestra como ejemplo gráfico la realización del proceso de inferencia siguiendo la aplicación directa de la Regla Composicional de Inferencia "Sup*" (Figura 1a) y mediante el proceso en tres fases que acabamos de describir (Figura 1b). En ella puede apreciarse la dependencia que presenta el primer caso frente a los respectivos universos de discurso. En el segundo método, el mecanismo es independiente de este parámetro, lo que permite tratar de forma homogénea todas las variables. Por simplicidad a la hora de expresar el ejemplo se ha realizado la representación tomando como relación de implicación la t-norma "mínimo", aun cuando ésta verifica la propiedad conmutativa y el proceso, en este caso, puede abordarse a partir del planteamiento alternativo discutido en el Apartado 1.

Tras este análisis, disponemos de un mecanismo que permite realizar el proceso de inferencia *Modus Tollens* en la gran mayoría de las situaciones. Además de la ventaja ya comentada del tratamiento homogéneo que se realiza de todas las variables (intervalo $[0,1]$) veremos a continuación la ventaja que también supone a la hora de considerar el encadenamiento entre reglas, al permitir tratar el proceso de implicación como la "propagación" de los GDV lingüísticos entre las reglas que constituyen el encadenamiento. El análisis del proceso de implicación *Modus Tollens* entendido de esta forma conduce a una compactación de la BCB que permite eludir en todo momento la utilización de las distribuciones asociadas a las distintas proposiciones. De esta

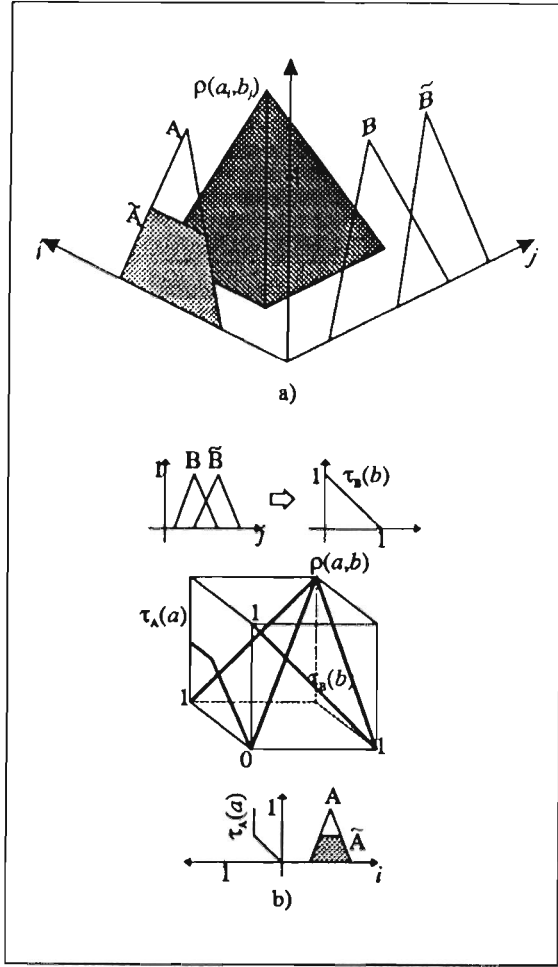


Figura 1: Ejecución Modus Tollens de una regla según la Regla Composicional de Inferencia. a) Ejecución sobre los universos de discurso. b) Ejecución en el dominio de la verdad.

forma se logra una reducción en el número de operaciones a realizar en el proceso global de ejecución de la BCB, tal y como veremos en el siguiente apartado.

2.2 Razonamiento Modus Tollens en BCB con encadenamiento

Vamos a considerar una BCB en la que una misma variable intermedia X figura en la Parte Consecuente de una regla R^S y en la Parte Antecedente de un conjunto de reglas R^t , $t=1, \dots, T$. Utilizando la notación descrita en 1, podemos caracterizar esta situación, sin pérdida de generalidad, como

$$\begin{aligned} &SI \ X_1^S \ ES \ A_1^S \ Y... \ ENTONCES \ X_{M_S+1}^S \ ES \ B_1^S \ Y... \\ &SI \ X_1^1 \ ES \ A_1^1 \ Y... \ ENTONCES \ X_{M_1+1}^1 \ ES \ B_1^1 \ Y... \end{aligned}$$

$$SI \ X_1^T \ ES \ A_1^T \ Y... \ ENTONCES \ X_{M_T+1}^T \ ES \ B_1^T \ Y... \quad (10)$$

con

$$X_{M_S+1}^S = X_1^1 = X_1^T \triangleq X \quad (11)$$

Diremos en esta situación que la regla R^S está encadenada con las reglas R^t , $t=1, \dots, T$ a través de la variable X .

Esta caracterización del encadenamiento entre reglas a nivel de variable [1, 2] resulta más flexible que otras descritas en la bibliografía [10, 9] y es, además, más coherente con la filosofía en que se ha basado tradicionalmente el razonamiento borroso.

El proceso de razonamiento *Modus Tollens* sobre esta BCB tendrá como objetivo la asignación de una distribución de posibilidad a cada una de las variables de la parte antecedente de R^S , siguiendo el mecanismo descrito en el apartado anterior, para lo cual será necesario disponer de los GDV lingüísticos de las proposiciones de la Parte Consecuente de dicha regla. La existencia de encadenamiento entre reglas impone que la información acerca de la variable $X_{M_S+1}^S$ se obtiene a partir de un proceso de inferencia *Modus Tollens* en el conjunto de reglas R^t , $t=1, \dots, T$, y la posterior agregación de las distribuciones inferidas, que conduce al cálculo del GDV lingüístico $\tau_{B_1^S}$ correspondiente a la proposición " $X_{M_S+1}^S \ ES \ B_1^S$ ". Veremos cómo este proceso puede simplificarse notablemente, y cómo la obtención de $\tau_{B_1^S}$ puede realizarse directamente a partir de los GDV lingüísticos de las proposiciones de la parte consecuente de R^t , $t=1, \dots, T$.

Teniendo en cuenta 9,

$$\tau_{B_1^S}(b) = \bigvee_{j_1/b_{1,j_1}^S = b} \tilde{b}_{1,j_1}^S, \quad b \in [0, 1] \quad (12)$$

Pero, dado el encadenamiento existente entre las reglas, la distribución asociada a la variable intermedia provendrá de la agregación de las distribuciones inferidas previamente en las reglas R^t , $t=1, \dots, T$. Es decir,

$$\tilde{b}_{1,i}^1 = \bigvee_{t=1}^T \tilde{a}_{1,i}^t. \quad \text{Por tanto, tendremos, finalmente, que}$$

$$\tau_{B_1^S}(b) = \bigvee_{t=1}^T \tau_{A_1^t}(\mu_{A_1^t, B_1^S}(b)), \quad b \in [0, 1] \quad (13)$$

donde $\tau_{A_1^t}$, $t=1, \dots, T$ son las respectivas Funciones de Razonamiento Hacia Adelante (ver expresión 8) y se define el Valor Lingüístico de Compactación entre las distribuciones A_1^t y B_1^S como

$$\mu_{A_1^t, B_1^S}(b) = \bigwedge_{i/b_{1,i}^S = b} a_{1,i}^t \quad (14)$$

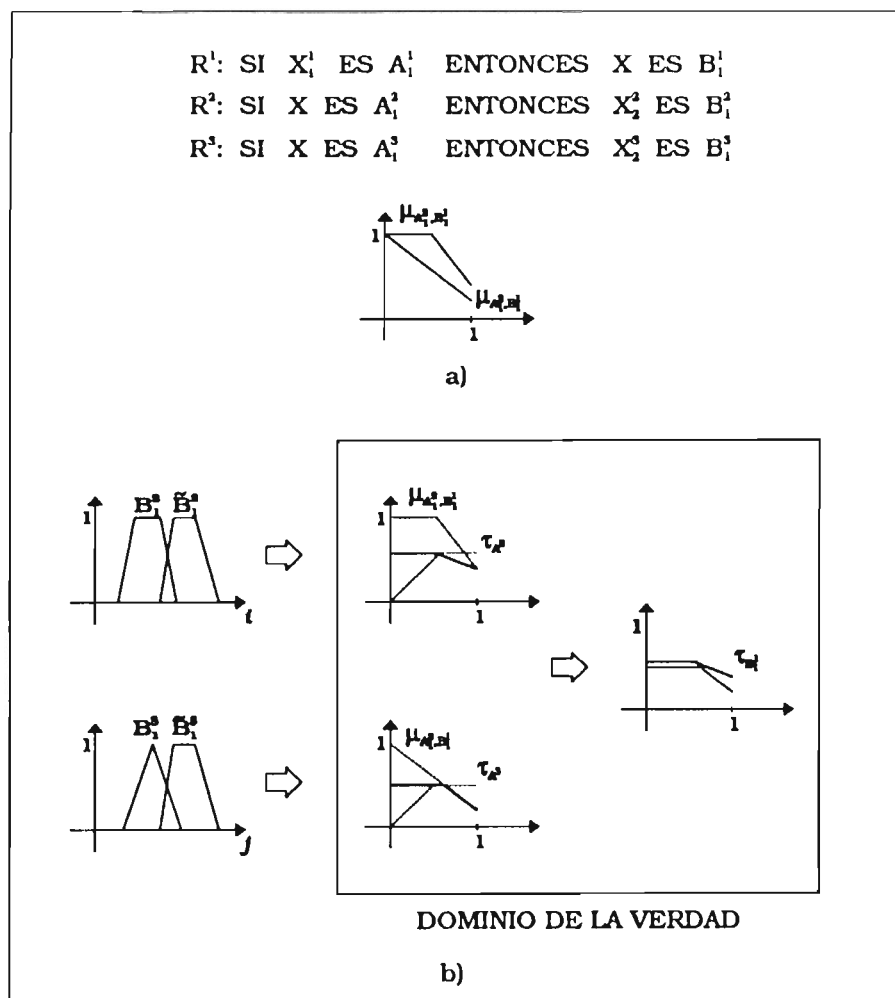


Figura 2: Obtención del GDV lingüístico de una proposición. Operaciones a realizar a priori (a) y durante la ejecución Modus Tollens (b).

que, como su propia definición indica, puede ser calculado a priori. La Figura 2 muestra una sencilla descripción gráfica de los resultados obtenidos. En ella se presenta el proceso de ejecución según el esquema *Modus Tollens* en una BCB formada por 3 reglas.

En el desarrollo aquí resumido únicamente se exige que las funciones $\tau_{A_i^t}$, $t=1, \dots, T$, sean monótonas decrecientes, lo cual se verifica si las relaciones de implicación correspondientes lo son frente a su primer parámetro. Tenemos, por tanto, que únicamente dos valores son relevantes en el cálculo del GDV lingüístico de una proposición asociada a una variable intermedia en la Parte Consecuente de una regla: los GDV lingüísticos de las proposiciones de la Parte Consecuente de las reglas con las que está encadenada y los valores lingüísticos de compactación correspondientes. Puede interpretarse, pues, el proceso de realización de inferencias en BCB con encadenamiento como la simple propagación de los GDV lingüísticos de ciertas proposiciones y, a partir de ellos, y aplicando 13, ir

obteniendo progresivamente los GDV lingüísticos de las restantes proposiciones mediante operaciones de inferencia y encadenamiento realizadas siempre en el intervalo $[0,1]$. Finalmente, y sólo para aquellas variables de interés, se realizará un proceso de conversión de los GDV lingüísticos en distribuciones de posibilidad, a través de la expresión 6; en el caso de que para estas variables las inferencias se hayan realizado a través de varias reglas, se desarrollará un proceso de agregación, ya sobre los universos de discurso donde se definen.

El ahorro computacional obtenido se centra, básicamente, en dos puntos principales: la realización a priori de todas las operaciones de obtención de los distintos valores de compactación y la realización del proceso global de inferencia en BCB como la "propagación" de GDV lingüísticos. En el primer caso, esto es posible porque las distribuciones de posibilidad que intervienen en el cálculo son conocidas con anterioridad al momento de ejecución de la BCB. Las

operaciones relacionadas con el encadenamiento entre reglas quedan recogidas en él, con lo cual se traslada parte de la carga computacional desde la fase de ejecución a la fase de definición de la BCB. La ejecución de ésta como "propagación" de GDV lingüísticos permite eludir, además, el cálculo de las distribuciones de posibilidad inferidas para las variables que no sean de interés, y la utilización de las distribuciones de posibilidad descritas en la BCB para dichas variables. La aplicación de la Regla Composicional de Inferencia en su formulación tradicional exigiría, por contra, la realización del proceso en sucesivas etapas de inferencia y agregación, y la consiguiente toma en consideración y almacenamiento de las distribuciones de posibilidad correspondientes.

3 Discusión

En este trabajo se ha planteado la ejecución del proceso de inferencias en BCB con encadenamiento siguiendo el esquema *Modus Tollens*. El mecanismo de encadenamiento a nivel de variables que hemos planteado permite una mayor flexibilidad en la realización de inferencias. Frente a otras caracterizaciones más estrictas del encadenamiento entre reglas [9, 4, 5] que exigen una total identificación en las proposiciones, nuestra solución adopta, a la hora de plantear encadenamiento entre reglas, la filosofía de no necesidad de una total adecuación entre las distribuciones de posibilidad. Asimismo, los operadores asumidos permiten, con alguna ligera restricción, un elevado grado de flexibilidad, a través de múltiples combinaciones cuya adecuación a cada problema concreto deberá ser objeto de un estudio previo. Según las propiedades lógicas que se deseen en cada caso, se deberán escoger unos operadores u otros, y un mecanismo de realización de inferencias concreto. Así, para relaciones de implicación que verifiquen las propiedades conmutativa o contrapositiva, puede resultar más conveniente reformular la BCB y adoptar una estrategia *Modus Ponens*. Para otros casos se describe un planteamiento general del problema, que aborda el proceso de ejecución mediante una estrategia de compactación. Esta estrategia permite un tratamiento homogéneo de todas las variables y traslada parte de la carga computacional al momento de definición de la BCB. La cuantificación de este ahorro computacional, junto con un análisis más detallado de la adecuación de cada estrategia para cada caso concreto, serán objeto de un estudio posterior.

Referencias

- [1] BUGARÍN, A. BARRO, S. Y REGUEIRO, C.V. "Saving Computation Time through Compaction of Chained Fuzzy Rules". *Proc. 1st European Conf. on Fuzzy and Intelligent Technologies (EUFIT '93)*, 3, (1993), 1543-1549.
- [2] BUGARÍN, A. Y BARRO, S. "Fuzzy Reasoning supported by Petri Nets". *IEEE Transactions on Fuzzy Systems*, 2, 2, (1994), 135-150.
- [3] BUGARÍN, A. *Soluciones para la computación eficiente en sistemas basados en lógica borrosa*. Tesis Doctoral, Universidad de Santiago de Compostela 1994.
- [4] CHEN, S.M., KE, J.S., Y CHANG, J.F. "Knowledge Representation Using Fuzzy Petri Nets". *IEEE Transactions on Systems, Man and Cybernetics*, 2, 3, (1990), 311-319.
- [5] CHUN, M.G., Y BIEN, Z. "Fuzzy Petri-Net Representation and Reasoning Methods for Rule-Based Decision Making Systems". *IEICE Transactions on Fundamentals*, E76-A, 6, (1993), 974-983.
- [6] EZAWA, Y. Y KANDEL, A. "Robust Fuzzy Inference". *International Journal of Intelligent Systems*, 6,
- [7] GODÓ, L., JACAS, J. Y VALVERDE, L. "Fuzzy Values in Fuzzy Logic". *International Journal of Intelligent Systems*, 6, (1991), 199-212.
- [8] GUPTA, M.M. Y QI, J. "Theory of T-norms and Fuzzy Inference Methods". *Fuzzy Sets and Systems*, 40, (1991), 431-450.
- [9] LOONEY, C.G. "Fuzzy Petri Nets for Rule-Based Decision-making". *IEEE Transactions on Systems, Man and Cybernetics*, 18, 1, (1988), 178-183.
- [10] UEHARA, K. Y FUJISE, M. "Multi-Stage Fuzzy Inference by Propagation of Linguistic Truth Values". *Proc. 4th IFSA World Congress*, (1991), 196-199.
- [11] ZADEH, L.A. "Outline of a New Approach to the Analysis of Complex Systems and Decision Processes". *IEEE Transactions on System, Man and Cybernetics*, 3, 1, (1973), 28-44.

El Uso de Conectivos Difusos en el Diseño de Algoritmos Genéticos*

F. Herrera, M. Lozano, J.L. Verdegay

Departamento de Ciencias de la Computacion e I.A

Universidad de Granada,

18071 - Granada, España

e-mail: herrera,lozano,verdegay@robinson.ugr.es

Resumen

Los Algoritmos Genéticos son procedimientos adaptativos para la búsqueda de soluciones en espacios complejos, inspirados en la evolución biológica y con operaciones basadas en la Genética.

Uno de los principales problemas que aparece en la aplicación de los Algoritmos Genéticos es la inseguridad de convergencia hacia el óptimo global. Distintas técnicas basadas en la lógica difusa se pueden utilizar para resolver este problema. Este trabajo presenta una de ellas, consistente en el diseño de operadores de cruce usando conectivos difusos para Algoritmos Genéticos con codificación real, y su extensión, basados en el uso de conectivos difusos parametrizados.

Palabras clave: Algoritmos Genéticos, Codificación Real, Conectivos Difusos.

1 Introducción. Algoritmos Genéticos. Codificación Real

Los Algoritmos Genéticos (AG) son procedimientos adaptativos para la búsqueda de soluciones en espacios complejos inspirados en la evolución biológica, con patrones de operaciones basados en el principio darwiniano de reproducción y supervivencia de los individuos que mejor se adaptan al entorno en el que viven ([5,9]).

Los mecanismos fundamentales bajo los cuales funciona un AG son: operar sobre una población de individuos (soluciones factibles), que inicialmente se genera de forma aleatoria, y cambiar en cada iteración t la población atendiendo a los tres siguientes pasos: (1) evaluación de los individuos de la población, (2) selección de un nuevo conjunto de individuos y reproducción sobre la base de su relativa conveniencia o adaptación, y (3) recombinación para formar una nueva población a partir de los operadores de cruce y

mutación. Los individuos resultantes de estas tres operaciones forman la siguiente población, iterando este proceso hasta que el sistema cese de mejorar, o se llegue a una iteración fijada T .

Generalmente, cada individuo de la población se representa mediante una cadena binaria de longitud fija llamada cromosoma que codifica los valores de las variables que intervienen en el problema. A cada una de las posiciones del cromosoma se le denomina gen.

Existen razones teóricas ([6]) que presentan al alfabeto binario como muy indicado para codificar las soluciones. Sin embargo, el buen comportamiento demostrado por los AG no radica en el uso de cadenas de ceros y unos ([1, 6]), por ello se abre un camino hacia el uso de alfabetos con un cardinal más alto, seguido del desarrollo de nuevos operadores genéticos sobre estos alfabetos. En particular, para problemas de optimización de parámetros con variables sobre dominios continuos, parece natural representar los genes como números reales directamente. De esta forma, un cromosoma será un vector de números reales con el mismo tamaño que el vector solución del problema. Cada gen va a representar una variable del problema. A estos algoritmos los llamaremos Algoritmos Genéticos con Codificación Real (AGCR).

El uso de parámetros reales nos posibilita usar grandes dominios (incluso dominios desconocidos) para las variables, tarea difícil en las implementaciones binarias donde incrementar el dominio significaría sacrificar precisión, supuesta una longitud fija para los cromosomas. Otra ventaja al utilizar parámetros reales es su capacidad para explotar la gradualidad de las funciones con variables continuas (donde el concepto gradualidad se refiere a que pequeños cambios en las variables se corresponden con pequeños cambios en la función). También desaparece la necesidad de que los rangos para las variables sean potencia de dos. Por último se permite un diseño más fácil de las herramientas que manejan restricciones no triviales, y la representación de las soluciones está muy cercana a la formulación natural de muchos problemas ([12]).

El operador de mutación permite explorar el espa-

* Trabajo financiado por el proyecto PB92-0933

cio de búsqueda: dado un cromosoma se escoge uno de sus genes aleatoriamente y se cambia de valor. En el caso de usar codificación binaria el cambio es obvio ("1" por "0" y viceversa). Sin embargo, para el caso de codificación real han aparecido diversas versiones ([7, 9]). Este operador aporta nuevo material genético a la población permitiendo visitar nuevas zonas del espacio.

El operador de cruce permite explotar la información actual disponible en la población sobre el espacio de búsqueda: a partir de dos cromosomas se generan dos nuevos a través de un intercambio de genes. La idea subyacente radica en la propia filosofía del proceso de búsqueda de un AG, el cual persigue obtener cada vez mejores soluciones. Una forma de hacerlo es combinando buenas soluciones ya conocidas. En [9] se presentan diversos modelos aparecidos en la literatura para el operador de cruce utilizando codificación binaria y real.

El operador de cruce juega un papel fundamental en los AG, de hecho puede considerarse como una de las características que definen y diferencian a los AG de otros algoritmos de evolución, siendo uno de los componentes que más se ha tenido en consideración a la hora de mejorar el comportamiento de los AG ([11]). El operador de mutación representa un complemento al operador de cruce, necesario para evitar que, debido a este último, se pierda información útil para la búsqueda ([5]).

Asociado a los operadores genéticos de cruce y mutación existe la probabilidad de cruce, P_c , y la probabilidad de mutación, P_m . La primera es la probabilidad de aplicar el operador de cruce sobre parejas de cadenas obtenidas en la reproducción, la segunda es la probabilidad de aplicar el operador de mutación sobre los caracteres de cada cadena.

El problema principal que aparece en la aplicación de los AG, es la inseguridad de convergencia hacia el óptimo global. Este problema es consecuencia de la posibilidad de una convergencia prematura hacia un óptimo local ocasionada por una falta de diversidad en la población, junto con una desproporcionada relación entre explotación y exploración mantenida por el AG. Bajo estas circunstancias el operador de cruce se convierte en poco eficaz, ya que los alelos de algunas posiciones de todos los cromosomas de la población son iguales. Una probabilidad de mutación baja no ayudaría a solucionar el problema, mientras que una probabilidad alta causa una pérdida de zonas del espacio prometedoras que puedan estar siendo tratadas.

Para resolver estos problemas se pueden utilizar distintas técnicas basadas en las lógicas difusas. Una de ellas consiste en integrar en el AG un sistema de control difuso que controle dinámicamente los parámetros del AG (P_m , P_c , etc), de forma que se presente en cada momento una correcta relación explotación/exploración, junto con un nivel de diversidad adecuado ([10, 8]). Otra va encaminada al diseño de operadores de cruce usando conectivos difusos ([7, 8]).

En este trabajo presentamos operadores de cruce para AGCR diseñados utilizando conectivos difusos ([7, 8]), y su extensión, basada en el uso de conec-

tivos difusos parametrizados para diseñar operadores de cruce dinámicos.

2 Diseño de Operadores de Cruce para AGCR utilizando Conectivos Difusos

Como ya se ha apuntado los cromosomas manipulados por un AGCR son vectores de números reales cuya precisión viene marcada por la del ordenador bajo el que se ejecuta el algoritmo. Su tamaño es idéntico al del vector solución del problema, de modo que, cada gen representa una variable del problema. Los valores de los genes de un cromosoma están forzados a pertenecer al intervalo establecido para la variable que representan, de modo que los operadores genéticos tendrán que preservar este requerimiento.

Según Booker ([3]) el operador de cruce es uno de los puntos claves para hacer frente a la convergencia prematura, de aquí que la solución a este problema pase por el diseño de nuevas alternativas de este operador. En esta línea, en [7] se presentaron nuevos operadores de cruce para AGCR basados en el uso de los conectivos difusos ([13, 14]) t-normas, t-conormas y funciones promedio, que proporcionan distintos niveles de diversidad en la población, extendiendo el estudio a funciones de compensación en [8].

El operador de cruce consiste en combinar los genes de la siguiente forma. Supuestos dos genes c_i^1 y c_i^2 que se van a cruzar, $c_i^1, c_i^2 \in [m_i, M_i]$, y siendo $x_i = \min(c_i^1, c_i^2)$ e $y_i = \max(c_i^1, c_i^2)$, el intervalo de actuación del gen i $[m_i, M_i]$ se puede dividir en tres regiones $[m_i, x_i]$, $[x_i, y_i]$, e $[y_i, M_i]$ donde se pueden obtener descendientes, incluso el considerar una región $[x'_i, y'_i]$ con $x'_i \leq x_i$ e $y'_i \geq y_i$ parece razonable, lo cual produciría diversidad en la búsqueda realizada. Gráficamente, los intervalos son los siguientes:



Así, proponemos un conjunto de operadores de cruce que permitan obtener descendientes en los anteriores intervalos. Para ello se utilizan cuatro funciones F , S , M y L definidas de $[a, b] \times [a, b]$ en $[a, b]$, $a, b \in \mathbb{R}$, y que cumplen:

$$(P1) \quad \forall x, y \in [a, b] \quad F(x, y) \leq \min(x, y)$$

$$(P2) \quad \forall x, y \in [a, b] \quad S(x, y) \geq \max(x, y)$$

$$(P3) \quad \forall x, y \in [a, b] \quad \min(x, y) \leq M(x, y) \leq \max(x, y)$$

$$(P4) \quad \forall x, y \in [a, b] \quad F(x, y) \leq L(x, y) \leq S(x, y)$$

$$(P5) \quad F, S, M, \text{ y } L \text{ son monótonas, no decrecientes.}$$

Supuesto que $Q \in \{F, S, M, L\}$ y que $C_1 = (c_1^1 \dots c_n^1)$ y $C_2 = (c_1^2 \dots c_n^2)$ son dos cromosomas que han sido seleccionados para aplicarles el operador de cruce. Se puede generar el cromosoma $H = (h_1 \dots h_n)$ como

$$H = Q(C_1, C_2), \quad h_i = Q(c_i^1, c_i^2), \quad i = 1, \dots, n$$

Utilizando los operadores t-normas, t-conormas, funciones promedio, y operadores de compensación generalizados que se utilizan como conectivos difusos ([13, 14]), vamos a asociar F a una t-norma, S a una t-conorma, M con un operador promedio, y L con un operador de compensación generalizado. Para ello, es necesario un conjunto de transformaciones lineales para poder aplicar estos operadores bajo los intervalos de definición de los genes.

Sea O un operador perteneciente al conjunto formado por las t-normas, t-conormas, funciones promedio y operadores de compensación generalizados, para cada posición $i \in \{1, \dots, n\}$ se realizarán las siguientes operaciones:

1. Transformar c_i^1 y c_i^2 en los valores $s_i^1, s_i^2 \in [0, 1]$ de forma que :

$$s_i^k = \frac{c_i^k - m_i}{M_i - m_i}, \quad k = 1, 2$$

Con este paso transformamos los valores de los genes para que se les pueda aplicar el operador.

2. Aplicar el operador $O(s_i^1, s_i^2)$ y calcular el valor h_i :

$$h_i = m_i + (M_i - m_i) \cdot O(s_i^1, s_i^2)$$

de modo que el gen h_i es el valor para el gen de la posición i del cromosoma resultante del cruce de C_1 y C_2 , y está en relación a sus límites originales, $h_i \in [m_i, M_i]$.

Partiendo de un conjunto de conectivos $(T_j, G_j, P_j, \hat{C}_j)$, $j = 1, \dots, k$, se pueden construir un conjunto de funciones F_j, S_j, M_j y L_j asociadas de la siguiente forma:

$$\begin{aligned} F_j(c_i^1, c_i^2) &= m_i + (M_i - m_i) \cdot T_j(s_i^1, s_i^2) \\ S_j(c_i^1, c_i^2) &= m_i + (M_i - m_i) \cdot G_j(s_i^1, s_i^2) \\ M_j(c_i^1, c_i^2) &= m_i + (M_i - m_i) \cdot P_j(s_i^1, s_i^2) \\ L_j(c_i^1, c_i^2) &= m_i + (M_i - m_i) \cdot \hat{C}_j(s_i^1, s_i^2) \end{aligned}$$

Estos operadores tienen la propiedad de ser continuos y no decrecientes.

3 Ejemplo

En [7, 8] se presentan distintas pruebas realizadas para comparar el comportamiento de AG clásicos con los AGCR utilizando operadores propuestos en la literatura y por último con un conjunto de familias de AGCR con operadores de cruce diseñados usando el conjunto de conectivos $(T_j, G_j, P_j, \hat{C}_j)$, $j = 1, \dots, 5$ que se muestran en la siguiente tabla:

t-norma	t-conorma
<i>Producto Lógico</i> $T_1(x, y) = \min(x, y)$ $P_1(x, y) = \lambda x + (1 - \lambda)y$	<i>Suma Lógica</i> $G_1(x, y) = \max(x, y)$
<i>Producto Hamacher</i> $T_2(x, y) = \frac{xy}{x+y-xy}$ $P_2(x, y) = \frac{1}{\frac{1}{\lambda x + (1-\lambda)y} + 1}$	<i>Suma Hamacher</i> $G_2(x, y) = \frac{x+y-2xy}{1-xy}$
<i>Producto Algebraico</i> $T_3(x, y) = xy$ $P_3(x, y) = x^\lambda y^{1-\lambda}$	<i>Suma Algebraica</i> $G_3(x, y) = x + y - xy$
<i>Producto Einstein</i> $T_4(x, y) = \frac{xy}{1+(1-x)(1-y)}$ $P_4(x, y) = \frac{2}{1+(\frac{2-x}{x})^\lambda (\frac{2-y}{y})^{1-\lambda}}$	<i>Suma Einstein</i> $G_4(x, y) = \frac{x+y}{1+xy}$
<i>Producto Acotado</i> $T_5(x, y) = 0 \vee (x + y - 1)$ $P_5(x, y) = \lambda x + (1 - \lambda)y$	<i>Suma Acotada</i> $G_5(x, y) = 1 \wedge (x + y)$

Tabla 1: Conjunto de operadores

Cada t-conorma es dual a la t-norma mostrada a su izquierda. La función promedio está calculada a partir de $P(x, y) = f^{-1}((1 - \lambda)f(x) + \lambda f(y))$, ($0 \leq \lambda \leq 1$) donde f es la función generadora aditiva de la t-norma colocada en el mismo nivel ([13, 14]), excepto para la primera, que al no ser Arquimediana no tiene función generadora, y utilizaremos $f(x) = x$. Para cada familia j ($j = 2, \dots, 5$) de operadores de la tabla 1 se considera un operador de compensación $\hat{C}_j$ definido de la siguiente forma:

$$\hat{C}_j = P_j(T_j, G_j)$$

Para la familia de operadores lógicos se considera $\hat{C}_1 = T_1^{1-\lambda} G_1^\lambda$.

Las familias de algoritmos se distinguen en como se realizan los dos pasos siguientes:

1. Generación de hijos utilizando los distintos operadores definidos.
2. Selección de los hijos fruto del cruce que formarán parte de la población.

Una propuesta es la siguiente: Para cada par de cromosomas de un total de $\frac{1}{2} * P_c * N$, siendo P_c la probabilidad de cruce y N el tamaño de la población, se generan cuatro descendientes, resultado de aplicarles las funciones F, S, M y L .

ciones las t-normas se acercarán al mínimo y las t-conormas al máximo, proporcionando convergencia.

Proponemos un conjunto de operadores de cruce basados en las familias de funciones: $\{F^p\}_{p=1,\dots,G}$, $\{S^p\}_{p=1,\dots,G}$, y $\{M^p\}_{p=1,\dots,G}$, $G \in \mathbb{N}$ definidas de $[a, b] \times [a, b]$ en $[a, b]$, $a, b \in \mathbb{R}$, cumpliendo las correspondientes propiedades P1-P5 junto con las siguientes:

$$(P6) \quad \forall x, y \in [a, b], \lim_{p \rightarrow G} F^p(x, y) \cong \min(x, y)$$

$$(P7) \quad \forall x, y \in [a, b], \lim_{p \rightarrow G} S^p(x, y) \cong \max(x, y)$$

$$(P8) \quad \forall x, y \in [a, b], \min(x, y) \leq M^p(x, y) \leq \frac{x+y}{2} \\ \text{ó } \frac{x+y}{2} \leq M^p(x, y) \leq \max(x, y), \quad \forall x, y \in [a, b], \lim_{p \rightarrow G} M^p(x, y) \cong \frac{x+y}{2}$$

Supongamos que en un AG con un máximo de generaciones τ los cromosomas $C_1^t = (c_{11}^t, \dots, c_{1n}^t)$ y $C_2^t = (c_{21}^t, \dots, c_{2n}^t)$ han sido seleccionados para aplicarles el operador de cruce en la generación t . Si $Q_p \in \{F^p, S^p, M^p\}$, $p = 1, \dots, \tau$ podemos generar el cromosoma $H^t = (h_1^t, \dots, h_n^t)$ como

$$H^t = Q_i(C_1^t, C_2^t), \quad h_i^t = Q_i(c_{1i}^t, c_{2i}^t), \quad i = 1, \dots, n$$

A continuación vamos a construir familias de funciones con las propiedades (P6) y (P7) usando las t-normas y t-conormas parametrizadas mostradas en la tabla 4. En la tabla 5 se muestran las propiedades que cumplen por las t-normas parametrizadas de esta tabla, siendo éstas, análogas para el caso de las t-conormas. Notar que T_6 es la t-norma drástica (la menor de las t-normas).

Tipo	t-norma y t-conorma
Yager ($q > 0$)	$T_1^q(x, y) = 1 - (1 \wedge \sqrt[q]{(1-x)^q + (1-y)^q})$ $G_1^q(x, y) = 1 \wedge \sqrt[q]{x^q + y^q}$
Frank ($q > 0$) ($q \neq 1$)	$T_2^q(x, y) = \log_q[1 + \frac{(q^x-1)(q^y-1)}{q-1}]$ $G_2^q(x, y) = 1 - \log_q[1 + \frac{(q^{1-x}-1)(q^{1-y}-1)}{q-1}]$
Dombi ($q > 0$)	$T_3^q(x, y) = \frac{1}{1 + \sqrt[q]{(\frac{1-x}{x})^q + (\frac{1-y}{y})^q}}$ $G_3^q(x, y) = 1 - \frac{1}{1 + \sqrt[q]{(\frac{x}{1-x})^q + (\frac{y}{1-y})^q}}$
Dubois ($0 \leq q \leq 1$)	$T_4^q(x, y) = \frac{xy}{x \vee y \vee q}$ $G_4^q(x, y) = 1 - \frac{(1-x)(1-y)}{(1-x) \vee (1-y) \vee q}$

Tabla 4: Operadores dinámicos

Tipo	T_6	T_5	T_4	T_3	T_2	T_1
Yager	$0 \leftarrow$	1				∞
Frank		∞		$\rightarrow 1$		$\rightarrow 0$
Dombi	$0 \leftarrow$				1	∞
Dubois				1		0

Tabla 5: Propiedades de los operadores dinámicos

Se puede comprobar que la relación de orden (de menor a mayor) en la velocidad de convergencia hacia

el mínimo de estas t-normas es: Yager, Frank, Dombi, y Dubois.

Podemos asociar a la familia de funciones $\{F^p\}_{p=1,\dots,\tau}$ una t-norma parametrizada T^q , $1 \leq q \leq \infty$, y a $\{S^p\}_{p=1,\dots,\tau}$ una t-conorma parametrizada G^q , $1 \leq q \leq \infty$, y de forma que $q = \delta(p)$ con $\lim_{p \rightarrow \tau} \delta(p) \cong \infty$.

A la familia de funciones $\{M^p\}_{p=1,\dots,\tau}$ le asociamos la función promedio parametrizada definida como

$$\forall x, y \in [0, 1], \quad P^q(x, y) = \sqrt[q]{\frac{x^q + y^q}{2}}, \quad -\infty < q < \infty,$$

obtenida de la expresión de la media generalizada $Mg(x, y) = f^{-1}((1-\lambda)f(x) + \lambda f(y))$, ($0 \leq \lambda \leq 1$) con $f(x) = x^q$, y $\lambda = \frac{1}{2}$, y cumpliendo que $\forall x, y \in [0, 1]$, $\lim_{q \rightarrow -\infty} P^q(x, y) = \min(x, y)$, $\lim_{q \rightarrow +\infty} P^q(x, y) = \max(x, y)$ y $\lim_{q \rightarrow 1} P^q(x, y) = \frac{x+y}{2}$. Asociando a $\{M^p\}$ consideramos $\{P^q\}$ con $q = \psi(p)$ donde

$\psi : [1, \tau] \rightarrow [-\infty, 1]$ es una función

monótona creciente ó

$\psi : [1, \tau] \rightarrow [1, +\infty]$ es una función

monótona decreciente.

Al igual que en la sección anterior hace falta un conjunto de transformaciones sobre los genes para poder aplicar estos operadores.

Se han implementado cuatro algoritmos AGD1-AGD4 basados en el uso de los operadores dinámicos presentados. Para cada pareja de cromosomas a cruzar se generan cuatro descendientes aplicando funciones F^p , S^p , y dos M^p cumpliendo cada una de ellas una parte distinta de la propiedad (P8), los dos más prometedores forman parte de la nueva población.

A continuación se muestran las funciones $\delta(\cdot)$ asociadas con cada uno de los operadores dinámicos utilizados.

Operadores	$\delta(p)$
Yager	$\frac{1}{\ln(\frac{p+1}{p})}$
Frank	$\ln(\frac{p+1}{p})$
Dombi	$\frac{1}{\ln(\frac{p+1}{p})}$
Dubois	$\frac{1}{\sqrt[p]{p}}$

Tabla 6: Funciones $\delta(\cdot)$ utilizadas

La función $\psi(p)$ utilizada es $1 + \ln(\frac{p}{e})$ como función monótona creciente y $1 + \ln(\frac{e}{p})$ como función monótona decreciente.

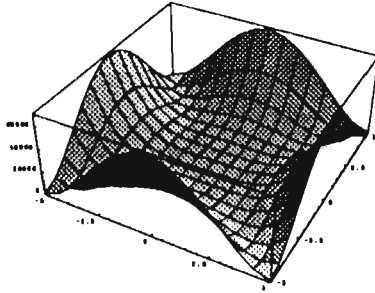
Los resultados medios obtenidos bajo las mismas condiciones que los anteriores experimentos son:

Algoritmos	100	500	5000
AGD1	5.149822e+00	3.004678e+00	8.060713e-01
AGD2	8.111676e+00	2.855141e+00	5.509746e-01
AGD3	8.175044e+00	2.343028e+00	8.944534e-01
AGD4	2.073980e+00	8.210735e-01	1.406304e-03

Tabla 6: Resultados para los AGD

Los resultados muestran que los algoritmos AGD1-AGD3 no mejoran los resultados de la tabla 3. Esto

A continuación se muestran los resultados obtenidos sobre la función de generalizada de Rosenbrock, con la formulación analítica y gráfica siguiente:



$$f(\vec{x}) = \sum_{i=1}^{n-1} (100 \cdot (x_{i+1} - x_i^2)^2 + (x_i - 1)^2)$$

$$-5.12 \leq x_i \leq 5.12$$

$$\min(f) = f(1, \dots, 1) = 0$$

Para realizar los experimentos se han usado un conjunto de AGCR (AGR1-AGR9) basados en distintos operadores de cruce presentados en la literatura. Las características de cada uno de estos algoritmos se muestran en la siguiente tabla:

Algoritmos	Mutación	Cruce
AGR1	Aleatoria	Simple [16,12]
AGR2	No Unif.	Simple [16,12]
AGR3	Aleatoria	Arit. Uniforme [12]
AGR4	No Unif.	Arit. Uniforme [12]
AGR5- α	No Unif.	BLX- α [4]
		($\alpha = 0, .15, .3, .5$)
AGR6	No Unif.	Discreto [15]
AGR7	No Unif.	Linear [16]
AGR8	No Unif.	Inter. Extendido [15]
AGR9	No Unif.	Lineal Extendido [15]

Tabla 2: AG con codificación reales (AGR)

Se ha incluido un AG con codificación binaria (AGB) que utiliza 30 genes por variable, cruce múltiple con dos punto y una probabilidad de selección proporcional.

Por medio de NAGR1, ..., NAGR5 notaremos la familia de AG utilizando los operadores de cruce y selección de descendientes propuestos. Todos los algoritmos utilizan el procedimiento de selección de J. Baker ([2]) junto al modelo elitista. Cada algoritmo se ha ejecutado tres veces distintas para 100, 500, y 5000 generaciones. Los resultados medios obtenidos se muestran en la tabla 3.

Algoritmos	100	500	5000
AGB	1.1262e+01	3.6069e+00	1.9045e+00
AGR1	3.0446e+01	2.0070e+01	6.0669e+00
AGR2	6.9230e+00	2.1448e+00	4.7343e-01
AGR3	4.1941e+00	8.6031e+00	6.3745e+00
AGR4	3.7915e+00	2.9624e+00	8.9244e-01
AGR5-0.0	4.5504e+00	2.3454e+00	9.1602e-01
AGR5-0.15	3.2157e+00	2.9556e+00	7.0929e-01
AGR5-0.3	3.7477e+00	1.5844e+00	4.8854e-01
AGR5-0.5	8.2653e+00	2.1099e+00	1.7329e+00
AGR6	5.5393e+00	2.8379e+00	3.5106e-01
AGR7	4.7115e+00	1.8487e+00	5.1499e-01
AGR8	4.6861e+00	3.5337e+00	5.3325e-01
AGR9	4.2196e+00	2.9374e+00	3.8014e-02
NAGR1	4.1431e+00	1.9687e+00	4.9364e-03
NAGR2	1.2378e+01	7.3868e-01	4.0848e-02
NAGR3	1.0031e+01	3.5037e+00	1.0099e+00
NAGR4	1.2425e+01	4.3985e+00	1.8347e+00
NAGR5	6.0121e+00	2.6930e+00	1.2150e+00

Tabla 3: Resultados Comparativos

4 Diseño de Operadores de Cruce Dinámicos usando Conectivos Difusos Parametrizados

Una medida inicial para solucionar el problema de la convergencia prematura consiste en intensificar la exploración en los comienzos del AG y la explotación al final. Con esto conseguimos una gran diversidad al principio, con lo que aumentamos las probabilidades de visitar las zonas del espacio conteniendo soluciones óptimas. Finalmente, suponiendo que la población contiene información sobre estas zonas, se propicia la convergencia (a través de la explotación) para conseguir aproximarse lo máximo posible a las soluciones óptimas.

La mutación no uniforme ([12]) está diseñada persiguiendo esta idea, de forma que, la medida en la que se muta un gen decrece conforme avanza la ejecución del AG. Dando diversidad al principio, y forzando la convergencia al final.

Los resultados presentados muestran a los conectivos lógicos como los más adecuados a la hora de construir operadores de cruce basados en el uso de conectivos difusos. Así, parece interesante la construcción de operadores de cruce que permitan dar diversidad al comienzo del AG, y al final reflejen un comportamiento semejante al de los operadores de cruce basados en el uso de los conectivos lógicos. De esta forma, en la finalización del AG mantendríamos un buen comportamiento, y la convergencia.

Por ello, nos planteamos generalizar el operador de cruce presentado anteriormente. El nuevo operador se define usando t-normas, t-conormas y funciones promedio parametrizadas, cuyo parámetro dependerá de la generación t en la que se aplica el operador de cruce. Al principio se utilizan t-normas y t-conormas alejadas del mínimo y máximo respectivamente, proporcionando diversidad, conforme avanzan las genera-

es debido a que los operadores dinámicos de Yager, Frank y Dombi convergen de forma lenta hacia los operadores lógicos que son los que mejor comportamiento han mostrado, por ello, el nivel de diversidad durante la ejecución del algoritmo es grande haciendo que la convergencia final sea pobre en afinidad. Por el contrario, el algoritmo AGD4 muestra los mejores resultados para 100 y 5000 generaciones y un segundo puesto en 500 generaciones, esto se explica por la rápida convergencia hacia los operadores lógicos que muestra el operador dinámico de Dubois, ya que incluso, de partida comienza con los operadores algebraicos (Tabla 5), de esta forma, al principio este algoritmo introduce diversidad y rápidamente muestra un comportamiento parecido al algoritmo NAGR1, el cual usa los operadores lógicos para definir su operador de cruce. El algoritmo AGD4 ha inducido una buena relación explotación y exploración propiciando un nivel de diversidad conveniente, y, hace pensar en la necesidad de un estudio de las distintas posibles funciones $\delta(\cdot)$, para descubrir los niveles de diversidad y las etapas del AG donde ha de propiciarse ésta, para que la relación entre explotación y exploración sea óptima y apropiada para mejorar la eficiencia de los AG.

5 Conclusiones

En este trabajo se ha estudiado el uso de conectivos difusos para la definición de nuevos operadores de cruce para AGCR. Los resultados obtenidos muestran que al usar este tipo de operadores se mejora la actuación de AGCR. Los operadores presentados muestran una actuación destinada a ofrecer distintos niveles de diversidad en la población, factor importante frente a la convergencia prematura y una correcta relación entre explotación y exploración. Con el uso de los operadores dinámicos se consigue que el nivel de diversidad cambie a lo largo de la ejecución del algoritmo proporcionando gran diversidad al comienzo y poca al finalizar.

Referencias

1. ANTONISSE J. A new interpretation of schema notation that overturns the binary encoding constraint. En: *Proceeding of the Third International Conference on Genetic Algorithms*, J. David Schaffer (Ed.), (Morgan Kaufmann Publishers, San Mateo, 1989) 86-91.
2. BAKER. J. E. Reducing bias and inefficiency in the selection algorithm. *Proc. 2nd Int. Conf. on genetic algorithms*, LEA Cambridge, MA, 1987, 14-21
3. BOOKER, L. Improving Search in Genetic Algorithms. En: *Genetic Algorithms and Simulated Annealing*, L. Davis (Ed.), (Morgan Kaufmann Publishers, Los Altos, 1987) 61-73.
4. ESHELMAN L.J., SCHAFFER J.D. Real-Coded Genetic Algorithms and Interval-Schemata. En *Foundations of Genetic Algorithms-2*, L.Darrell Whitley (Ed.), (Morgan Kaufmann Publishers, San Mateo, 1993) 187-202.
5. GOLDBERG, D.E. *Genetic Algorithms in Search, Optimization, and Machine Learning*. Addison-Wesley, New York, 1989.
6. GOLDBERG, D.E. Real-Coded Genetic Algorithms, Virtual Alphabets, and Blocking. *Complex Systems* 5 (1991), 139-167.
7. HERRERA, F., LOZANO, M., VERDEGAY, J.L. Algoritmos Genéticos con Parámetros Reales. V *Conferencia de la Asociación Española para la Inteligencia Artificial* (1993), 41-50.
8. HERRERA, F., LOZANO, M., VERDEGAY, J.L. Fuzzy Tools to Improve Genetic Algorithms. Aceptado en *EUFIT'94*, 1994.
9. HERRERA, F., LOZANO, M., VERDEGAY, J.L. Algoritmos Genéticos: Fundamentos, Extensiones y Aplicaciones. *Technical Report #DECSAI-94105*, Dpto. de Ciencias de la Computación e Inteligencia Artificial, Universidad de Granada (1994) Aparecerá en la revista ARBOR.
10. LEE, M.A., TAKAGI, H. Dynamic Control of Genetic Algorithms using Fuzzy Logic Techniques. *Proc. of Fifth Int. Conf. on Genetic Algorithms*, Urbana-Champaign (1993), 76-83.
11. LIEPINS, G.E., VOSE, M.D. Characterizing Crossover in Genetic Algorithms. *Annals of Mathematics and Artificial Intelligence* 5, 1 (1992) 27-34.
12. MICHALEWICZ, Z. *Genetic Algorithms + Data Structures = Evolution Programs*. Springer-Verlag, New York, 1992.
13. MIZUMOTO M. Pictorial Representations of fuzzy connectives, part I: Cases of t-norms, t-conorms and averaging operators. *Fuzzy Sets and Systems* 31 (1989), 217-242.
14. MIZUMOTO M. Pictorial Representations of fuzzy connectives, part II: Cases of Compensatory Operators and Seft-dual Operators. *Fuzzy Sets and Systems* 32 (1989), 45-79.
15. MÜHLENBEIN H., SCHLIERKAMP-VOOSEN D. Predictive Models for the Breeder Genetic Algorithm I. Continuous Parameter Optimization. *Evolutionary Computation* 1 1 (1993), 25-49.
16. WRIGHT A. Genetic Algorithms for Real Parameter Optimization. En: *Foundations of Genetic Algorithms-1*, G.J.E Rawlin(Ed.), (Morgan Kaufmann, San Mateo, 1991) 205-218.

A Temporal Reasoning System based on Fuzzy Temporal Constraints*

Lluís Godo and Lluís Vila

IIIA – CSIC

Camí de Sta. Bàrbara, s/n, 17300 Blanes, Catalonia, Spain

{godo,vila}@ceab.es

Abstract

In this paper we propose a propositional temporal language based on fuzzy temporal constraints which turns out to be expressive enough for domains –like many coming from medicine– where knowledge is of propositional nature and an explicit representation of time and uncertainty are required. The language is provided with a natural possibilistic semantics to account for the uncertainty issued by the fuzziness of temporal constraints. Also we present an inference system based on specific rules dealing with the temporal constraints and a general fuzzy modus ponens rule which behaviour is shown to be sound. The analysis of the different choices as fuzzy operators lead us to identify the well-known *Lukasiewicz implication* as very appropriate to define the notion *possibilistic entailment*, a central element of our inference system.

Keywords. temporal knowledge-based systems, uncertainty management, possibilistic logic, temporal constraints.

1 Introduction

Representation and Reasoning about time is an important representational issue to be considered in all those reasoning tasks in Artificial Intelligence which take account of a dynamic¹ domain [15]. Most of AI systems incorporating an explicit temporal representation are based in some manner on constraint-reasoning techniques [1, 13, 17, 3]. Dechter, Meiri & Pearl [4] provide a formal basis for some of these various schemes by defining the networks of metric temporal constraints based point-to-point constraints.

Metric temporal constraints are metric constraints among pairs of time-points representing the possible values for the temporal distance between time-points. A particular case is the *Simple Temporal Problem* (STP), each constraint $tc_{ij} = [a_{ij}, b_{ij}]$ -where a_{ij} and b_{ij} are values belonging to the variables domain- restricts the possible values for distance between a pair

of variables by

$$a_{ij} \leq X_j - X_i \leq b_{ij}$$

This sort of “simple” temporal constraint networks are a suitable support for temporal knowledge encoding. Metric temporal constraints have been recently used as the “temporal specialist” for general purpose reasoning systems. Schwalb *et al* [14] propose a propositional language of events and fluents which incorporates point-to-point, point-to-interval and interval-to-interval temporal constraints, although the characterized tractable core of this language includes only those constraints that are pointizable, i.e. can be transformed to point-to-point convex relations.

Escalada & Vila developed a Horn-like first-order language called the *Temporal Token Calculus* (TTC) [16], which incorporates conjunctions of simple point-to-point metric temporal constraints. TTC is provided with a complete inference system which includes a set of inference rules that completely axiomatize the simple temporal constraints. This work partially builds upon the TTC.

Metric point-to-point temporal constraints represent the set of possible values for the temporal distance between two points. The larger is the set of possible values, the higher is the degree of uncertainty about the actual value of the temporal distance. These values are all “equally possible values”.

Nonetheless, this is not enough for domains where knowledge about time is highly pervaded with vagueness and uncertainty. Let's illustrate it with an example taken from a medical domain. *Brucellosis* is one infectious pathology which may be the origin of serious lumbalgia problems. It has an evolution pattern which can be regarded as a particular instance of the common infectious evolution pattern. It is usually composed of an *inoculation event*[Inoculation], an *Initial Period*[IP] (with the common symptoms of headache, epistaxi, ...), a period of *Ondulating Fever*[OFP] and, finally, it reaches the state of an *Intervertebral Affection* [IA].² There is some (vague) knowledge about the temporal evolution of *brucellosis* cases:

*This work has been partially supported by Spanish MEC grant FPI-PN90 77909683 and by CICYT project ARREL (TIC92-0579-C02-01).

¹A dynamic domain is a domain where changes occur that modify the state of the world.

²This knowledge has been supplied by an expert on lumbalgia pathologies, the physician Antoni Llovera

- The *initial period* usually starts at a time between one and three weeks after the *inoculation event*, although extreme cases range from starting at the very inoculation time up to four weeks after;
- The *initial period* uses to be short: it last few days, although in unusual cases it may last up to two weeks at most;
- The *ondulating fever* period always occurs after the *initial period*. It uses to last between 20 and 25 days though other less possible cases range from 4 days as the lowest bound up to 45 days the uppest one.
- Finally, the *intervertebral affection* usually starts between 15 and 20 days after the start of the *ondulating fever* period. Nevertheless –supported by spine mechanical arguments– there are cases where the *intervertebral affection* did not appear after a long period which may last up to 22 months. Finally the *intervertebral affection* lasts more or less between 3 and 6 months. Extreme cases may range from the shortest case of 40 days and longer periods up to 12 months in cases where it is not properly treated.

Several pieces of work considered the representation of approximate temporal knowledge. Among them we would stress those based on the *possibility theory* [18, 10, 7, 12, 5, 2]. In particular, Barro *et al* propose an straight forward definition of the metric temporal constraints above mentioned based on fuzzy sets. This promising approach is called *Fuzzy Temporal Constraint Networks* [2].

We aim at defining an approximate temporal logic which permits to be provided with a computationally manageable inference engine. Thus, we propose a representation of uncertain temporal knowledge based on fuzzy temporal constraints. It allows us to separate the temporal component from the rest of the system to use temporal constraint satisfaction algorithms for it.

The goal of this paper is to investigate the embedding of fuzzy temporal constraints into a logical framework. For the sake of clarity, we have chosen a simple propositional language as the latter one. Dealing with fuzzy temporal constraints induces many-valued interpretations of our formulas, but inference from fuzzy constraints also induces uncertainty as soon as they represent a kind of incomplete information. This leads to extending the whole language to handle both fuzziness and uncertainty. Any well-formed formula will be associated with a certainty degree. Therefore, the natural framework where to model fuzziness and uncertainty in a unified way is the possibilistic approach, in accordance with using fuzzy sets for representing the temporal constraints.

The paper is organized as follows. In next section, we describe the syntax –which is illustrated by formalizing the example above introduced– and semantics of our basic language. In the third section an extended language and semantics to handle uncertainty is presented. In section 4 we present an inference system

and prove its soundness. Finally, we sketch future lines of work. Further motivation and proofs can be found in the long report [11].

2 The Basic Temporal Language

The atemporal part is simply made over a set of classical crisp propositions. The temporal part is based on fuzzy temporal constraints over a set of temporal propositions. It consists of the introduction of a single predicate FUZZDIST which states a fuzzy temporal constraint between a pair of time points. Whereas in temporal constraints networks a constraint is represented as an interval, now is represented by a (convex) fuzzy set of time points, inducing a possibility distribution. The link between the temporal propositions and the temporal constraints is performed through the duration-valued functions BEGIN(p) and END(p) that specify the initial and final instants of the period, the proposition p holds throughout.

Regarding the underlying time structure, and for the sake of simplicity, we take a fix interpretation of the set of duration symbols $\mathcal{DU}$ as the set of real numbers $\mathcal{R}$. Accordingly, the set $\mathcal{FDU}$ of fuzzy durations will be taken as the set of fuzzy subsets of $\mathcal{R}$. However, nothing would prevent to take other particular either *discrete* or *dense* group structures of time. For instance, sets of integers or rational numbers could also be possible models.

2.1 Syntax

Our propositional language $\mathcal{L}$ is a set of sentences over propositions and fuzzy temporal constraints. The following symbols are used:

$\mathcal{P}_\infty$, a set of primitive atemporal propositional variables,

$\mathcal{P}_t$, a set of primitive temporal propositional variables, and

t_0 , a special time point symbol.

We have two sorts of atomic formulas:

- **Atemporal Propositions:** the elements of $\mathcal{P}_\infty$.
- **Temporal constraints:** consist of a set of propositional variables (indexed by $\mathcal{P}_t \cup \{t_0\} \times \mathcal{P}_t \cup \{t_0\} \times \mathcal{FDU}$) of the form

$$\text{FUZZDIST}(t, t', \pi),$$

where t and t' are the temporal expressions BEGIN(p), END(q) or t_0 , being p and q some temporal propositional variables of $\mathcal{P}_t$, and $\pi \in \mathcal{FDU}$ is a fuzzy duration intended to represent a soft constraint on the duration temporal variable " $t' - t$ ".

The well-formed formulas (wffs) of our basic language $\mathcal{L}$ are of the form:

$$B_1 \text{ and } \dots B_i \dots \text{ and } B_m \implies H$$

with $m \geq 0$, where B_i and H are either literals from $\mathcal{P}_\infty$ or temporal constraints.

Like in *definite clauses*, the conjunction of B_i is called the *body* or *antecedent* and H is called the *head* or *conclusion*. We distinguish between two types of formula according to the form of the *body*, *fact*, when $m = 0$ (empty body) and written simply H , and *rule*, otherwise.

Figure 1 presents a graphical representation of the example presented in the introductory section.

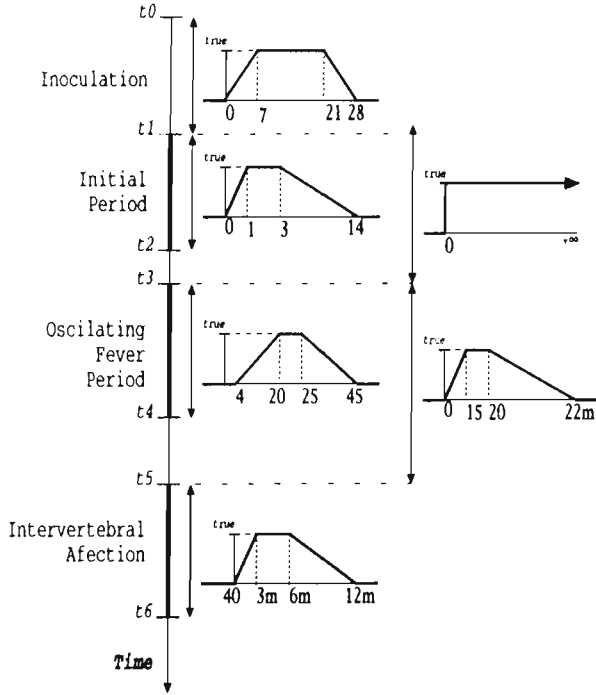


Figure 1: Brucellosis Evolution Pattern.

The *rule* for such a piece of knowledge would be as follows:

```

FUZZDIST(BEGIN(Inoculation), BEGIN(IP), [0, 7, 21, 27])
and
FUZZDIST(BEGIN(IP), END(IP), [0, 1, 3, 15])
and
FUZZDIST(BEGIN(OFP), END(OFP), [4, 20, 25, 45])
and
FUZZDIST(BEGIN(IP), BEGIN(OFP), [0, 0, +∞, +∞])
and
FUZZDIST(BEGIN(IA), END(IA), [40d, 3m, 6m, 12m])
and
FUZZDIST(BEGIN(OFP), BEGIN(IA), [0, 15, 20, 22m])
⇒
Brucellosis

```

2.2 Semantics

The semantics of our propositional language $\mathcal{L}$ involves, for each model, first the assignment of intervals

of time points to the temporal propositional variables, and second the interpretation in terms of truth-values of the formulas of the language, taking into account that atemporal propositional variables will be assigned to 1 (*True*) or 0 (*False*), while the temporal constraints expressions can be assigned any value of $[0, 1]$.

Definition 1 An $\mathcal{L}$ -model is a pair $\omega = \langle \omega_t, \omega_\infty \rangle$, where

$\omega_t : \mathcal{P}_t \rightarrow \text{PAIRS}(\mathcal{DU})$, where $\text{PAIRS}(\mathcal{DU})$ is the set of ordered pairs (x, y) from $\mathcal{DU}$ such that $x \leq y$, and

$\omega_\infty : \mathcal{P}_\infty \rightarrow \{0, 1\}$

Definition 2 Let Ω denote the set of $\mathcal{L}$ -models. The satisfaction relation $\models$ is a mapping

$$\models : \Omega \times \mathcal{L} \rightarrow [0, 1]$$

defined as:

$$\models (\omega, p) = \omega_\infty(p), \text{ if } p \in \mathcal{P}_\infty$$

$$\models (\omega, \neg p) = 1 - \omega_\infty(p), \text{ if } p \in \mathcal{P}_\infty$$

$$\models (\omega, \text{FUZZDIST}(t, t', \pi)) = \pi(M_\omega(t') - M_\omega(t)),$$

where M_ω is the interpretation function of temporal terms defined as

$$M_\omega(\text{BEGIN}(p)) = t_1, \text{ if } \omega_t(p) = (t_1, t_2)$$

$$M_\omega(\text{END}(p)) = t_2, \text{ if } \omega_t(p) = (t_1, t_2)$$

$$M_\omega(t_0) = 0$$

$$\models (\omega, P \text{ and } P') = \min(\models (\omega, P), \models (\omega, P')), \text{ and}$$

$$\models (\omega, B \Rightarrow H) = I_G(\models (\omega, B), \models (\omega, H)), \text{ where } I_G \text{ is the Gödel's implication function:}$$

$$I_G(x, y) = \begin{cases} 1 & x \leq y \\ y & x > y \end{cases}$$

From now on, and for the sake of a simpler notation, we will write $\omega(\phi)$ for $\models (\omega, \phi)$.

2.3 Fuzzy Temporal Constraint Inference Rules.

Now, we are interested in presenting a set of inference rules that deal exclusively with temporal constraints and that capture the completion operations performed in temporal constraints networks. Actually, they are the extension for fuzzy constraints of those proposed in [16] and, together with modus ponens, proved to be complete for crisp constraints. We will show their soundness w.r.t. the following notion of logical consequence in accordance to the well-known Zadeh's entailment principle in fuzzy logic.

Definition 3 A well-formed formula ϕ is valid, written $\models \phi$, iff $\omega(\phi) = 1$ for all $\omega \in \Omega$.

Definition 4 A well-formed formula ϕ is a logical consequence of a set of other well-formed formulas $\phi_1, \dots, \phi_n$, written $\{\phi_1, \dots, \phi_n\} \models \phi$ iff for any $\omega \in \Omega$, $\min(\omega(\phi_1), \dots, \omega(\phi_n)) \leq \omega(\phi)$.

Proposition 1 Let $\phi_1, \dots, \phi_n$ and ϕ be fuzzy temporal constraints. It holds that

$$\{\phi_1, \dots, \phi_n\} \models \phi \text{ iff } \models \phi_1 \text{ and } \dots \text{ and } \phi_n \Rightarrow \phi,$$

which is a kind of restricted deduction theorem.

The definition of the inference rules make use of the following two special fuzzy duration sets:

$$\text{Null duration function, } \pi_0(x) = \begin{cases} 1 & \text{if } x = 0 \\ 0 & \text{otherwise} \end{cases}$$

$$\text{Universal duration function, } \pi_U = 1$$

and of the following three operations on the set of fuzzy durations:

$$\text{Opposed, } \pi^-(x) = \pi(-x)$$

$$\text{Conjunction, } (\pi \cap \pi')(x) = \min\{\pi(x), \pi'(x)\}$$

$$\text{Composition, } (\pi \oplus \pi')(x) = \sup_{z=y+z} \{\min\{\pi(y), \pi'(z)\}\}$$

Next we present the inference rules.

$\overline{\text{FUZZDIST}(t, t, \pi_0)}$	{Reflexivity}
$\overline{\text{FUZZDIST}(t, t', \pi_U)}$	{Universal constraint}
$\frac{\text{FUZZDIST}(t_1, t_2, \pi_{d_{12}})}{\text{FUZZDIST}(t_2, t_1, \pi_{d_{12}})}$	{Symmetry}
$\frac{\text{FUZZDIST}(t_1, t_2, \pi_{d_{12}}), \text{FUZZDIST}(t_2, t_3, \pi_{d_{23}})}{\text{FUZZDIST}(t_1, t_3, \pi_{d_{12} \oplus d_{23}})}$	{Transitivity}
$\frac{\text{FUZZDIST}(t_1, t_2, \pi_{d_{12}}), \text{FUZZDIST}(t_1, t_2, \pi'_{d_{12}})}{\text{FUZZDIST}(t_1, t_2, \pi_{d_{12}} \cap \pi'_{d_{12}})}$	{Intersection}
$\frac{\text{FUZZDIST}(t_1, t_2, \pi_{d_{12}})}{\text{FUZZDIST}(t_1, t_2, \pi'_{d_{12}}), \text{ if } \pi_{d_{12}} \leq \pi'_{d_{12}}}$	{Inclusion}

Proposition 2 The above inference rules are sound w.r.t. the previous semantics.

3 The Possibilistic Temporal Language

To fully exploit the use of fuzzy expressions for temporal durations in the language, it seems very natural to also allow for partial degrees of matching between fuzzy expressions, and thus, leading to deal with partial degrees of certainty of propositions involving fuzzy duration constraints. On the other hand, even with non-fuzzy propositions, it is always useful for a representation language to allow dealing with uncertain knowledge. *Possibility theory* [6] offers a unified framework where to express both uncertainty and fuzziness. We present below an extension of the language in the

previous section where lower bounds of a necessity-like degree are attached to formulas, with a semantics based on ideas in [8, 9] that extend the one of Possibilistic Logic for crisp propositions.

We make more precise the above claim. Temporal constraint networks provide the tightest constraints between durations of events, that may be used as inputs in heuristic decision rules that may help in turn to complete the network. Therefore, when trying to apply decision rules, we are interested in certainty qualifying the conditions of such rules given for granted the constraints provided by the network. Next subsection is devoted to discuss how such certainty evaluation can be performed.

3.1 Certainty evaluation of fuzzy constraints

In technical terms, for a given duration variable X on $\mathcal{DU}$, we aim at finding the certainty evaluation of a fuzzy proposition $X \text{ is } A$ (the condition of a rule), being A a fuzzy subset of $\mathcal{DU}$, knowing that the values of X are restricted by a possibility distribution π (the constraint induced by the network). Dubois & Prade have discussed this issue in [8, 9] and they propose to use the following measure:

$$E(\mu_A|\pi) = \inf_{u \in \mathcal{DU}} \mathcal{I}_G^R(\pi(u), \mu_A(u))$$

where $\mathcal{I}_G^R$ is the reciprocal of the Gödel's implication, i.e.:

$$\mathcal{I}_G^R(x, y) = \mathcal{I}_G(1 - y, 1 - x)$$

and μ_A is the membership function of the fuzzy set A . It is remarkable to notice that $E(\mu_A|\pi) = 1$ iff $\mu_A \geq \pi$, and that $E(\mu_A|\pi)$ reduces, when A is non-fuzzy, to $N_\pi(A) = \inf\{1 - \pi(u) | u \notin A\}$, the necessity measure of A based on π and used in Possibilistic Logic. In fact, Dubois & Prade discuss in [8, 9] the inverse problem, that is, which possibility distribution corresponds to the semantical interpretation of " $X \text{ is } A$ " is α -certain", and to which evaluation of A is identified. This is also of much interest since it will allow us to use uncertain constraints derived from the decision rules as inputs in the network. An interesting line of argumentation leads them to represent the above qualified proposition by the following inequalities

$$\pi(u) \leq \max(\mu_A(u), 1 - \alpha), \forall u \in \mathcal{DU}$$

which, as expected, turn to be equivalent to

$$\mathcal{E}(\mu_A|\pi) \geq \alpha$$

However, the certainty degree $E(\mu_A|\pi)$ turns to provide not very natural results in very common situations. In particular, the existence of only one element u in the domain for which $\pi(u) = 1$ and $\mu_A(u) < 1$ causes the certainty degree $E(\mu_A|\pi)$ to be 0, independently whether the value $\mu_A(u)$ is close to 0 or close to 1.

This counter-intuitive behaviour has led us to look for an alternative definition of the certainty degree. If

one wants to keep the property that $E(\mu_A|\pi) = 1$ iff $\pi \leq \mu_A$, one is forced to stay either with residuated implications or with their reciprocals. Residuated implications, in general, share the problem that the resulting certainty degree does not collapse to the necessity degree in the non-fuzzy case (actually it becomes a trivial $\{0, 1\}$ -valued measure), and thus the resulting semantics is not an extension of that of Possibilistic Logic. On the other hand, the reciprocal implications, in general again, share the above mentioned problem of the Gödel's reciprocal implication. However, among these two families of implication functions, there is only one exception (up to isomorphisms), the well-known Łukasiewicz implication

$$\mathcal{I}_L(x, y) = \min(1, 1 - x + y),$$

that avoids the above problems. Namely, defining

$$E(\mu_A|\pi) = \inf_{u \in \mathcal{D}U} \mathcal{I}_L(\pi(u), \mu_A(u))$$

we keep most of the interesting properties of the previous definition while solving the main problem with it. Now the interpretation of "given π , (X is A) is (at least) α -certain" as $E(\mu_A|\pi) \geq \alpha$ is semantically equivalent to

$$\pi(u) \leq \mathcal{I}_L(\alpha, \mu_A(u)), \forall u \in \mathcal{D}U$$

This representation can be provided with practically the same argumentation used in [9] to justify their proposal, only a slight modification in one step is needed. The agreement of this proposal with the original one in the non-fuzzy case is easy to establish by noticing that $\mathcal{I}_L(\alpha, \mu_A(u)) = \max(1 - \alpha, \mu_A(u))$ when A is non-fuzzy.

3.2 Possibilistic Semantics

Now, we are prepared to define our *Possibilistic Temporal Language* and show that captures the above requirements.

Definition 5 *The set of possibilistic temporal formulas is the set $\mathcal{L}_T = \{\phi[\alpha] \mid \phi \in \mathcal{L} \text{ and } \alpha \in [0, 1]\}$.*

Definition 6 *A Possibilistic Temporal model Π is a possibility distribution over the set Ω of $\mathcal{L}$ -models $\Pi : \Omega \rightarrow [0, 1]$ ³*

Definition 7 (Possibilistic Entailment) *A possibilistic temporal model Π satisfies a wff ϕ with a certainty degree α , written $\Pi \models_T \phi[\alpha]$, iff*

$$E(\phi|\Pi) = \inf_{\omega \in \Omega} \mathcal{I}_L(\Pi(\omega), \mu_\phi(\omega)) \geq \alpha,$$

where we have identified the $\mathcal{L}$ -formula ϕ with its corresponding fuzzy subset of the set of $\mathcal{L}$ -models Ω defined as $\mu_\phi(\omega) = \omega(\phi)$. The notion of logical consequence is the natural one, i.e., a possibilistic temporal formula G is a logical consequence of a set of possibilistic temporal formulas $F_1, \dots, F_n$, written $\{F_1, \dots, F_n\} \models_T G$, iff for any possibilistic model Π , $\Pi \models_T F_1, \dots, \Pi \models_T F_n$ imply $\Pi \models_T G$.

³These possibilistic models differ from those of Possibilistic Logic in that the possibility distributions are defined on $[0, 1]$ -valued $\mathcal{L}$ -models, rather than on $\{0, 1\}$ -valued $\mathcal{L}$ -models.

The possibilistic entailment $\models_T$ in $\mathcal{L}_T$ is related to the entailment relation $\models$ of the basic language $\mathcal{L}$ without uncertainty as follows.

Proposition 3 *Let $\phi_1, \dots, \phi_n$ and ϕ well-formed formulas of $\mathcal{L}$. Then it holds that $\{\phi_1, \dots, \phi_n\} \models \phi$ iff $\{\phi_1[1], \dots, \phi_n[1]\} \models_T \phi[1]$*

This proposition shows that the possibilistic entailment $\models_T$ actually extends $\models$ in a natural way, as it could be expected. In particular, the set of inference rules for fuzzy constraints presented in section 2.3 are then also sound w.r.t. for $\models_T$, once the fuzzy temporal constraints appearing in those rules are attached the certainty value 1.

A *possibilistic temporal knowledge base* is a pair $KB = \langle \mathcal{FB}, \mathcal{RB} \rangle$ of a set of weighted facts $\mathcal{FB}$ and a set of weighted rules $\mathcal{RB}$.

4 General Inference

Since temporal inference rules deal exclusively with temporal constraints which certainty degree is 1, for the sake of completeness, additional inference rules are needed to state the *degree of fulfillment* for a temporal proposition and viceversa.

Basic Constraint Inference Rules. As already mentioned, the *Reflexivity*, *Universal Constraint*, *Symmetry*, *Transitivity*, *Intersection* and *Inclusion* inference rules, with the certainty value [1] attached to premises and conclusions, are sound rules w.r.t. to the possibilistic semantics, and they capture constraint network procedures

Possibilistic Constraint Inference Rules. The following inference rules show how uncertainty influences fuzzy temporal constraints, and thus how they provide a kind of bridge between knowledge from the temporal constraint network and knowledge from a decision rule set. In other words they provide a way to infer certain fuzzy constraints from uncertain ones, and viceversa.

$$\frac{\text{FUZZDIST}(t_1, t_2, \pi_{d_{12}})[1]}{\text{FUZZDIST}(t_1, t_2, \pi'_{d_{12}})[\alpha], \text{ where } \alpha = E(\pi_{d_{12}}|\pi'_{d_{12}})} \{R1\}$$

And the converse inference rule is defined as follows:

$$\frac{\text{FUZZDIST}(t_1, t_2, \pi_{d_{12}})[\alpha]}{\text{FUZZDIST}(t_1, t_2, \pi'_{d_{12}})[1], \text{ where } \pi'_{d_{12}} = \mathcal{I}_L(\alpha, \pi_{d_{12}})} \{R2\}$$

As a matter of example, consider the piece of knowledge that the duration of the *Ondulating Fever Period* of a patient has been between 2 and 3 weeks, but in any case not less than 17 days and not more than 32. This knowledge can be represented by the proposition $\text{FUZZDIST}(\text{BEGIN}(\text{OFP}), \text{END}(\text{OFP}), \Pi)$ with certainty 1, being π the trapezoidal function presented graphically in figure 2. The rule

in the previous example which has as one of conditions $\text{FUZZDIST}(\text{BEGIN}(\text{OFP}), \text{END}(\text{OFP}), A)$ being A the fuzzy duration whose membership function is also shown in the figure 2. Therefore to use that rule in this context we must evaluate, among other things, the certainty of that condition. Applying the $R1$ inference rule we can conclude $\text{FUZZDIST}(\text{BEGIN}(\text{OFP}), \text{END}(\text{OFP}), \pi_2)$, with certainty $E(\mu_A|\pi) = 0.9$. This certainty value would be used after to conclude *Brucellosis* from the rule when applying the modus ponens inference rule introduced below.

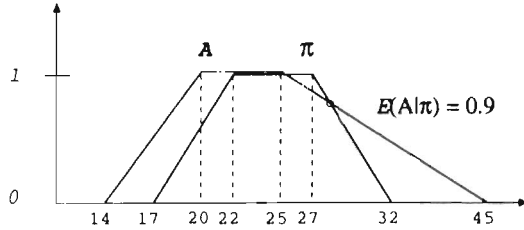


Figure 2: Using the entailment measure.

Proposition 4 *The above two inference rules $R1$ and $R2$ are sound.*

General Inference Rule. General inference is performed through a single inference rule which is a sort of *modus ponens* (PMP) and extends the one for necessity-valued formulas in Possibilistic Logic:

$$\frac{B_1 \text{ and } \dots \text{ and } B_m \xrightarrow{[\alpha]} H, \quad B_1[\alpha_1], \dots, B_m[\alpha_m]}{H[\alpha'], \text{ where } \alpha' = \min(\alpha, \alpha_1, \dots, \alpha_m)} \{PMP\}$$

Proposition 5 (soundness) *The above possibilistic modus ponens inference rule PMP is sound*

5 Conclusions and Future Work

We have presented in this paper a propositional temporal language based on fuzzy temporal constraints, and able to deal also with uncertainty within the possibilistic framework. Although this is a very restricted language, it turns out to be expressive enough for a large set of applications in the medical domain and, eventually, in other domains where knowledge is of propositional nature, yet explicit account of temporality and uncertainty are required.

Acknowledgments. We would like to thank Francesc Esteva, Pere García, Jordi Levy and Carles Sierra for comments on an earlier version of this paper. We are also thankful to the physician Antoni Llovera for his valuable cooperation as medical expert.

Bibliography

- [1] J. F. ALLEN. Maintaining knowledge about temporal intervals. Technical Report TR86, Department of Computer Science University of Rochester, 1983.
- [2] S. BARRO, R. MARIN, J. MIRA, and A. PATON. A model and a language for the fuzzy representation and handling of time. *Fuzzy Sets and Systems*, (to appear), 1993.
- [3] T. L. DEAN and D. V. McDERMOTT. Temporal data base management. *Artificial Intelligence*, 32:1987, 1987.
- [4] R. DECHTER, I. MEIRI, and J. PEARL. Temporal constraint networks. *Artificial Intelligence*, 49:61-95, 1991.
- [5] D. DUBOIS, J. LANG, and H. PRADE. Timed possibilistic logic. In Z. Ras, editor, *Fundamenta Informaticae Special Issue on Artificial Intelligence*. 1992.
- [6] D. DUBOIS and H. PRADE. *Possibility Theory*. Plenum Press, New York, 1988.
- [7] D. DUBOIS and H. PRADE. Processing fuzzy temporal knowledge. *IEEE Transactions on Systems, Man and Cybernetics*, 19(4), July/August 1989.
- [8] D. DUBOIS and H. PRADE. Resolution principles in possibilistic logic. *Intl. Journal of Approximate Reasoning*, 4(1):1-21, 1990.
- [9] D. DUBOIS and H. PRADE. Fuzzy rules in knowledge-based systems -modelling gradedness, uncertainty and preference. In R. Yager and L. Zadeh, editors, *An Introduction to Fuzzy Logic Applications in Intelligent Systems*. Kluwer Acad. Publ., Dordrecht, 1992.
- [10] S. DUTTA. An event-based fuzzy temporal logic. In *18th IEEE Int. Symp. on Multi-valued Logics*, pages 64-71, 1988.
- [11] L. GODO and L. VILA. A temporal reasoning system based on fuzzy temporal constraints. Report de Recerca forthcoming, IIIA - CSIC, 1994.
- [12] J. KOHLAS and P.-A. MONNEY. Temporal reasoning under uncertainty with belief functions. In *IPMU'90*, pages 67-73, 1990.
- [13] J. MALIK and T. O. BINFORD. Reasoning in time and space. In *IJCAI'83*, pages 343-345, 1983.
- [14] E. SCHWALB, K. KASK, and R. DECHTER. Temporal reasoning with constraints on fluents and events. In *to appear in KR'94*, 1994.
- [15] L. VILA. A survey on temporal reasoning in artificial intelligence. *AI Communications*, 7(1):4-28, March 1994.
- [16] L. VILA and G. ESCALADA-IMAZ. Temporal token calculus: a temporal reasoning approach for knowledge-based systems. In *GULP-PRODE'94: 1994 Joint Conference on Declarative Programming*, 1994. (in press).
- [17] M. VILAIN and H. KAUTZ. Constraint propagation algorithms for temporal reasoning. In *AAAI'86*, pages 377-382, 1986.
- [18] M. VITEK. Fuzzy information and fuzzy time. In *Proc. IFAC Symp. Fuzzy Information, Knowledge Representation and Decision Analysis*, pages 159-162, 1983.

Relaciones de indistinguibilidad descomponibles

J. Jacas

J. Recasens

E.T.S. d'Arquitectura de Barcelona
Univ. Politècnica de Catalunya
Diagonal 649. 08028 Barcelona.
e-mail: jacas@ea.upc.es

E.T.S. d'Arquitectura del Vallés
Univ. Politècnica de Catalunya
Sant Cugat del Vallés.
e-mail: recasens@ea.upc.es

RESUMEN¹

Las relaciones descomponibles son una generalización de la llamada implicación de Mandani utilizada en control. En este trabajo estudiamos en primer lugar la relación entre la propiedad de ser descomponible y las propiedades de simetría, transitividad y reflexividad. Finalmente se caracteriza las relaciones descomponibles que son indistinguibilidades.

Palabras clave: t-norma, relación difusa, implicación de Mandani, indistinguibilidad.

1 Introducción

Una de las cuestiones más importantes en el estudio de las relaciones de indistinguibilidad [2,3] es su generación. Un método usual (en procesos de aprendizaje) es partir de una relación reflexiva y simétrica y calcular su clausura transitiva. Otro método (útil en reconocimiento de formas) proviene del Teorema de representación que permite generar cualquier relación de indistinguibilidad a partir de una familia de conjuntos borrosos.

Algunas familias de relaciones de indistinguibilidad pueden ser generadas por otros métodos más acordes a su naturaleza. Tal es el caso de las indistinguibilidades de Mandani [5] que son un caso particular de relaciones descomponibles. En este artículo se estudian las relaciones descomponibles simétricas y de Mandani. En particular se demuestra que las relaciones de indistinguibilidad descomponibles (en el caso arquimediano estricto) satisfacen una "relación tetraédrica" similar a las relaciones de "estar entre" y que las relaciones que satisfacen una relación tetraédrica son *múltiplos* de relaciones descomponibles.

Recordemos las siguientes definiciones.

Definición 1.1 Fijada una t-norma T , una relación borrosa R en un conjunto X es **T-descomponible** si, y sólo si, existen dos subconjuntos borrosos μ, ν de X tales que para todo $x, y \in X$

$$R(x, y) = T(\mu x, \nu y)$$

¹Financiado parcialmente por el proyecto número PB91-0334-C03 de la DGICYT

Notación. A lo largo de todo el artículo μx y νy representarán $\mu(x)$ y $\nu(y)$ respectivamente.

Definición 1.2 Fijada una t-norma T , una T-indistinguibilidad en un conjunto X es una relación borrosa E en X que para todo x, y, z de X satisface

$$1.2.1. \text{ (Reflexividad)} \quad E(x, y) = 1$$

$$1.2.2. \text{ (Simetría)} \quad E(x, y) = E(y, x)$$

$$1.2.3. \text{ (T-transitividad)} \quad T(E(x, y), E(y, z)) \leq E(x, z).$$

El siguiente resultado es inmediato:

Proposición 1.3 Dada una t-norma T , toda relación R en un conjunto X T-descomponible es T-transitiva.

Demostración.

$$\begin{aligned} T(R(x, y), R(y, z)) &= T(T(\mu x, \nu y), T(\mu y, \nu z)) \\ &\leq T(\mu x, \nu z) = \\ &= R(x, z). \blacksquare \end{aligned}$$

2 Relaciones descomponibles simétricas

Una relación R descomponible en un conjunto X definida con un solo subconjunto difuso μ de X (i.e. $R(x, y) = T(\mu x, \mu y)$) es simétrica. En este apartado estudiamos el recíproco.

Proposición 2.1 Sea T una t-norma arquimediana estricta con generador aditivo t y $\hat{T}$ su pseudoinversa. La relación T-descomponible $R(x, y) = T(\mu x, \nu y)$ en X es simétrica si, y sólo si,

$$a) \mu \leq \nu \text{ y } \hat{T}(\nu|\mu) \text{ es una constante } k,$$

o bien

$$b) \nu \leq \mu \text{ y } \hat{T}(\mu|\nu) \text{ es una constante } k.$$

Demostración.

$\Rightarrow$) De la simetría de R se sigue

$$t\mu x + t\nu y = t\mu y + t\nu x, \quad \forall x, y \in X$$

Si $\mu x \leq \nu x$, entonces $t\mu x - t\nu x = t\mu y - t\nu y$ y, por tanto, $\mu y \leq \nu y$ (es decir $\mu \leq \nu$) y $\hat{T}(\nu | \mu) = k$ y estamos en el caso a).

Si $\mu x \geq \nu y$, un razonamiento análogo nos lleva al caso b).

$\Leftarrow$
Si $\mu \leq \nu$ y $\hat{T}(\nu | \mu) = k$, entonces

$$\begin{aligned} t^{-1}(t\mu x - t\nu y) &= k \\ t\mu x - t\nu y &= tk \quad \forall x \in X \end{aligned}$$

$$\begin{aligned} T(\mu x, \nu y) &= t^{-1}(t\mu x + t\nu y) = \\ &= t^{-1}(tk + t\nu x + t\mu y - tk) = \\ &= t^{-1}(t\mu y + t\nu x) = \\ &= T(\mu y, \nu x) \end{aligned}$$

La demostración para el caso $\nu \leq \mu$ y $\hat{T}(\mu | \nu) = k$ es análoga. ■

Notación. Dada una t-norma arquimediana estricta T y $a, b \in [0, 1]$ escribiremos $b = a^{1/2}$ si, y sólo si, $T(b, b) = a$.

Lema 2.2 $b = a^{1/2}$ si, y sólo si, $tb = \frac{ta}{2}$.

Demostración.

$$T(b, b) = a \Leftrightarrow t^{-1}(tb + tb) = a \Leftrightarrow 2tb = ta \Leftrightarrow tb = \frac{ta}{2}. \quad \blacksquare$$

Corolario 2.3 (de la proposición 2.1.) Una relación R T -descomponible en X con T una t-norma arquimediana estricta es simétrica si, y sólo si,

- a) $\mu \leq \nu$ y $R(x, y) = T(\rho x, \rho y)$ con $\rho = T(\nu, \hat{T}(\nu | \mu)^{1/2})$, o bien
- b) $\nu \leq \mu$ y $R(x, y) = T(\rho x, \rho y)$ con $\rho = T(\mu, \hat{T}(\mu | \nu)^{1/2})$.

Demostración.

$\Rightarrow$
Caso a) $\mu \leq \nu$.

$$\begin{aligned} \text{Por el lema, } \rho x &= t^{-1}(t\nu x + \frac{tk}{2}) \text{ y, por todo,} \\ R(x, y) &= T(\mu x, \nu y) = t^{-1}(t\mu x + t\nu y) = t^{-1}(tk + t\nu x + t\nu y) \\ &= t^{-1}(t\nu + \frac{tk}{2} + t\nu y + \frac{tk}{2}) = t^{-1}(t\rho x + t\rho y) = T(\rho x, \rho y) \end{aligned}$$

Caso b) Si $\nu \leq \mu$ el razonamiento es análogo al anterior

$\Leftarrow$
Trivial, puesto que T es simétrica. ■

Corolario 2.4 Si R es una relación T -descomponible métrica en X , con T una t-norma arquimediana estricta, se puede suponer, sin pérdida de generalidad, que $\mu = \nu$

Vamos a estudiar las relaciones descomponibles simétricas con la t-norma mínimo, para ello definimos $M = \sup_{x \in X} \{\mu x | \mu x = \nu x\}$ si este conjunto no es vacío y $M = 0$ en caso contrario.

Proposición 2.5 La relación Min -descomponible R en X definida por $R(x, y) = \text{Min}(\mu x, \nu y)$ es simétrica si, y sólo si,

- a) $\mu \leq \nu$ en los puntos x de X en los que $\mu x \neq \nu x$, μ es constante mayor o igual que M .
- b) $\nu \leq \mu$ en los puntos $x \in X$ en los que $\nu x \neq \mu x$, ν es constante mayor o igual que M .

Demostración.

$\Rightarrow$ Veamos primero que $\mu \leq \nu$ o $\nu \leq \mu$:

En efecto, si existieran $x, y \in X$ con $\mu x > \nu x$ y $\mu y < \nu y$, la igualdad

$$\text{Min}(\mu x, \nu y) = \text{Min}(\mu y, \nu x). \quad (2.5.1.)$$

contradiría los siguientes dos casos (i), (ii):

- (i) Si $\mu x \geq \mu y$, por 2.5.1. $\text{Min}(\mu x, \nu y) = \nu y$. Pero de $\nu y > \mu y$ se seguiría $\text{Min}(\mu x, \nu y) > \text{Min}(\mu y, \nu x)$.
- (ii) Si $\mu x \leq \mu y$, por 2.5.1. $\text{Min}(\mu y, \nu x) = \nu x$ y, puesto que $\mu y < \nu y$, para no contradecir 2.6.1 debería ser $\text{Min}(\mu x, \nu y) = \mu x$. Pero entonces se tendría $\mu x = \nu x$ contradiciendo la hipótesis de $\mu x > \nu x$.

Por tanto $\mu \leq \nu$ ó $\nu \leq \mu$.

(a) Caso $\mu \leq \nu$.

Si no existe $x \in X$ con $\mu x \neq \nu x$, entonces $\mu = \nu$ y hemos terminado.

Si existen $x, y \in X$ con $\mu x < \nu x$ y $\mu x < \nu y$, entonces $\text{Min}(\mu x, \nu y) \leq \mu x < \nu x$, y por 2.5.1. $\text{Min}(\mu y, \nu x) = \mu y$. De forma simétrica, $\text{Min}(\mu x, \nu y) = \mu x$ y por 2.5.1 $\mu x = \mu y$. Es decir, μ es constante en los puntos $x \in X$ donde $\mu x \neq \nu x$.

Si existe $z \in X$ con $\mu z = \nu z$, sea $x \in X$ con $\mu x < \nu x$. De $\text{Min}(\mu x, \nu z) = \text{Min}(\mu z, \nu x)$ se sigue entonces $\text{Min}(\mu z, \nu x) = \mu z$ y, por tanto, $\mu x \geq M$.

b) El caso $\nu \leq \mu$ es análogo al anterior.

$\Leftarrow$

Caso a):

Si $\mu x = \nu x$ y $\mu y = \nu y$, entonces $\text{Min}(\mu x, \nu y) = \text{Min}(\mu y, \nu x)$ trivialmente.

Si $\mu x = \nu x$ y $\mu y \neq \nu y$, entonces $\mu x = \nu x \leq \mu y < \nu y$ y por tanto $\text{Min}(\mu x, \nu y) = \text{Min}(\mu y, \nu x)$.

Si $\mu x \neq \nu x$ y $\mu y \neq \nu y$, entonces $\mu x = \mu y \leq \nu x$ y $\mu x = \mu y \leq \nu y$. Por tanto $\text{Min}(\mu x, \nu y) = \text{Min}(\mu y, \nu x)$.

Caso b) Análogo al anterior ■

Corolario 2.6 Si la relación Min-descomponible R en X definida por $R(x, y) = \text{Min}(\mu x, \nu y)$ es simétrica y $\mu \leq \nu$ ($\nu \leq \mu$), entonces $R(x, y) = \text{Min}(\mu x, \mu y)$ ($R(x, y) = \text{Min}(\nu x, \nu y)$).

3 Relaciones de indistinguibilidad descomponibles

En esta sección estudiamos las relaciones de indistinguibilidad descomponibles. El lema 3.1 nos obliga a modificar ligeramente la definición de descomponibilidad y nos lleva a considerar las indistinguibilidades de Mamdami [5].

Lema 3.1 Dada una t -norma T , una relación R T -descomponible en X , $R(x, y) = T(\mu x, \nu y)$, es reflexiva si, y sólo si $\mu = \nu \equiv 1$.

Demostración. Trivial. ■

Definición 3.2 Dada una t -norma T , una relación R definida en un conjunto X diremos que es T' -descomponible si, y sólo si, existen dos subconjuntos borrosos μ, ν de X tales que para cada x, y de X

$$R(x, y) = T'(\mu x, \nu y)$$

$$\text{donde } T'(a, b) = \begin{cases} T(a, b), & \text{si } a \neq b \\ 1, & \text{si } a = b \end{cases}$$

Proposición 3.3 Dado un subconjunto borroso μ de X y una t -norma R , la relación borrosa $E(x, y) = T'(\mu x, \mu y)$ es una T -indistinguibilidad en X .

Demostración. Trivial. ■

Lema 3.4 Una T -indistinguibilidad E en X T' -descomponible con $E(x, y) = T'(\mu x, \mu y)$ es separadora (i.e. $E(x, y) = 1 \Rightarrow x = y$) si, y sólo si, μ es inyectiva.

Demostración.

$$E(x, y) = 1 \Leftrightarrow T'(\mu x, \mu y) = 1 \Leftrightarrow \mu x = \mu y. \quad \blacksquare$$

En [3] se demuestra que dada una T -indistinguibilidad separadora E en un conjunto X con $E(x, y) = 0 \forall x, y \in X$ y T arquimediana estricta, la relación clásica ternaria B en X definida por $(x, y, z) \in B$ si, y sólo si, $T(E(x, y), E(y, z)) = E(x, z)$ es una relación de “estar entre” métrica en el sentido de Menger [4].

Las relaciones de “estar entre” definidas por T -indistinguibilidades separadoras descomponibles son de dos tipos: o bien $B = \emptyset$ o bien existe un único elemento de X que está entre cualesquiera otros dos

Proposición 3.5 Sea T una t -norma arquimediana estricta y μ un subconjunto difuso inyectivo en X con $\mu x \neq 1$ o $\forall x \in X$. La relación de “estar entre” B definida por la T -indistinguibilidad $E(x, y) = T'(\mu x, \mu y)$ es vacía

Demostración.

Dados $x, y, z \in X$ con $x \neq y \neq z \neq x$, si $T(E(x, y), E(y, z)) = E(x, z)$, entonces $T(T'(\mu x, \mu y), T'(\mu y, \mu z)) = T(\mu x, \mu x)$ lo que implica $\mu y = 1$.

Corolario 3.6 Sea T una t -norma arquimediana estricta y μ un subconjunto difuso inyectivo en X con $\mu x \neq 0 \forall x \in X$. Sea E la T -indistinguibilidad definida por $E(x, y) = T'(\mu x, \mu y)$.

Si existe un elemento $y \in X$ con $\mu y = 1$, entonces y está entre cualesquiera otros dos elementos de X y estas son las únicas relaciones de “estar entre” definidas por E .

Por tanto, las indistinguibilidades del corolario anterior dotan al conjunto X de una estructura radial en la que todos los objetos están interconectados a través de un objeto central privilegiado. Las indistinguibilidades descomponibles generan un tipo de relaciones similares a las de “estar entre” que se pueden asociar a estructuras tetraédricas.

Proposición 3.7 Dada una t -norma arquimediana estricta T y un subconjunto borroso inyectivo μ en X , la T -indistinguibilidad separadora R en X definida por $E(x, y) = T'(\mu x, \mu y)$ genera la siguiente relación tetraédrica: dados cuatro elementos x, y, z, t de X distintos dos a dos,

$$T(E(x, y), E(z, t)) = T(E(x, z), E(y, t)). \quad (3.7.1)$$

Demostración.

$$\begin{aligned} & t^{-1}(tt^{-1}(t\mu x + t\mu y) + tt^{-1}(t\mu z + t\mu t)) = \\ & t^{-1}(tt^{-1}(t\mu x + t\mu z) + tt^{-1}(t\mu y + t\mu t)). \quad \blacksquare \end{aligned}$$

La Proposición anterior tiene el siguiente recíproco parcial:

Proposición 3.8 Sea T una t -norma arquimediana estricta y E una T -indistinguibilidad separadora en C que satisfaga la relación tetraédrica 3.7.1. Entonces existe una T -indistinguibilidad E' en X T' -descomponible y una constante $k \in [0, 1]$ con $E(x, y) = \hat{T}(k \mid E'(x, y)) \forall x, y \in X, x \neq y$.

Para la demostración usamos el siguiente lema:

Lema 3.9 Dados $a, b, x \in [0, 1]$ con $a \geq b$ y T arquimediana estricta, se tiene

$$\hat{T}(a \mid b) = \hat{T}(T(a, x), T(b, x)).$$

Demostración de la proposición 3.8.

Fijemos $a \in X$ y definamos $\rho x = E(a, x)$.

Dados y, z, t de X diferentes dos a dos y diferentes de a ,

$$T(E(a, y), E(z, t)) = T(E(a, z), E(y, t))$$

$$T(\rho y, E(z, t)) = T(\rho z, E(y, t))$$

Podemos suponer $\rho y \leq \rho z$ y la última igualdad equivale a

$$\hat{T}(\rho z \mid \rho y) = \hat{T}(E(z, t), E(y, t))$$

y, por el lema 3.9

$$\hat{T}(T(\rho t, \rho z) \mid T(\rho t, \rho y)) = T(E(z, t) \mid E(y, t))$$

Por tanto,

$$T(\rho y, \rho t) = T(E(y, t), k)$$

Definiendo

$$\rho'x = \begin{cases} \rho x, & \text{si } x \neq a \\ k, & \text{si } x = a, \end{cases}$$

$$E' = T'(\rho', \rho'). \blacksquare$$

Nota 1. Si $k = 1$, se tiene $E' = E$, $\rho' = \rho$ y el elemento $a \in X$ está entre cualesquiera otros dos de X . Nos encontramos por lo tanto en el caso del corolario 3.6.

La constante k se puede pensar como el grado de adecuación de tomar a como parte central del corolario 3.6.

Nota 2. En la demostración anterior sólo hemos usado la relación tetraédrica para un elemento a de X .

De esto se siguen los dos siguientes corolarios:

Corolario 3.10 Sea T una t -norma arquimediana estricta y E una T -indistinguibilidad separadora en X . Si existe $a \in X$ con $T(E(a, y), E(z, t)) = T(E(a, z), E(y, t)) \forall x, y, z \in X$ distintos dos a dos y distintos de a , entonces existe una T -indistinguibilidad T' -descomponible E' y una constante $k \in [0, 1]$ con $E(x, y) = \hat{T}(k \mid E'(x, y)) \forall x, y \in X, x \neq y$.

Corolario 3.11 Sea T una t -norma arquimediana estricta y E una T -indistinguibilidad separadora en X . Si existe $a \in X$ con $T(E(a, y), E(z, t)) = T(E(a, z), E(x, y)) \forall x, y, z \in X$ distintos dos a dos y distintos de a , entonces $T(E(x, y), E(z, t)) = T(E(x, z), E(y, t)) \forall x, y, z, t \in X$ distintos dos a dos.

Las similitudes (Min-indistinguibilidades) descomponibles coinciden con las similitudes unidimensionales en el sentido de [2]. Aplicando las técnicas allí expuestas, se demuestra la siguiente proposición:

Proposición 3.12 Una similitud E en un conjunto X separadora es Min'-descomponible si, y sólo si, existe $a \in X$ con $E(a, \cdot)$ inyectiva.

Demostración. [2]

4 Conclusiones

En este artículo hemos estudiado las relaciones simétricas descomponibles y las indistinguibilidades de Mandani para t -normas estrictas y el mínimo.

En el caso de las t -normas estrictas se ha visto que, sin pérdida de generalidad, podemos suponer que los dos conjuntos generadores coinciden (i.e. $R = T(\mu \times \mu)$) y que las indistinguibilidades de Mandani satisfacen una "relación tetraédrica" que las caracteriza salvo un factor de escala k .

Para el Mínimo, las similitudes descomponibles coinciden con las ya estudiadas similitudes unidimensionales y se caracterizan a partir de resultados conocidos.

Bibliografía

- [1] FERNÁNDEZ, M.J. AND SUAREZ, F. Eigen Fuzzy sets asociados a relaciones simétricas y descomponibles. *Actas del 3r. Congreso Español sobre Tecnologías y Lógica Fuzzy*. Santiago de Compostela (1993). 41-50.
- [2] JACAS, J., On the generators of a T -indistinguishability operator. *Stochastica* 12, (1988), 46-63.
- [3] RECASENS, J. *Sobre la representació i generació de relacions d'indistingibilitat*. Tesis doctoral U.P.C., (1992).
- [4] TRILLAS, E. AND ALSINA, C. Introducción a los Espacios Métricos Generalizados. *Fund. J. March. Serie Universitaria* 49 Madrid, (1978).
- [5] TRILLAS, E., CUBILLO, S. AND RODRIGUEZ, A. Preórdenes borrosos, indistinguibilidades y estados lógicos borrosos. *Actas del 3r. Congreso Español sobre Tecnologías y Lógica Fuzzy*. Santiago de Compostela (1993), 57-64.

Una aproximación al razonamiento temporal borroso en medicina

R. Marín¹, R. Ruiz¹, F. Martín¹

S. Barro², F. Palacios³

¹Dept. de Informática y Sistemas
Universidad de Murcia
Campus de Espinardo
30071 Murcia, España
e-mail: roque@fac1.dif.um.es

²Dept. de Electrónica e Computación
Universidade de Santiago
15706 Santiago, España
e-mail: senen@gaes.usc.es
³Hospital General de Elche

Resumen

En este trabajo proponemos una solución para la aplicación de un modelo de restricciones temporales borrosas a tareas de monitorización inteligente de pacientes. En particular, analizamos las necesidades del dominio y describimos la aproximación seguida para representar la información temporal y para resolver las consultas sobre relaciones temporales planteadas durante la aplicación del conocimiento sobre el dominio. Estos procesos exigen la cooperación de un especialista temporal con los mecanismos de razonamiento sobre el dominio y son el punto clave para la integración entre ambos. La solución para la integración presentada aquí corresponde a una fase de la implementación de FRESH, una herramienta de desarrollo de sistemas inteligentes para monitorización en entornos dinámicos.

Palabras clave:

Razonamiento temporal, Sistemas expertos, Sistemas de monitorización, Cuidados Intensivos Coronarios.

1 Introducción

En trabajos anteriores [3, 4, 13, 14, 15] hemos presentado un modelo de restricciones temporales borrosas (FTCN o Fuzzy Temporal Constraint Network), como punto de partida para manejar la imprecisión de las referencias al tiempo que aparecen con frecuencia en dominios médicos. Estamos interesados en la aplicación del modelo teórico a tareas de monitorización de pacientes. Describimos aquí algunos de los problemas encontrados y las soluciones propuestas.

Una estrategia general para la aplicación práctica de técnicas de razonamiento temporal consiste en establecer una separación explícita entre las tareas de razonamiento sobre el *dominio* y de razonamiento sobre el *tiempo*. Las primeras se resuelven mediante la arquitectura

convencional de un sistema experto (SE) basado en conocimiento del dominio, mientras las segundas las implementa un *especialista temporal*, un sistema basado en algún modelo formal de restricciones temporales que representa y gestiona la información referente a la localización de un conjunto de eventos sobre el eje del tiempo. Las tareas básicas de un especialista temporal son: (1) inferir nuevas relaciones temporales a partir de las inicialmente representadas y reducir al mínimo la imprecisión y la ambigüedad en las relaciones conocidas, aplicando para ello algún mecanismo de propagación de restricciones temporales; (2) analizar la consistencia global del conjunto de relaciones parciales introducidas; y (3) computar posibles escenarios temporales (ordenaciones totales de los eventos) que sean compatibles con la información de partida. Esta última es una tarea básica en problemas de planificación, pero es irrelevante en problemas de diagnóstico e interpretación de datos de monitorización. En estos casos, el SE del dominio debe detectar patrones temporales de síntomas que sustenten o rechacen diferentes hipótesis diagnósticas, y para ello consultará al especialista temporal sobre la existencia de relaciones temporales concretas entre hechos concretos; el especialista temporal responderá a estas consultas tras aplicar los algoritmos de análisis correspondientes a las tareas (1) y (2) a las etiquetas temporales de los hechos.

Puesto que los objetos manejados son las etiquetas temporales asociadas a los eventos, el especialista temporal se limita a realizar razonamiento de sentido común *sobre* el tiempo, y en principio es independiente del dominio. Esta división de tareas fue propuesta por primera vez por Kahn y Gorry [11], reconociendo que la absoluta independencia del dominio puede ser difícil de conseguir en la práctica. La integración de un especialista temporal con un mecanismo de razonamiento sobre el dominio, en forma suficientemente flexible para adaptarse

a distintas aplicaciones, es aún un problema abierto. Cada modelo de restricciones temporales proporciona una solución de compromiso entre expresividad y eficiencia, que no es necesariamente válida para todos los dominios [1, 2, 8, 17]. Una solución suficientemente flexible para un rango dado de aplicaciones debe permitir expresar distintos tipos de etiquetas temporales y gestionar cada tipo en la forma más eficiente. Pero el punto más crítico para establecer la separación de tareas es el de la gestión de consultas sobre existencia de relaciones temporales. Este último problema es el que analizamos en este trabajo. En la siguiente sección describimos los conceptos básicos del modelo FTCN. Después nos centraremos en un dominio específico y haremos algunas consideraciones sobre el tipo de información temporal manejada. Finalmente, describiremos la solución propuesta para los procesos de consultas sobre eventos.

2 Representación de relaciones temporales borrosas

Una red de restricciones temporales borrosas (FTCN o Fuzzy Temporal Constraint Network) es un conjunto finito de $n+1$ variables temporales $\{X_0, X_1, \dots, X_n\}$ y un conjunto finito de restricciones binarias borrosas entre ellas $\{L_{ij}\}$. Cada variable representa a un instante de tiempo desconocido y cada restricción binaria L_{ij} describe los posibles valores para la duración transcurrida entre los instantes X_i y X_j , y se define mediante una distribución de posibilidad π_{ij} sobre el universo de las duraciones enteras. Denominamos *extensiones temporales* a estas distribuciones de posibilidad, porque representan la distancia temporal existente entre dos puntos del eje del tiempo: dada una duración entera precisa p , $\pi_{ij}(p)$ representa la posibilidad de que la duración transcurrida de X_i a X_j sea precisamente p unidades de tiempo [3, 4, 13, 14, 15].

Consideremos el problema de la representación de entidades y relaciones temporales en el modelo FTCN. Sea un evento e_j que corresponde a la ocurrencia de algún hecho del dominio en un tiempo dado, y que representaremos por una terna $e_j = (a_j, v_j, t_j)$, en la que a_j es alguna variable del dominio (e.g. "tensión arterial diastólica"), v_j es el valor preciso o impreciso que toma esta variable (e.g. "alta"), y t_j es el tiempo en el que el hecho es válido. Si se trata de un evento instantáneo, la entidad temporal t_j será un instante de tiempo p_j , y en el caso de un evento con duración t_j será igual a un intervalo i_j .

Un instante puede representarse directamente en el modelo mediante una pareja $p_j = (X_j, R_j)$, donde X_j es una variable temporal y R_j es un conjunto de restricciones

temporales $\{L_{j_1}, L_{j_2}, \dots\}$ que establecen la posición relativa del instante respecto a los instantes correspondientes a otras variables temporales $X_{j_1}, X_{j_2}, \dots$. Por ejemplo, un evento instantáneo e_1 , del que sabemos que ocurrió en el año 1990 aproximadamente 15 días antes del evento instantáneo e_2 , tendrá asociada una etiqueta temporal t_1 dada por una variable X_1 y un conjunto de dos restricciones temporales $R_1 = \{L_{10}, L_{12}\}$; L_{10} es una restricción respecto al origen de tiempos que expresa la fecha absoluta imprecisa ("el año 1990"), y en la que la imprecisión está asociada al uso de una granularidad temporal alta; L_{12} es una restricción temporal imprecisa ("aproximadamente 15 días antes") respecto a la variable temporal X_2 asociada al instante de ocurrencia del evento de referencia e_2 .

Por otra parte, un intervalo i_j se representa en el modelo mediante una terna $i_j = (p_{j_a}, p_{j_b}, R_{j_{ab}})$, en la que p_{j_a} y p_{j_b} corresponden, respectivamente, a los instantes inicial y final del intervalo, que se representan en la forma indicada anteriormente, y $R_{j_{ab}}$ es una relación temporal entre las variables temporales correspondientes a esos instantes, que representa la duración del intervalo.

Lo anterior permite representar en la red la localización temporal de cualquier evento. Sin embargo, interesa disponer de una forma de expresión más general y cercana al usuario. Para el usuario, cada entidad temporal se define mediante un conjunto de etiquetas temporales lingüísticas, $p_j = \{l_{j_1}, l_{j_2}, \dots\}$ o $i_j = \{l_{j_1}, l_{j_2}, \dots\}$. En otro trabajo [3] hemos definido las reglas sintácticas y semánticas de un *lenguaje de representación de entidades temporales* que traduce etiquetas temporales lingüísticas a la representación de la red. El lenguaje permite expresiones del tipo " *p_2 es un instante que ocurre aproximadamente 10 minutos después del instante p_1 y pocos minutos antes del intervalo i_3* ", donde las expresiones en cursiva corresponden a etiquetas temporales lingüísticas. En líneas generales, las etiquetas lingüísticas pueden venir dadas en forma absoluta (e.g. una fecha de calendario) o relativa. Las etiquetas relativas definen relaciones temporales entre instantes, entre intervalos, y entre instantes e intervalos.

Consideramos tres tipos de relaciones entre instantes: $\{d\}$ -antes, $\{d\}$ -después y $\{d\}$ -igual, donde d es una duración borrosa que cuantifica la distancia temporal imprecisa existente entre los instantes. Etiquetas de este tipo se representan directamente en el modelo como restricciones entre parejas de variables temporales. Hay diez tipos de relaciones entre instantes e intervalos (e.g. $\{d\}$ -comienzo, $\{d_s, d_f\}$ -pertenece, ...) Cada etiqueta de este tipo se traduce sobre la red en tres variables temporales (la del instante y las dos del intervalo) y en dos restricciones temporales (entre el instante y los puntos

inicial y final del intervalo); si el intervalo había sido definido previamente existirá además una restricción entre sus extremos). Análogamente, las relaciones entre intervalos pueden representarse mediante 4-tuplas de variables temporales y 4-tuplas de restricciones binarias entre estas variables. En el caso particular de relaciones cualitativas (*antes, después, ...*), las duraciones borrosas *d* se reducen a distribuciones universales, de forma que cualquier valor preciso positivo de duración es completamente posible.

En definitiva, cualquier expresión lingüística es traducida a un "trozo" de red que constituye su representación interna. En la siguiente sección justificaremos que este esquema de representación es válido para un dominio clínico particular, analizando el tipo de etiquetas temporales que aparecen asociadas a los hechos clínicos del dominio. Veremos también como el problema del razonamiento sobre el dominio implica la realización de consultas sobre relaciones temporales. En otra sección posterior veremos cómo este tipo de consultas se resuelven sobre en relación al modelo.

3 Razonamiento temporal en dominios clínicos

Cualquier proceso fisiopatológico puede describirse como una evolución temporal de un sistema dinámico, el cuerpo humano, cuyo resultado observable es una determinada secuencia de manifestaciones clínicas. El significado clínico de un síntoma aislado es frecuentemente ambiguo, y debe establecerse dentro del contexto temporal en el que ocurre. Un contexto temporal es definido por un intervalo de tiempo en el que, para que el síntoma sea asignable a una categoría diagnóstica dada, deben haber ocurrido otras determinadas manifestaciones en un cierto orden temporal. Este orden no tiene que ser necesariamente un orden total: los síntomas no se presentan siempre exactamente en la misma secuencia; lo significativo es la ordenación relativa de algunos subconjuntos de síntomas y, particularmente, que la distancia temporal entre síntomas sea concordante con las relaciones de causalidad del proceso fisiológico subyacente.

Consideremos un dominio clínico particular, el seguimiento de pacientes con isquemia ingresados en una UCIC (Unidad de Cuidados Intensivos Coronarios). El objetivo básico es detectar la aparición de nuevos episodios isquémicos en pacientes que ya han sufrido infarto agudo de miocardio y están sometidos a monitorización electrocardiográfica y hemodinámica.

La isquemia se manifiesta a través de cambios en el nivel eléctrico del segmento ST de la señal ECG. Estos cambios sólo son significativos si vienen precedidos por

la existencia de factores desencadenantes que provocan un aumento en el consumo de oxígeno y que se manifiestan poco después como alteraciones hemodinámicas y de la frecuencia cardíaca. La detección de estos eventos en los minutos previos a los cambios del ST es importante para descartar falsos positivos y evitar el disparo de alarmas irrelevantes. En ocasiones, estas manifestaciones vienen acompañadas de síntomas subjetivos (e.g. dolor), pero en un gran porcentaje de casos el episodio isquémico es silente y, por tanto, difícil de detectar sin la ayuda de un sistema de monitorización. Una primera corroboración del diagnóstico se realiza observando las modificaciones en la presión arterial diastólica y en la presión de enclavamiento que acompañan y siguen a los cambios del ST. Entre 3 y 8 horas después aparecen alteraciones analíticas que permiten confirmar el diagnóstico y decidir si se mantiene la medicación. Los antecedentes clínicos del paciente y la existencia de tratamiento previo pueden alterar este proceso y condicionar la interpretación de resultados.

Un sistema de monitorización convencional maneja sólo los datos de bajo nivel resultantes del procesado de señales fisiológicas. Sin embargo, la correcta interpretación clínica de los eventos de señal exige analizarlos en un nivel de abstracción superior, en correlación con datos de otro tipo. Los parámetros más frecuentemente usados en una UCI corresponden a observaciones clínicas (síntomas y signos), antecedentes, medicación y analítica, mientras que los parámetros de monitorización corresponden sólo a un 13 % del total [10].

Las etiquetas temporales asociadas a los distintos tipos de datos pueden ser de muy distinta naturaleza. En general, los antecedentes clínicos vienen fechados en forma aproximada y con un nivel de granularidad temporal superior al de otros datos, lo que equivale a imprecisión temporal (e.g. "insuficiencia renal moderada diagnosticada hace poco más de un año"). Tanto en los antecedentes como durante la monitorización hay que representar síntomas subjetivos definidos mediante etiquetas temporales relativas a algún punto de referencia (e.g. "palpitaciones desde unos 5 meses antes del ingreso", "dolor poco después de la última comida"). Los datos de analítica y de exploración clínica tienen etiquetas temporales absolutas y precisas. Finalmente, los parámetros extraídos de las señales monitorizadas tienen también etiquetas absolutas y precisas, generadas a partir de los índices de las muestras, que son relativos al origen de la monitorización. En consecuencia, el especialista temporal tiene que manejar etiquetas temporales absolutas y relativas, precisas e imprecisas. Todas ellas pueden ser representadas en nuestro modelo FTCN.

Consideremos ahora el problema del razonamiento

sobre el dominio. Desde el punto de vista operativo, la identificación del significado clínico de un evento se realiza examinando su posición temporal relativa a otros síntomas precedentes, es decir, comprobando la existencia de determinadas relaciones temporales entre los hechos. Ejemplos de consultas sobre relaciones temporales entre hechos del dominio médico para determinar su significado clínico aparecen en [12, 18]. Sin embargo, el conocimiento sobre procesos fisiopatológicos es imperfecto y la variabilidad entre pacientes es grande, de forma que no es posible hacer una especificación temporal precisa de los eventos genéricos que caracterizan un determinado problema clínico. En la práctica, ello implica que las reglas de producción incluirán condiciones sobre relaciones temporales borrosas (e.g. "si los últimos análisis de función renal fueron normales y ocurrieron hace menos de aproximadamente 15 días entonces descartar una agudización de la insuficiencia renal y sugerir una patología cardíaca"). El interfaz entre los mecanismos de razonamiento sobre el dominio y el especialista temporal deberá separar las consultas estrictamente temporales y transmitir las al especialista temporal para su resolución. Este proceso es el que esquematizamos en la siguiente sección.

4 Consulta sobre relaciones temporales borrosas en FRESH

FRESH (Fuzzy Temporal Reasoning Expert System Shell) es una herramienta de desarrollo de sistemas inteligentes para monitorización en entornos dinámicos que nuestro grupo está implementando. El trabajo se plantea en dos fases. En la fase actual, hemos abordado la construcción de un sistema experto para el dominio particular descrito previamente. El trabajo en este dominio operacional nos servirá para probar la validez de nuestra aproximación antes de generalizar a otros dominios. Describimos aquí sólo la solución adoptada para responder a las consultas sobre relaciones temporales borrosas que los mecanismos de razonamiento de FRESH plantean al especialista temporal. Algunos de los mecanismos de razonamiento particulares que se integrarán en FRESH están descritos en otros trabajos [5, 6, 7, 16].

En FRESH, los eventos se almacenan en una base de datos orientada a objetos. El objeto correspondiente a cada parámetro clínico a contiene una historia de eventos $H(a) = \{e_1, e_2, \dots\}$, en la que cada evento e_j es definido por un par valor/entidad temporal (v_j, t_j) . Las entidades temporales t_j apuntan hacia los nodos de una FTCN, y su correspondencia sobre la red se obtiene mediante el lenguaje de representación de entidades temporales en la forma descrita en la segunda sección. El especialista temporal descompone la red en un conjunto de topologías

particulares, y aplica a cada una algoritmos específicos de propagación de restricciones y análisis de consistencia que se ejecutan en tiempo polinómico. Dadas dos variables temporales X_i y X_j , llamamos M_{ij} a la restricción existente entre ellas después de aplicar los algoritmos de propagación y π_{ij}^M a la correspondiente distribución de posibilidad. M_{ij} será una restricción mínimamente imprecisa resultante de combinar todas las piezas de información que afectan a la distancia temporal entre ambas variables.

Las preguntas planteadas por los mecanismos de razonamiento del dominio son de la forma (Q, S, C_v, C_t) . S es una selección de eventos de la forma:

$S = \{e \mid \text{tipo}(e) \wedge R_v(e.\text{valor}, \text{vref}) \wedge R_t(e.\text{tiempo}, \text{tref})\}$, que corresponde al conjunto de eventos de un tipo dado cuyos valores y tiempos verifican ciertas condiciones. R_v es una condición de valor que compara el valor del evento con un valor de referencia. Análogamente, R_t es una condición de tiempo que compara el tiempo del evento con un tiempo de referencia, y viene dada por cualquiera de las expresiones lingüísticas generables mediante el mismo lenguaje de representación de entidades temporales descrito en la segunda sección. Aisladamente, S es una fórmula no cuantificada con una variable libre. Q es un operador de agregación que se aplica a los eventos del conjunto selección, y puede ser un cuantificador ($\exists, \forall$), o un operador del tipo máximo, mínimo, promedio, primero, último, intervalo maximal, ... Q se aplica al valor o al tiempo de cada uno de los eventos del conjunto S . C_v y C_t son condiciones de valor y tiempo, respectivamente, similares a las incluidas dentro del conjunto S . Esta sintaxis permite plantear preguntas del tipo: ¿Todos los {eventos de tensión arterial sistólica mayor que 100 mmHg y durante el intervalo del último episodio de isquemia} han ocurrido menos de aproximadamente 3 minutos antes del último cambio ST? La expresión entre corchetes es la que corresponde al conjunto selección; para cada uno de los eventos seleccionados hay que comprobar la condición de tiempo final, que corresponde a una relación temporal con el intervalo de tiempo asociado al evento de referencia "último episodio de isquemia". Preguntas como ésta son las que figuran en la parte de condiciones de las reglas.

Un intérprete de consultas descompone estas expresiones en preguntas a la base de datos (sobre el valor de los eventos) y preguntas al especialista temporal (sobre la localización temporal relativa de los eventos). En particular, el especialista temporal sólo recibe una secuencia de consultas sobre condiciones de tiempo $C_i = R_i(t_1, t_2)$ donde t_1 es la entidad temporal asociada a algún evento concreto del conjunto selección y t_2 es la entidad temporal asociada al evento de referencia (quizá definida

recursivamente mediante otro conjunto selección y localizada previamente). R_i puede ser cualquier relación temporal entre instantes, entre instantes e intervalos, o entre intervalos. El especialista temporal descompone cada una de estas consultas sobre condiciones de tiempo en consultas temporales elementales. Cada consulta temporal elemental comprueba si se verifica una cierta restricción entre variables de la red. Por ejemplo, una consulta sobre una relación R_i entre dos intervalos es descompuesta por el especialista temporal en cuatro consultas elementales sobre las restricciones entre los variables inicial y final de los intervalos.

Una consulta elemental es de la forma $\{X_i, D_{ij}, X_j\}$. Para resolverla hay que comprobar si entre los instantes correspondientes a los nodos X_i y X_j ha transcurrido una duración D_{ij} . Mediante el mismo lenguaje de representación de entidades temporales usado para la entrada de información se obtiene la distribución de posibilidad π_{ij}^D asociada a esa duración por la que se pregunta. Por otra parte, se conoce la duración "real" transcurrida entre los nodos, que corresponde a la restricción mínima M_{ij} y está representada por la distribución de posibilidad π_{ij}^M . La restricción M_{ij} habrá sido obtenida por los algoritmos de propagación a partir de las piezas de información de entrada. Para responder a la consulta es necesario comprobar la similitud entre π_{ij}^D y π_{ij}^M . Formalmente, se trata de comprobar en qué medida el contenido de la proposición "la distancia temporal entre X_i y X_j es D_{ij} " es compatible con la proposición de referencia "la distancia temporal entre X_i y X_j es M_{ij} ". El especialista temporal aplica las expresiones de filtraje borroso de Teoría de la Posibilidad propuestas por Dubois y Prade [9]. En el contexto de nuestro modelo, estas expresiones se particularizan a:

$$\Pi(X_i, D_{ij}, X_j) = \max_{m \in I} \min\{\pi_{ij}^D(m), \pi_{ij}^M(m)\},$$

$$N(X_i, D_{ij}, X_j) = \min_{m \in I} \max\{\pi_{ij}^D(m), 1 - \pi_{ij}^M(m)\}.$$

y obtienen una medida de posibilidad y una medida de necesidad, respectivamente, que definen una banda de estimación del grado en que la proposición consultada es cierta. En particular, la relación consultada será posiblemente cierta si y sólo si hay una solución de la red (una asignación de tiempos a los eventos) que satisfice la relación con grado de posibilidad no nulo.

5 Conclusiones

En este trabajo hemos discutido la aplicabilidad de nuestro modelo de restricciones temporales borrosas (FTCN) a las necesidades de gestión de información temporal presentes en un dominio clínico particular, la monitorización de pacientes en una UCIC. La solución

propuesta se basa en un especialista temporal que incluye un lenguaje de representación de entidades temporales, una red de restricciones temporales borrosas, algoritmos de descomposición de la red en topologías particulares, y algoritmos de propagación de restricciones y análisis de consistencia. Hemos mostrado las líneas básicas de los procesos de representación de información temporal y de resolución de las consultas planteadas por los mecanismos de razonamiento del dominio. El especialista temporal es una pieza básica de la herramienta FRESH para desarrollo de sistemas de monitorización inteligente.

Referencias

- [1] ALLEN, J. Time and time again: the many ways to represent time. *International Journal of Intelligent Systems* 6 (1991), 341-355.
- [2] BARBER, F. A metric time-point and duration-based temporal model, *SIGART*, 4, 3 (1993), 30-49.
- [3] BARRO, S., MARÍN, R., MIRA, J., AND PATON, A. R. A model and a language for the fuzzy representation and handling of time, *Fuzzy Sets and Systems*, 61 (1994), 153-175.
- [4] BARRO, S., MARIN, R., OTERO, R. P., RUIZ, R., AND MIRA, J. On the handling of time in intelligent monitoring of CCU patients. *Proc. of the 14th Annual International Conference of the IEEE Engineering in Medicine and Biology Society* (1992), 871-873.
- [5] BUGARIN, A. *Soluciones para la computación eficiente en sistemas basados en lógica borrosa*. Universidad de Santiago, Santiago de Compostela, (1994). (Tesis Doctoral).
- [6] BUGARIN, A., AND BARRO, S. Fuzzy reasoning supported by Petri nets. *IEEE Transactions on Fuzzy Systems*, (1994). En prensa.
- [7] BUGARIN, A., BARRO, S., AND RUIZ, R. Fuzzy control architectures. *Journal of Intelligent and Fuzzy Systems*, (1994). En prensa.
- [8] DECHTER, R., MEIRI, I., AND PEARL, J. Temporal constraint networks, *Artificial Intelligence* 49 (1991), 61-95.

- [9] DUBOIS, D., AND PRADE, H. *Possibility theory. An approach to computerized processing of uncertainty*. Plenum Press, 1988.
- [10] GARDNER, R. *Patient Monitoring Systems*. In *Medical Informatics: Computer Applications in Health Care*. H. Shortliffe and L. Perreault (eds.) Addison-Wesley, 1990.
- [11] KAHN, K., AND GORRY, A. Mechanizing temporal knowledge. *Artificial Intelligence* 9 (1977), 87-108.
- [12] KERAVNOU, K. Medical temporal reasoning. *Artificial Intelligence in Medicine* 3 (1991), 289-290.
- [13] MARIN, R., BARRO, S., BOSCH, A., AND MIRA, J. Modeling time representation from a fuzzy perspective. *Proc. of EUROCAST'93* (1993), 81-84.
- [14] MARIN, R., BARRO, S., BOSCH, A., NAVARRETE, I., CARDENAS, M. BUGARIN, A., AND MIRA, J. Un modelo para la representación de información temporal borrosa. *Actas del Tercer Congreso Español sobre Tecnologías y Lógica Fuzzy* (1993), 209-216.
- [15] MARIN, R., BARRO, S., BOSCH, A., AND MIRA, J. Modeling time representation from a fuzzy perspective. *Cybernetics and Systems* 25, 2 (1994), 207-215.
- [16] PRESEDO, J. *Monitorización Inteligente de Isquemia en SUTIL*. Universidad de Santiago, Santiago de Compostela, 1994. (Tesis Doctoral).
- [17] VAN BEEK, P. Approximation algorithms for temporal reasoning. *Proc. of the IJCAI-89* (1989) 1291-1296.
- [18] VAN BEEK, P. Temporal query processing with indefinite information. *Artificial Intelligence in Medicine* 3 (1991), 325-339.

Análisis de los rasgos característicos de las funciones de implicación según su comportamiento en un sistema borroso.

José Villar Collado

Universidad Pontificia Comillas
Instituto de investigación tecnológica
Fernando El Católico 63 Dup
28015 Madrid
e-mail: jose@iit.upco.es

Miguel Angel Sanz Bobi

Universidad Pontificia Comillas
Instituto de investigación tecnológica
Fernando El Católico 63 Dup
28015 Madrid
e-mail: masanz@iit.upco.es

Resumen

En este artículo hemos realizado un estudio del comportamiento de distintas funciones de implicación en un sistema borroso. Este sistema se ha utilizado como modelo matemático con objeto de analizar el error de estimación cometido, observar el método de desborrosificación más adecuado y el número de muestras apropiadas a la función de implicación escogida. A partir de este estudio hemos clasificado las funciones de implicación según cuatro formas de comportamiento, identificando rasgos de comportamiento con propiedades geométricas, permitiendo de este modo generar nuevas familias de implicaciones con comportamientos conocidos a priori y adaptados al problema concreto al que se aplican. Finalmente y a modo de ejemplo proponemos una nueva familia de implicaciones.

Palabras clave:

Sistema borroso, implicación borrosa, inferencia borrosa, modelado borroso.

1 Introducción

Uno de los problemas fundamentales que se plantean en el diseño de un sistema borroso es la elección de los distintos operadores que intervienen en él de entre la multitud de alternativas de que se dispone, y en particular de la función de implicación, sobre la que centraremos este trabajo.

Aunque la formulación teórica que soporta el proceso de inferencia no suele ser discutida, ésta propone un marco de trabajo tan amplio y flexible que no aporta ninguna orientación a la hora de seleccionar los operadores, durante el diseño de un sistema borroso.

Debido a esta carencia, multitud de investigadores han realizado y presentado estudios experimentales con vistas a poder obtener unos mínimos criterios de selección en

problemas muy concretos. Existen muchas de referencias a trabajos experimentales concretos donde se obtienen soluciones a problemas particulares, pero no existen muchos trabajos de carácter más general, que faciliten la elección ([1], [4], [7], etc).

La referencia [8] nos sirvió como punto de partida para proponer una metodología sencilla y general que permitiese probar y comparar operadores sin restringirse a casos particulares.

4 Análisis del error de estimación y clasificación de las implicaciones

4.1 Descripción del problema

Para comparar las funciones de implicación entre sí hemos utilizado un sistema borroso como estimador funcional de $y=x$. Tenemos por tanto dos variables, x e y con sus correspondientes universos de discurso X e Y , donde definimos un conjunto de reglas (X_i, Y_i) , con $i = 1:N$. Hemos utilizado $N=5$. El método de evaluación del estimador consiste en calcular para cada valor de entrada de la variable x , x_0 , el valor estimado para la variable y , y_0 , y compararlo con el valor real y_0 .

El cálculo del valor estimado de y está basado en el siguiente esquema de inferencia:

$$y_0 = DEFZ_{nmy} \left\{ \bigcup_i I(X_i(x_0), Y_i(y)) \right\}$$

donde $I(a,b)$ es la función de implicación, y $DEFZ_{nmy}$ es el método de desborrosificación empleado, que depende del número de muestras nmy utilizado para discretizar el universo de discurso de la variable y . Analizaremos el comportamiento del centro de gravedad, que notaremos como $cgrd$, y el de la media de los valores con máximo grado de pertenencia, que notaremos $maxp$, en relación con la función de implicación que se escoja.

La sencillez del esquema de inferencia permite eliminar la conocida composición *sup-T*, de forma que las funciones de implicación han podido ser comparadas sin preocuparnos de la optimalidad de dicha composición para la implicación escogida, que como se sabe están estrechamente relacionadas.

El operador utilizado para la agregación de reglas simbolizada por la unión en la ecuación anterior, ha sido el siguiente. Siendo $I(a,b)$ la función de implicación:

-operador *min* si $I(0,0)=1$. (reglas con antecedentes que no se cumplen generan distribuciones resultantes iguales al universo de discurso del consecuente).

-operador *max* si $I(0,0)=0$ (reglas con antecedentes que no se cumplen generan distribuciones resultantes iguales al conjunto vacío).

Obviamente este criterio está fundamentado en la distinta procedencia de las implicaciones, llamadas así en sentido amplio, ya que en el primer caso estamos modelando mediante una conjunción de implicaciones, mientras que en el segundo lo hacemos mediante una unión de conjunciones.

Las etiquetas utilizadas han sido distribuciones triangulares, todas ellas con la misma forma y distribuidas de manera que la suma de los grados de pertenencia de un punto cualquiera a las etiquetas de su universo es igual a 1, y dicho punto nunca pertenece con grado superior a 0 a más de dos etiquetas.

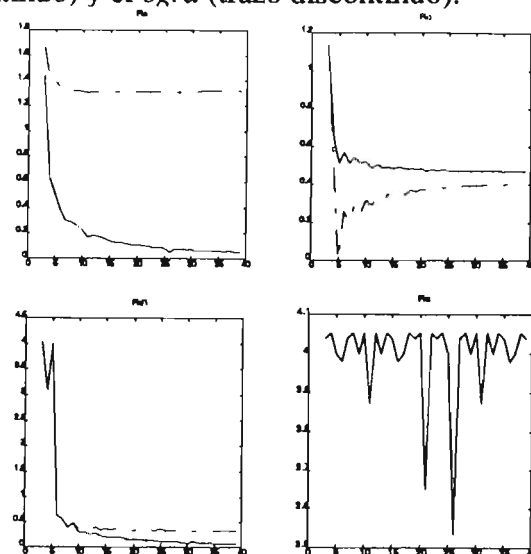
El índice de comparación utilizado ha sido la raíz del error cuadrático medio de estimación. El estimador ha sido evaluado en un número elevado de puntos, para evitar la dependencia artificial del error con el número de muestras del universo X .

Cuando una distribución resultante del proceso de inferencia tiene todos sus grados de pertenencia iguales a 0, el criterio de desborrosificación adoptado ha sido el cálculo del punto medio del universo de discurso.

4.2 Clasificación de las funciones de implicación.

Los resultados obtenidos al analizar los errores cometidos en función del número de muestras de la variable y , quedan resumidos en las figuras que a continuación aparecen. Como pretendemos mostrar en este trabajo, las funciones de implicación se pueden agrupar en cuatro grandes grupos, de los cuales hemos tomado un representante (Ra , Rc , $Rd1$ y Rs). En las figuras aparece la evolución de la raíz del error cuadrático medio en función del número

de muestras utilizado para discretizar el universo de discurso de la variable y , utilizando el método de desborrosificación *maxp* (trazo continuo) y el *cgrd* (trazo discontinuo).



Vamos a continuación a caracterizar cada uno de los tipos señalando aquellas características que consideramos más relevantes. Posteriormente intentaremos justificar los resultados aquí observados.

Implicaciones tipo 1

$Ra(x,y)=\min(1,1-x+y)$
$Ra1(x,y)=\min(1,1-(x-y)^2)$
$Ra2(x,y)=\min(1,1-\text{abs}(x-y))$
$Ra11(x,y)=Ra(x,y)$ si $y > x$, $Ra1(x,y)$ si $y < x$
$Rb(x,y)=\max(1-x,y)$
$Rast(x,y)=1-x+x.y$
$Rnum(x,y)=\min(\max(1-x,y),\max(x,1-x),\max(y,1-y))$
$Ra3(x,y)=Ra(x,y)$ si $y > x$, $\min(1,1-\text{sqrt}(\text{abs}(x-y)))$ si $x > y$
$Rm(x,y)=\max(\min(x,y),1-x)$

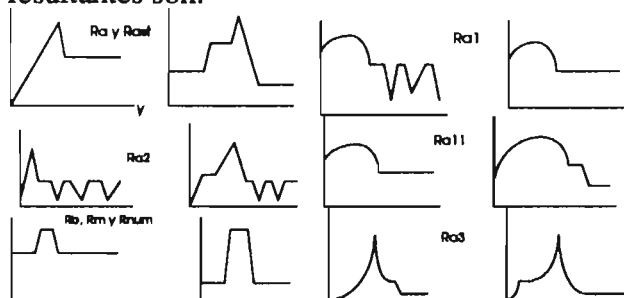
Las implicaciones de tipo 1 están caracterizadas por un error pequeño y de evolución suave y decreciente cuando se utiliza el método *maxp*, y por un error similar pero superior con el método *cgrd*.

Todas ellas presentan una superficie continua tal que $I(0,0)=1$, $I(0^+,0)=1$, $I(0,0^+)=1$.

Estas implicaciones deben usarse con el método *maxp*, y el número de muestras n_{my} debe ser superior a 20 ó 30 muestras para que el error sea suficientemente pequeño. Si el número de muestras tiende a infinito el error cometido en este caso concreto tiende a 0. Se comportan mejor aquellas implicaciones cuyo valor en la recta $a=b$ se acerca tanto más a 1. Respecto al método *cgrd* se comportan mejor aquellas implicaciones cuyo valor en la zona $a < b$ sea

menor, sin que dicho valor afecte al comportamiento del método *maxp*.

Las formas típicas de las distribuciones resultantes son:



Rb, *Rm* y *Rnum* generan distribuciones aparentemente apropiadas para el método *maxp* pero cuando varían los valores de *x* la distribución no cambia de sitio, solo varían sus grados de pertenencia, por lo que los resultados finales no son buenos.

Las demás implicaciones generan distribuciones aproximadamente centradas sobre el valor *y* deseado, por lo que el método *maxp* se comporta de forma correcta. Por último señalar que la implicación que genera distribuciones más acotadas para grados de pertenencia pequeños es la implicación *Ra3*, que es la que tiene menor valor en la zona $a < b$. Además la simetría de sus distribuciones resultantes para grados de pertenencia suficientemente grandes hace a la implicación bastante insensible al número de muestras que se utilice ante ambos métodos de desborrosificación.

Implicaciones tipo 2

$Rc(x,y) = \min(x,y)$
$Rprod(x,y) = x \cdot y$
$Rbp(x,y) = \max(0, x+y-1)$

Las implicaciones de tipo 2 presenta un error de evolución suave y decreciente con el método *cgrd*, y error algo mayor pero similar con el método *maxp*. El error presenta un mínimo con el método *cgrd* (en este ejemplo de valor 0) cuando *nmy* es igual al número de etiquetas de la partición de la variable *y*.

Las superficies de las implicaciones de tipo 2 son también continuas y relativamente suaves, y tales que $I(0,0)=0$, $I(0^+,0)=I(0,0^+)=0$.

Estas implicaciones deben de utilizarse con el método *cgrd*, y el número de muestras debería ser igual al número de etiquetas de la variable *y*. Cuanto menor es el valor de la implicación mejor ha sido su comportamiento en los ensayos realizados en este trabajo.

Las formas típicas de las distribuciones que resultan de la inferencia son:



Para este tipo de implicaciones, a medida que el grado de cumplimiento de una regla va disminuyendo y el de la regla adyacente va aumentando, el pico de la distribución no se desplaza sino que va apareciendo progresivamente un apéndice lateral en la distribución cuya altura aumenta en detrimento del pico inicial. Por ello el método *maxp* no funciona de forma correcta.

Implicaciones tipo 3

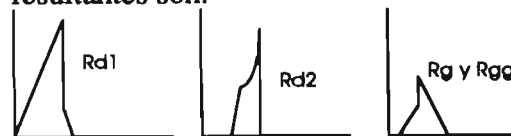
$Rd1(x,y) = 1$ si $x \notin y$, 0
$Rd2(x,y) = \min(1, y/x, (1-x)/(1-y))$ si $x, (1-y) > 0$, 1
$Rg(x,y) = 1$ si $x \in y$, y
$Rgg(x,y) = \min(Rg(x,y), Rg(1-x, 1-y))$

Las implicaciones de tipo 3 presentan errores pequeños con ambos métodos, ligeramente superiores a los obtenidos con los tipos 1 y 2, y con evoluciones también suaves. Ambos métodos de desborrosificación ofrecen resultados muy semejantes aunque parece ser ligeramente mejor el método *maxp*.

Las superficies de estas implicaciones presentan todas una discontinuidad en el punto (0,0), cumpliéndose $I(0,0)=1$, $I(0,0^+)=1$, $I(0^+,0)=0$.

Estas implicaciones pueden usarse casi indistintamente con uno u otro método de desborrosificación, con un número de muestras aconsejablemente superior a 15 para ambos métodos.

Las formas típicas de las distribuciones resultantes son:



No presentan zonas de indeterminación (es decir zonas horizontales) como ocurría con las implicaciones de tipo 1, lo que hace que su comportamiento frente al *cgrd* sea bueno. Por otro lado el pico de la distribución resultante se centra aproximadamente sobre el valor esperado de la *y* y exceptuando *Rd2* suelen tener una simetría razonable lo que garantiza su buen comportamiento con el método *maxp* y su robustez frente al efecto del muestreo.

Implicaciones tipo 4

$Rs(x,y) = 1$ si $x \in y$, 0
$Rsg(x,y) = \min(Rs(x,y), Rg(1-x, 1-y))$
$Rgs(x,y) = \min(Rg(x,y), Rs(1-x, 1-y))$
$Rss(x,y) = \min(Rs(x,y), Rs(1-x, 1-y))$
$Rcuad(x,y) = 1$ si $x < y$ o $y=1$, 0

Hemos agrupado dentro de las implicaciones de *tipo 4* todas aquellas cuyo mal comportamiento las hace desaconsejables en el ejemplo que nos ocupa. Están caracterizadas por errores que, independientemente del método de desborrosificación empleado, son elevados y con evoluciones bruscas e impredecibles.

Rsg, *Rgs* y *Rss* presentan todas ellas una discontinuidad que provoca que las distribuciones resultantes sean conjuntos críps con un solo elemento, siendo entonces desastroso el efecto del muestreo. *Rcuad* genera distribuciones resultantes indeterminadas (es decir el universo de discurso) por lo que el valor estimado es siempre el mismo.

4.3 Generalidad de los resultados anteriores

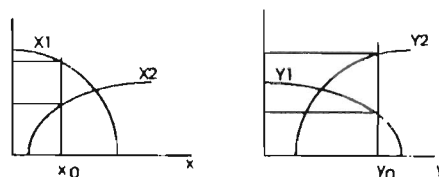
Las simulaciones se repitieron utilizando estimadores borrosos para otras funciones, en concreto funciones tan dispares como x^2 y $\cos(x)$. Las reglas se obtuvieron tomando una a una las etiquetas de la variable x y asociándole la etiqueta de la variable y para la que los pares (x_i, y_i) habían dado mayores grados de pertenencia conjunta, $X_j(x_i), Y_j(y_i)$.

Los resultados obtenidos, aunque numéricamente peores debido a la mala calidad de los estimadores, fueron cualitativamente idénticos a los hasta ahora relatados. Entre las diferencias cuantitativas importantes podemos señalar que para las implicaciones de *tipo 1* y *3* el error con el método *maxp* no tiende a 0 aunque sí a un mínimo al aumentar el número de muestras *nmy*, y que para las implicaciones de *tipo 2* el error se hace mínimo pero no 0 cuando *nmy* es igual al número de etiquetas de la variable sobre la que se concluye.

Las conclusiones siguen manteniéndose también cuando se varía el número de etiquetas utilizado en el modelo, cumpliéndose como era de esperar que a mayor número menor es el error cometido.

4.4 Justificación de los resultados anteriores

Vamos a intentar en este apartado justificar algunos de los resultados anteriores, estudiando cómo se desarrolla la inferencia cuando tenemos dos reglas adyacentes, $X1 \rightarrow Y1$ y $X2 \rightarrow Y2$, y cuando las observaciones son de tipo crisp, es decir parejas de datos (x_i, y_i) . Por reglas adyacentes entendemos que $X1$ se solapa con $X2$ e $Y1$ con $Y2$. Podemos representar gráficamente esta situación mediante la siguiente figura:



La distribución de posibilidad resultante para la variable y , dado un valor x_0 de la variable x está dada por la ecuación:

$$Y'(y) = \min(I(X_1(x_0), Y_1(y)), I(X_2(x_0), Y_2(y)))$$

y el valor $Y'(y_0)$ se obtiene particularizando la ecuación anterior para $y=y_0$ y $x=x_0$.

Supondremos, sin pérdida de generalidad, que las etiquetas están ordenadas como aparecen en la figura anterior.

Con estas hipótesis, para $y < y_0$ se tiene:

$$I(X_1(x_0), Y_1(y)) \geq I(X_1(x_0), Y_1(y_0))$$

$$I(X_2(x_0), Y_2(y_0)) \leq I(X_2(x_0), Y_2(y))$$

Esto se cumple debido a los requerimientos que hemos exigido a las etiquetas de las particiones (si dichas etiquetas no fuesen unimodales y convexas, las desigualdades anteriores no se cumplirían) y supuesto que la implicación cumple la siguiente desigualdad, que coincide con uno de los requerimientos lógicos habituales:

$$I(a, b) \geq I(a, b') \text{ si } b \geq b'$$

ya que entonces se tiene:

$$Y_1(y_0) < Y_1(y) \text{ e } Y_2(y) < Y_2(y_0)$$

Por tanto la distribución de posibilidad resultante, en la región $y < y_0$ está dada por:

$$Y'(y) = I(X_2(x_0), Y_2(y))$$

Análogamente, para $y > y_0$ se tiene:

$$Y'(y) = I(X_1(x_0), Y_1(y))$$

Una condición que permite que el método *maxp* funcione de forma óptima es que la distribución de posibilidad tenga máximo grado de pertenencia en y_0 , y grados menores para cualquier $y \neq y_0$, esto es:

$$Y'(y_0) > Y'(y) \text{ para } y \neq y_0$$

Podemos ver que esta condición se cumple si se cumplen las siguientes condiciones:

$$I(a, b) = 1 \text{ para } b = a, I(a, b) < 1 \text{ para } a > b$$

$$Y_i(y = f(x)) = X_i(x)$$

ya que entonces se tiene:

$$Y_i(y_0) = X_i(x_0) \text{ e } Y_i(y) < X_i(x) \text{ para } (x, y) \neq (x_0, y_0)$$

y por tanto:

$$\min(I(X_1(x_0), Y_1(y_0)), I(X_2(x_0), Y_2(y_0))) =$$

$$I(X_1(x_0), Y_1(y_0)) =$$

$$I(X_2(x_0), Y_2(y_0)) = 1 > I(X_2(x_0), Y_2(y))$$

y análogamente ocurre para $y > y_0$.

En el caso del estimador para $y=x$, la propiedad $Y_i(y=f(x))=X_i(x)$ se cumple con exactitud. En el caso de los otros estimadores esta propiedad se cumple solo para algunos puntos aislados, debido a que no se ha realizado

ningún ajuste de las formas de las funciones de pertenencia.

Podemos decir por tanto que las implicaciones que cumplen las propiedades:

$$I(a, b=a)=1$$

$$I(a, b) < 1 \text{ si } a > b$$

$$I(a, b) \geq I(a, b') \text{ si } b \geq b'$$

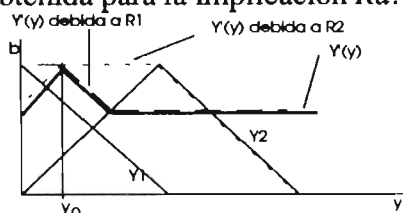
funcionan de forma idéntica con el método maxp, cuando las particiones reúnen las condiciones de solape mencionadas y se cumple $Y_i(y=f(x))=X_i(x)$, supuesto además un número infinito de muestras en el proceso de cálculo del método maxp. Por ello la forma de la función de implicación en la zona $b > a$ no influye en el resultado de la estimación.

Respecto al cálculo de la desborrosificación, cuanto más estrecha sea la distribución resultante, mayor será el número de muestras necesario para garantizar un error pequeño. Para minimizar el efecto del muestreo podemos imponer una nueva condición para que el ancho de la distribución resultante sea algo mayor:.

$$I(a, b)=1 \text{ e } I(a, b) < 1 \text{ para } b < a \text{ y } b=a$$

Con este nuevo requisito el ancho de las distribuciones resultantes aumenta, principalmente en la zona donde se tienen los máximos grados de pertenencia, de forma que un muestreo menos fino puede conducir a resultados suficientemente satisfactorios.

Vamos ahora a centrarnos en las asimetrías que aparecen en las distribuciones resultantes. Son estos efectos los que disminuyen la eficacia del método cgrd como desborrosificador. Si analizamos el origen de las asimetrías, podemos ver que éstas aparecen principalmente debido a las zonas de indeterminación que se originan en las distribuciones resultantes, en las zonas donde los grados de pertenencia no son elevados, como se puede ver en la siguiente figura obtenida para la implicación R_a :



Las zonas de indeterminación aparecen debido a que, para la regla R_i definida por las etiquetas X_i e Y_i , y para $b=0$ (es decir $Y_i(y)=0$) se tiene en general $I(a, b)=I(X_i(x_0), 0) \neq 0$ si $X_i(x_0) \neq 0$. Podemos por tanto imponer una nueva condición que equivale a la condición de simetría pero muy relajada al aplicarse solo a ciertos valores de a y b , para conseguir disminuir las asimetrías de las distribuciones resultantes, y que podemos formular como:

$$I(a, b) \rightarrow 0 \text{ para } a \rightarrow 0 \text{ y } b=0$$

Podemos comprobar que las implicaciones de tipo 3 cumplen esta propiedad, razón por la cual presentan tan buenos comportamientos con el método cgrd. Sin embargo al no obtenerse una total simetría el método cgrd siempre arroja peores resultados que el método maxp.

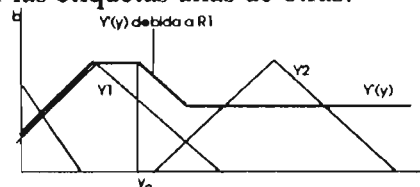
Por último señalar que para que la implicación siga respetando los requerimientos lógicos habituales basta imponer dos condiciones más, perfectamente compatibles con las hasta ahora enunciadas, y que son:

$$I(1, b)=b$$

$$I(a, b)=1 \text{ para } b > a$$

4.5 Optimalidad de las particiones borrosas.

Supongamos que existen zonas donde la variable solo pertenece a una etiqueta y con grado menor que 1. Esta condición equivale a separar las etiquetas unas de otras:



En las zonas donde las etiquetas se solapan el sistema sigue funcionando correctamente. Sin embargo, y como se puede apreciar en la figura, en las zonas donde el solape es 0, se generan distribuciones con grados de pertenencia 1 para valores de y distintos de y_0 , que además no están centradas en y_0 por lo que el método maxp genera un valor estimado para y_0 desplazado respecto al real. Esto no parece por tanto deseable. Cuando el solape es mayor, se puede comprobar que el método maxp sigue funcionando correctamente (es decir con error 0 para un infinito número de muestras). Por tanto no interesa que las variables pertenezcan a una única etiqueta, salvo en puntos aislados en donde debe cumplirse $X_i(x_0)=Y_i(y_0)$.

Cuando las etiquetas presentan zonas de más de un elemento con grado de pertenencia 1, por ejemplo etiquetas trapezoidales las distribuciones resultantes presentan α -cortes de valor 1 con más de un elemento, no centrados en el valor esperado de y_0 , y por tanto aumenta el error cuando se usa el método maxp. Para etiquetas trapezoidales, el error depende del ancho de la base superior de los trapecios. Cuanto menor sea menor es el error de estimación.

4.6 Familia de funciones de implicación

Las propiedades siguientes caracterizan a una implicación con un buen comportamiento ante

ambos métodos de desborrosificación, y robusta frente al muestreo.

$$\begin{aligned} I(a,b) &= 1 \text{ para } b \geq a \\ I(a,b) &< 1 \text{ si } a > b \\ I(a,b) &\geq I(a,b') \text{ si } b \geq b' \\ I(a,b) &\approx 1 \text{ e } I(a,b) < 1 \text{ para } b < a \text{ y } b \approx a \\ I(a,b) &\rightarrow 0 \text{ para } a \rightarrow 0 \text{ y } b = 0 \end{aligned}$$

La siguiente familia de funciones de implicación cumple, para algunas funciones $n(a)$ particulares las propiedades anteriores, en mayor o menor medida:

$$Ro(a,b) = \begin{cases} 1 & \text{si } a \leq b \\ \frac{(2 \times b + a)^{n(a)}}{2} & \text{si } b \leq \frac{a}{2} \\ 1 - \frac{(2 \times (a - b) + a)^{n(a)}}{2} & \text{si } b \geq \frac{a}{2} \end{cases}$$

$$n(a) = \frac{h}{a^k}$$

con h y k reales positivos o nulos.

Podemos ver que para $h=1$ y $k=0$ obtenemos una de las implicaciones que aparece típicamente en la literatura, $Rd1$:

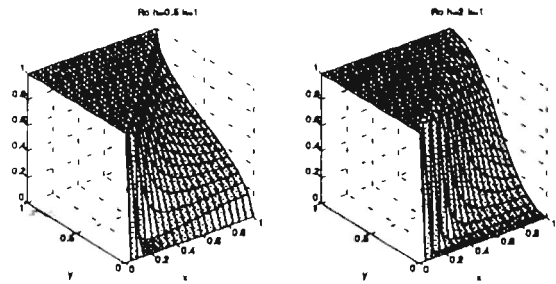
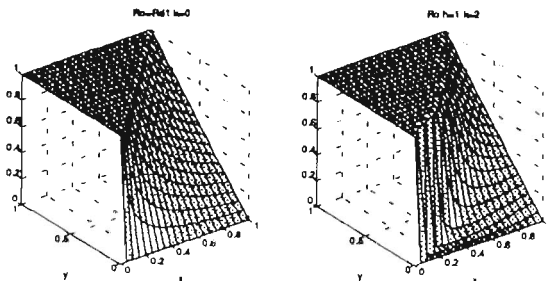
$$\text{si } n(a) = 1 \text{ (} h=1 \text{ y } k=0 \text{)} \Rightarrow Ro = Rd1 \begin{cases} 1 & \text{si } a \leq b \\ \frac{b}{a} & \text{si } a > b \end{cases}$$

El parámetro k está relacionado con la propiedad:

$$I(a,b) \approx 1 \text{ e } I(a,b) < 1 \text{ para } b < a \text{ y } b \approx a$$

Cuanto mayor es el valor de k , mejor se cumple esta propiedad, para valores tales que $b < a$, $a \neq 1$ y $a, b \neq 0$.

El parámetro h permite variar la función de implicación en la región dada por $a=1$. Para $h=1$, se tiene $Ro(1,b)=b$. Para valores de h mayores que 1 $Ro(1,b)$ adopta formas similares a las de una sigmoidea. Las siguientes figuras ilustran lo dicho:



El comportamiento de esta implicación es casi idéntico frente a los dos métodos de desborrosificación y notablemente bueno. Se puede además controlar el ancho de las distribuciones que genera, y escoger de forma acorde el número de muestras para la discretización. Resultados interesantes pueden obtenerse también modificando ligeramente implicaciones ya conocidas, en base a alguna de las propiedades enumeradas.

Referencias

- [1] M. Mizumoto y H.J. Zimmermann, Comparison of Fuzzy Reasoning Methods, Fuzzy Sets and Systems 8 (1982) 253-283, North-Holland.
- [2] D. Dubois y H. Prade, Fuzzy Logics and the Generalized Modus Ponens Revisited, Cybernetics and Systems, 15:293-331, 1984.
- [3] E. Trillas y L. Valverde, On Mode and Implication in Approximate Reasoning, M.M. Gupta et al. Eds, Approximate Reasoning in Expert Systems, Elsevier Science Publishers Amsterdam, 1985.
- [4] J. B. Kiska, M. E. Kochanska y D. S. Sliwinska, The Influence of Some Fuzzy Implication Operators on the Accuracy of a Fuzzy Model, Part I y Part II, Fuzzy Sets and Systems 15 (1985) 111-128, North-Holland.
- [5] T. Terano, K. Asai y M. Sugeno, Fuzzy Systems Theory and its Applications, Academic Press, INC., 1987.
- [6] S. Heilpern, The Expected Value of a fuzzy number, Fuzzy Sets and Systems 47 (1992) 81-86, North-Holland.
- [7] Z. Cao y A. Kandel, Applicability of Some Fuzzy Implication Operators, Fuzzy Sets and Systems 31 (1989) 151-186, North-Holland.
- [8] D. Park, Z. Cao y A. Kandel, Investigations on the applicability of fuzzy inference, Fuzzy Sets and Systems 49 (1992) 151-169, North-Holland.

Query-Answering in Fuzzy Temporal Constraints*

Lluís Vila and Lluís Godo

IIIA – CSIC

Camí de Sta. Bàrbara, s/n, 17300 Blanes, Catalonia, Spain

{vila,godo}@ceab.es

Abstract

Temporal Constraint Networks are a well-known powerful formalism for encoding temporal knowledge based on the CSP techniques. It has been recently proposed a redefinition of *Metric temporal constraints* to cope with vagueness of temporal relations based on fuzzy sets, called *Fuzzy Temporal Constraint Networks*.

In this paper we identify those queries on a fuzzy temporal constraint network relevant to a knowledge-based reasoning system, we provide definitions for their answering and explore the functions on which they are based by discussing the different choices for them. These results can be useful to those interested in defining a system for temporal reasoning under uncertainty. For instance, they are satisfactorily applied to the definition of a *possibilistic temporal reasoning system* [19].

Keywords. temporal knowledge-based systems, temporal constraints, fuzzy sets.

1 Introduction

Enriching knowledge-based systems with some sort of temporal representation has been an active area of research in the AI community. *Metric Temporal Constraint Networks* defined by Dechter, Meiri & Pearl [3] is a constraint-based approach which provides a formal basis for representing temporal relations. On the other hand, it is well-known that in many real AI problems belong to domains where information and knowledge is not precisely known but is *vague* and/or *uncertain*. It can be illustrated with an example coming from a real application medical domain. *Brucellosis* is one infectious pathology that may be the origin of serious *lumbalgia* problems. Consider the following description of a particular real case of *brucellosis*: A certain patient complains about an acute pain in the back which started three days ago. During the interview he explains that he suffered a period of *ondulating fever* during three or four weeks more or less. He remembers having discomfort for few days just before the fever.

About two weeks ago the fever ceased. He is positive about do not having fever at the time the pain in the back started. Answering a specific question from the physician, the patient recalls he ate goat cheese a little longer than one month ago. He is sure that it was not more than 45 days.

We are interested in describing uncertain and vague temporal information such as “a little longer than”, “three or four weeks more or less”, “few days” or “about three weeks ago”, within a common formalism where precise temporal knowledge can be naturally expressed. We also wish answering queries stated in similar terms.

Recently, several formalisms have been proposed for approximate temporal knowledge representation [9, 2, 7, 6, 11, 4, 13], some of which are based on possibility theory. A straight forward yet very promising idea is redefining metric temporal constraints above mentioned using fuzzy sets yielding the *Fuzzy Temporal Constraint Networks* [13]. We aim at using fuzzy temporal constraints as the formalism to encode temporal knowledge within a general-purpose reasoning system for problem-solving in dynamic domains [20]. Such a goal determines the set of relevant queries and operations to be performed on the fuzzy temporal constraint network. In this paper builds upon previous work on temporal constraint networks [13]. We make use of some previous definitions and explore, here, fuzzy set operations to define the relevant queries.

The structure of this paper is as follows. First we sketch the formalism of metric temporal constraints and their generalization using fuzzy sets. Then, we identify the relevant queries and operations over a fuzzy temporal constraint network, define their answering and discuss advantages and shortcomings of the different choices for the functions on which they are based.

2 Metric Temporal Constraints

Metric temporal constraints are a formalism based on metric constraints among pairs of time-points. A temporal knowledge base is represented by a set of variables and a set of constraints on these variables. Although there are certain cases where higher order constraints are required, the expressiveness of binary

*This work has been partially supported by Spanish MEC grant FPI-PN90 77909683 and by CICYT project ARREL (TIC92-0579-C02-01).

constraints is satisfactory for many interesting applications. We shall stay within the binary case. A temporal knowledge base is represented by a set of variables $\mathcal{X}$ and a set of constraints $\mathcal{T}$ on these variables. Each variable X_i represents a time point. Each binary constraint $T_{ij}^k = [a_{ij}^k, b_{ij}^k]$ —where a_{ij}^k and b_{ij}^k are values belonging to the variables domain—restricts the possible values for distance between a pair of variables by

$$a_{ij}^k \leq X_j - X_i \leq b_{ij}^k$$

A set of variables and a set of binary constraints on them form a *temporal constraint network*.

An interesting case is when the set of variables:

1. Include an *initial reference* time point t_0 —for simplicity $t_0 = 0$ is assumed.
2. Take values over a continuous domain extended with $+\infty$, $-\infty$, $+\epsilon$ and $-\epsilon$.

Such a case turns out to be specially suitable for temporal representation related to the occurrence of events and states [8, 18]. A nice feature of such a case, is that it enables to reduce every temporal knowledge—either concerning absolute or relative, metric or qualitative relations—into a conjunction of temporal constraints.

A tuple $S = (t_1, \dots, t_n)$ is called a *solution* of a constraint network if the assignment $\{X_1 := t_1, \dots, X_n := t_n\}$ satisfies all the constraints. A value v is a *feasible value* for the variable X_i if there exists a solution in which $X_i = v$. A constraint network is *consistent* if at least one solution exists. Regarding query answering, the concept of *minimal network* is central. Two networks are *equivalent* if they have the same variables set and they represent the same solution set. A constraint network is *minimal* if it is *tighter* than any other equivalent network. Intuitively the constraints of the minimal network summarize the information about the possible solutions of the network. Thus, from a computational point of view it supports efficient query-answering about consistency checking, feasible values and distances, consistent scenarios generation and constraint addition. There are graph processing techniques to compute the *minimal network*.

An interesting subcase is the one where every constraint is *simple*, in other words it can be expressed as a single interval. A set of simple temporal constraints is called a *Simple Temporal Problem* (STP). In such a case the queries above can be computed in polynomial time. For the general case, called the *Temporal Constraint Satisfaction Problem* (TCSP), where constraints are made of a canonical collection of intervals, checking consistency is a NP-hard problem and CSP techniques—like *path-consistency* and *network consistency*—can be applied as problem pre-processing. Afterwards, metric temporal constraints have been extended by incorporating qualitative and/or interval constraints [16, 12, 10, 14, 17]. In this paper we consider the simple case only.

3 Fuzzy Temporal Constraints

Temporal Constraint Networks build upon the notion of binary metric temporal constraint. In turn, the generalization to the fuzzy case is based on the definition of fuzzy temporal constraint. A *Fuzzy temporal constraint* is defined in [13, 1] as a constraint between a pair of time-points represented as a possibilistic function over distances instead of an interval of distance values.

Definition 1 (Fuzzy temporal constraint) Given two variables X_i, X_j and a metric value domain for them $\mathcal{D}$, a fuzzy temporal constraint C_{ij} between them is represented as a possibility distribution function $\pi_{ij} : \mathcal{DU} \rightarrow [0, 1]$, where $\mathcal{DU}$ is the union of $\mathcal{D}$ and its complementary.

The notion of fuzzy temporal constraint is a direct generalization of a metric temporal constraint. The notion of constraint network is directly generalized also.

Definition 2 A Fuzzy Temporal Constraint Network (FTCN) $\mathcal{N} = \langle \mathcal{X}, \mathcal{C} \rangle$ consists of a set of variables $\mathcal{X} = X_1, \dots, X_n$ and a set of fuzzy temporal constraints $\mathcal{C} = \{C_{ij} | i, j \leq n\}$ between them.

In figure 1 the different kinds of single temporal constraint are graphically presented.

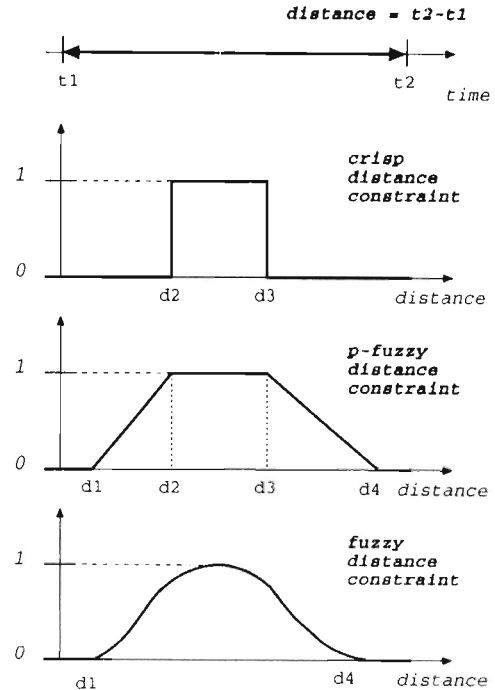


Figure 1: Temporal Constraints.

3.1 Using FTCN for Temporal Knowledge Representation

We are interested in a particular sort of FTCN which is specially adequate for representing temporal knowledge as well as allows efficient query answering.

Definition 3 A “normalized” fuzzy temporal constraint network (FTCN) is a FTCN which exhibits the following features:

1. “Initial time point”. $\mathcal{X}$ includes the special variable X_0 which value is fixed to 0.
2. “Universal duration”. $\forall X_i, X_j \in \mathcal{X}. C_{ij} = \pi_U \in \mathcal{C}$.
3. “Null duration”. $\forall X_i \in \mathcal{X}. C_{ii} = \pi_0$.
4. “Canonical constraints”. For each $X_i, X_j \in \mathcal{X}$ all constraints C_{ij} are reduced to a single constraint equivalent to them by applying constraint intersection.

To stay within the simple temporal case we must impose the condition for any fuzzy constraint π to be *unimodal*:

Definition 4 A normalized FTCN is said to be *simple* if any constraint $C_{ij} = \pi$ is unimodal, i.e.

$$\forall d_i \leq d_j \leq d_k \in \mathcal{DU}. \pi(d_j) \geq \min(\pi(d_i), \pi(d_k))$$

Simple normalized FTCN turn out to be an useful formalism to express vague temporal knowledge. Figure ?? illustrates it by formalizing the example presented in the introduction.

3.2 Solution of a FTCN

We define the notion of *solution* of a Temporal Constraint Network which is basic for defining the interesting queries.

Definition 5 (solution) Given a n -tuple $S = \{t_1, \dots, t_n\}$ where $t_i \in \mathcal{D}$ and a fuzzy temporal constraint network $\mathcal{N}$ then S is a possible solution of $\mathcal{N}$ at degree α iff $\mathcal{N}(S) = \alpha$, where

$$\mathcal{N}(S) = T(\dots, \pi_{ij}(t_j - t_i), \dots)$$

where T is a T -norm function and $T(\alpha_1, \dots, \alpha_n)$ denotes $T(\alpha_1, T(\alpha_2, \dots, T(\alpha_{n-1}, \alpha_n) \dots))$.

Akin to the notion of minimal network, we define the *minimal* fuzzy temporal constraint network. Two networks are equivalent if they represent the same solution set. A fuzzy temporal constraint $C = \pi$ is *tighter* than another one $C' = \pi'$ if $\forall d \in \mathcal{DU}. \pi(d) \leq \pi'(d)$. A FTCN $\mathcal{N}$ is tighter than another one $\mathcal{N}'$ with the same set of variables, if $\forall i, j \leq n. C'_{ij} \in \mathcal{N}'$ there exists a constraint $C_{ij} \in \mathcal{C}$ such that is tighter than C'_{ij} . The relation “tighter than” establishes a partial order among FTCNs. Given a certain FTCN there is one equivalent network which is the tightest, called the *minimal FTCN*. Henceforth, whenever we refer to

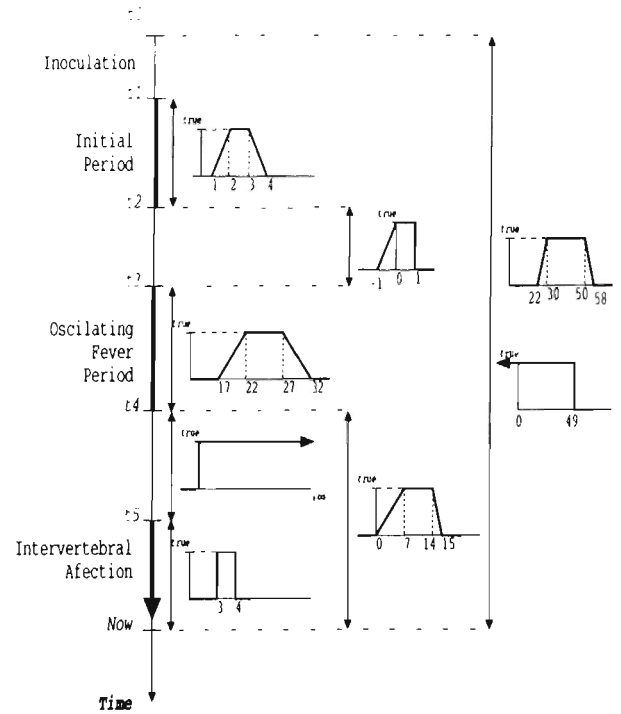


Figure 2: Brucellosis evolution case.

the possibility distributions for a FTCN network we shall mean those belonging to the equivalent minimal network.

The minimal network can be computed by using a classical *path consistency* algorithm. Let's previously define a special duration fuzzy set, the *Impossible duration* $\pi_0(x) = 0, \forall x$, and some basic binary operations on fuzzy constraints:

1. *Union*,

$$(\pi \cup \pi')(d_i) = \max\{\pi(d_i), \pi'(d_i)\}$$

2. *Conjunction*,

$$(\pi \cap \pi')(d_i) = \min\{\pi(d_i), \pi'(d_i)\}$$

3. *Composition*,

$$(\pi \circ \pi')(d_i) = \sup_{d_j + d_k = d_i} \{\min\{\pi(d_j), \pi'(d_k)\}\}$$

These operations on fuzzy sets can be directly used into the general scheme of the path consistency algorithm [?]:

```

function FPC-1 is
1   for  $k := 1$  to  $n$  do
2     for  $i, j := 1$  to  $n$  do
3        $T_{ij} \leftarrow \pi_{ij} \cap \pi_{ik} \circ \pi_{kj}$ 
4       if  $T_{ij} = \pi_{\emptyset}$  then Exit
          (the FTCN is inconsistent)
5     endfor
6   endfor
endfunction

```

Figure 3 presents the minimal constraint representation of our example.

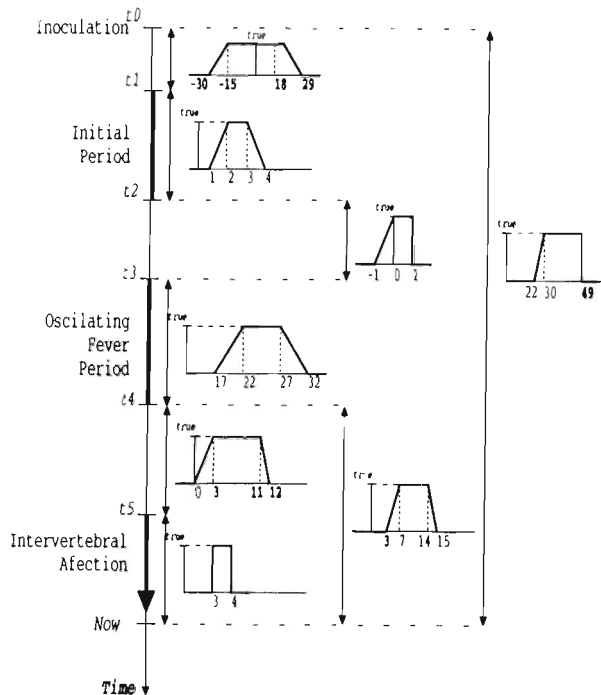


Figure 3: Minimal representation of the brucellosis evolution case.

4 Queries on Fuzzy Temporal Constraint Networks

Once we enjoy the definition of fuzzy temporal constraint network we are interested in identifying and analyzing the relevant queries regarding it. On the one hand, we consider the relevant queries for a fuzzy temporal constraint network independently of what it is going to be used for. The most basic one is checking whether a set of constraints is consistent or not. We propose the following possibilistic generalization:

Definition 6 ("Network consistency degree" – Cn-) *The degree of consistency of a fuzzy temporal constraint network $\mathcal{N}$ composed of n variables is*

$$Cn(\mathcal{N}) = \max_{S \in \mathcal{D}^n} N(S)$$

Given a network with a certain degree of consistency one may wish to know whether a certain value is a feasible one or to answer queries about the set of all - up to a certain degree - possible solutions.

Definition 7 (feasible value) *A value t is a feasible value at degree α for a variable X_i , if there exists a solution S in which $X_i = t$ and $N(S) = \alpha$.*

On the other hand, we consider queries relevant to knowledge-based temporal reasoning. We aim at using fuzzy temporal constraints as the formalism to encode temporal knowledge within a general-purpose reasoning system for problem-solving in dynamic domains. In this paper we consider inference in a propositional language. Results here will be the basis for a straight forward extension to the first-order case.

Our goal determines the set of relevant queries and operations which will be performed on the fuzzy temporal constraint network. From a knowledge-based system viewpoint, the fuzzy temporal constraint network is viewed as an efficient encoding of the *temporal facts* base, i.e. the base composed of those facts formed by a temporal relation -the temporal constraints. An inference system *retrieves* the fact base to find out temporal constraints which allow the application of a certain inference rule by fulfilling the temporal constraints as well as *asserts* those facts inferred through the rule application. In particular, a *Modus Ponens*-like inference rule, retrieves an implication and a set of facts which satisfy its left-hand side to assert the right-hand side.

Assertion. It involves checking whether the constraint(s) at hand is (are) *compatible* with the constraint network.

Retrieval. It usually involves checking whether a set of temporal constraints being part of a certain form are *entailed* by those that can be retrieved from the constraint network. In the case of *Modus Ponens*, for instance, the retrieval of the temporal constraint network is intended to fulfill the antecedent of an implicative form.

In the fuzzy context, boolean answers are generalized to a possibility degree. *Compatibility* (C) and *entailment* (E) are captured in *possibility theory* by the notions of *possibility* and *necessity* measures respectively [5].

Proposition 1 (constraint compatibility C) *Given a fuzzy temporal constraint $C_{ij}^* = \pi_{ij}^*$ and a fuzzy temporal constraint network $\mathcal{N}$ - where π_{ij} represents the possibility distribution associated to the constraint C_{ij} of the equivalent minimal network -, the degree of compatibility of C_{ij}^* with $\mathcal{N}$ is*

$$C(C_{ij}^*, \mathcal{N}) = \max_{d \in \mathcal{D}^n} T(\pi_{ij}^*(d), \pi_{ij}(d))$$

where T is a T -norm function.

There is a whole range of particular functions one can take as T , being the *minimum* the most “optimistic” choice and *Lukaziewicz* the most “pessimistic” one.

Example 1 Consider the constraint on the distance from the Inoculation event to the begin of the Initial Period. Let's assume that a new finding about such a distance is known with a certain fuzziness. An interesting matter to know is the degree at which the finding is “compatible” with the constraint. An instance of such situation is presented in figure 4. The T -norm function used there is the minimum. The degree obtained from the figures in the example is $C = 0.875$.

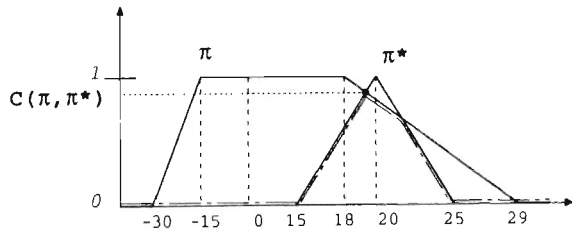


Figure 4: Using the Compatibility measure.

The definition of entailment is a bit more controversial. In general, it is defined by means of an implication function.

Proposition 2 (constraint entailment E) Given a fuzzy temporal constraint $C_{ij}^* = \pi_{ij}^*$ and a fuzzy temporal constraint network $\mathcal{N}$ —where π_{ij} represents the possibility distribution associated to the constraint C_{ij} of the equivalent minimal network—, the degree of entailment of C_{ij}^* by $\mathcal{N}$ is

$$E(C_{ij}^*|\mathcal{N}) = \inf_{d \in \mathcal{DU}} I(\pi_{ij}(d), \pi_{ij}^*(d))$$

where I is an implication function.

In order to define the appropriate implication function, we have several choices. One way is simply taking the dual definition of the compatibility degree:

$$E(C^*|\mathcal{N}) = 1 - C(\overline{C^*}, \mathcal{N}) = \inf_{d \in \mathcal{DU}} S(1 - \pi_{ij}(d), \pi_{ij}^*(d))$$

, where S is a T -conorm function dually defined from T as $\forall x, y. S(x, y) = 1 - T(1 - x, 1 - y)$. In such a case the implication function is stated as $I(x, y) = S(1 - x, y)$ which is known as *strong implication*. Though it is an acceptable *implication function*, there are some “intuitive” properties that are not guaranteed. In particular, one could expect $E(C_{ij}^*|\mathcal{N}) = 1$ whenever π_{ij}^* is equal to π_{ij} . To account for this property, alternative implication functions have been proposed [15]:

- *Residuated implications*: $I(x, y) = \sup\{c \mid T(x, c) \leq y\}$ where T is a T -norm function.

- *Reciprocal residuated implications*: $I(x, y) = I'(1 - y, 1 - x)$ where I' is a residuated implication.

The entailment measures obtained from these families are not necessarily dual with respect to the compatibility measure. A particular case of them has been identified for which duality is guaranteed. Namely, taking $T(x, y) = \max(0, x + y - 1)$ as T -norm leads to the well-known *Lukaziewicz implication*

$$I(x, y) = \min(1, 1 - x + y)$$

Example 2 Consider the constraint on the distance from the Inoculation event to now. Let's assume that there is knowledge about the usual pattern of evolution of the Brucellosis pathology. In particular we may assume that the mentioned distance is part of the premise of a condition within a rule intended to hypothesize about the holding of brucellosis. In such a case we would be interested in assessing the degree at which such a condition is entailed by the constraint obtained from the patient history. An instance of it is presented in figure 5. The implication function used there is based on the above mentioned Lukaziewicz implication. The degree obtained from the figures in the example is $E = 0.333$.

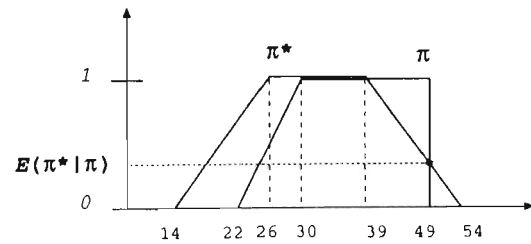


Figure 5: Using the Entailment measure.

Compatibility and *entailment* measures are employed to define additional interesting queries within the temporal context. This is a matter of our current research.

5 Conclusions

Fuzzy temporal constraint networks suitably support the representation of vague or uncertain temporal knowledge. We have used the straight forward generalizations of metric temporal constraints to the fuzzy to analyze the query answering on them. The notion of *network consistency degree* is introduced and constraint compatibility and entailment are captured by applying the notions of possibility and necessity measures usual in possibility theory. We identified the well-known *Lukaziewicz implication* as the only one which allows a definition of entailment dual with respect the compatibility measure whilst guarantees the expected intuitive property of $E(C_{ij}^*|\mathcal{N}) = 1$ whenever π_{ij}^* is equal to π_{ij} .

The measures of constraint compatibility and entailment are fundamental notions to define a general purpose system for reasoning about time under uncertainty. We are currently analyzing optimizations for the adapted path consistency used for obtaining the minimal network as well as the definition of additional fuzzy propagation algorithms derived from it to incorporate a further constraint on a minimal network and re-compute it.

Bibliography

- [1] S. BARRO, R. MARIN, J. MIRA, and A. PATON. A model and a language for the fuzzy representation and handling of time. *Fuzzy Sets and Systems*, (to appear), 1993.
- [2] T. DEAN and K. KANAZAWA. Probabilistic temporal reasoning. In *AAAI'88*, pages 524-528, 1988.
- [3] R. DECHTER, I. MEIRI, and J. PEARL. Temporal constraint networks. *Artificial Intelligence*, 49:61-95, 1991.
- [4] D. DUBOIS, J. LANG, and H. PRADE. Timed possibilistic logic. In Z. Ras, editor, *Fundamenta Informaticae Special Issue on Artificial Intelligence*. 1992.
- [5] D. DUBOIS and H. PRADE. *Possibility Theory*. Plenum Press, New York, 1988.
- [6] D. DUBOIS and H. PRADE. Processing fuzzy temporal knowledge. *IEEE Transactions on Systems, Man and Cybernetics*, 19(4), July/August 1989.
- [7] S. DUTTA. An event-based fuzzy temporal logic. In *18th IEEE Int. Symp. on Multi-valued Logics*, pages 64-71, 1988.
- [8] G. ESCALADA-IMAZ, C. SIERRA, and L. VILA. Defining a reasoning system based on temporal tokens. In C. Sierra and L. Vila, editors, *1st Spanish Workshop on Temporal Representation and Reasoning TARRAT'93*. IIIA, 1993.
- [9] T. C. FALL. Evidential reasoning with temporal aspects. In *AAAI'86*, pages 891-895, 1986.
- [10] H. A. KAUTZ and P. B. LADKIN. Integrating metric and qualitative temporal reasoning. In *AAAI'91*, pages 241-246, 1991.
- [11] J. KOHLAS and P.-A. MONNEY. Temporal reasoning under uncertainty with belief functions. In *IPMU'90*, pages 67-73, 1990.
- [12] P. LADKIN. Constraint reasoning with intervals: A tutorial, survey and bibliography. Technical Report TR-90-059, ICSI, Berkeley, 1990.
- [13] R. MARIN, S. BARRO, A. BOSCH, and J. MIRA. Modeling the representation of time from a fuzzy perspective. *Cybernetics and Systems: an International Journal*, (to appear), 1993.
- [14] I. MEIRI. Combining qualitative and quantitative constraints in temporal reasoning. In *AAAI'91*, 1991.
- [15] E. TRILLAS and L. VALVERDE. On mode and implication in approximate reasoning. In M. Gupta, A. Kandel, W. Bandler, and J. Kiszka, editors, *Approximate Reasoning in Expert Systems*, pages 157-166. North-Holland, Amsterdam, 1985.
- [16] P. van BEEK. Approximation algorithms for temporal reasoning. In *IJCAI'89*, pages 1291-1296, 1989.
- [17] L. VILA. Constraints on distances between temporal distances. Report de Recerca forthcoming, IIIA, 1993.
- [18] L. VILA. Reformulación del cálculo de eventos en términos de restricciones temporales métricas. *Novática*, 108:23-29, Marzo-Abril 1994.
- [19] L. VILA, L. GODO, and G. ESCALADA-IMAZ. A temporal reasoning system based on fuzzy temporal metric constraints. In *FUZZY'94: IV Congreso Español de Lógica y Tecnologías Fuzzy*. IIIA - CSIC, 1994.
- [20] L. VILA, L. GODO, and G. ESCALADA-IMAZ. A temporal reasoning system based on fuzzy temporal metric constraints. Report de Recerca forthcoming, IIIA - CSIC, 1994.

Fuzzy Logics as Families of Bivaluated Logics

J.L. Castro

Depto de Ciencias de la Computación e Inteligencia Artificial.

ETSI-Informática. Universidad de Granada.

18071-Granada. Spain.

Telf: 58.244017, Fax: 58.243317, E-mail: jcastro@ugr.es.

Abstract

In this paper we will prove the equivalence between L -valued logics and L -indexed families of bivaluated logics. It is also proved that this equivalence can be used as a tool for checking compactness and decidability properties of a L -valued logic.

Keywords: Fuzzy Logic, Consequence operator, decidability and compactness.

1 Introduction

In recent years a conspicuous interest in fuzzy logic has brought and many interesting applications has been developed within Artificial Intelligence. The present paper tries to study those logics from an abstract point of view, in order to obtain general properties that can be useful for applications.

In section 2, after introducing the concept of L -closure system, we establish the equivalence between L -consequence operators Pavelka's L -closure systems. The associated classical abstract logics are presented in section 3. Finally, in section 4, that result is used in order to study multivalued logic from the properties of their bivaluated logic associated.

2 Closure Systems and Consequence Operations

Pavelka [1979] introduced the notion of L -consequence operator in order to study the L -valued logic from an general point of view. In 1991, Murali [1991] introduced the concept of L -closure system in fuzzy set theory and proved that both concepts are in a bijective correspondence. We will prove that the correspondence is a lattice anti-isomorphism.

Let L a complete lattice with the order $\leq$ and the operations infimum and supremum $\wedge$ and $\vee$ respectively (inf and sup for infinite case), minimum 0 and maximum 1.

Definition 1 [7] A family $\tilde{C}$ of L -subsets of a non-empty set S is said to be a L -closure system on S if $\tilde{C}$ is closed under arbitrary intersections, i.e. $\bigcap \tilde{F}$ is in $\tilde{C}$ if $\tilde{F} \subseteq \tilde{C}$, and $S \in \tilde{C}$, i.e. the L -subset $S(a) = 1$ for each $a \in S$ is in $\tilde{C}$.

If s is an L -subset of S , s_α will denote the α -cut of s , $s_\alpha = \{a \in S : s(a) \geq \alpha\}$.

Theorem 1. If $\tilde{C}$ is an L -closure system on S , for each $\alpha \in L$ the family $\tilde{C}_\alpha$ defined by $\tilde{C}_\alpha = \{s_\alpha : s \in \tilde{C}\}$ is a classical closure system.

Theorem 2. The family of all L -closure systems on S is a complete lattice under

$$\inf \tilde{C}_i = \bigcap_i \tilde{C}_i = \{s \in L^S \mid s \in \tilde{C}_i \text{ for all } i\}$$

$$\sup \tilde{C}_i = \inf \{\tilde{C} : \tilde{C}_i \subseteq \tilde{C}, \text{ for all } i\}.$$

L -set inclusion, minimum $\{S\}$ and maximum L^S .

Remark 1. This lattice is an extension of the complete lattice of all closure system on S , since a closure system is also an L -closure system, being infimum and supremum operations the same in both lattice.

Pavelka [10] introduced the notion of an L -consequence operator in order to give an abstract notion of L -multivalued logic.

Definition 2. A L -consequence operator on S is defined as a mapping from the power set L^S into itself, $\tilde{C} : L^S \rightarrow L^S$, such that

C1') fuzzy inclusion: $A \subseteq \tilde{C}(A)$,

C2') fuzzy monotony: If $A \subseteq B \subseteq S$, then $\tilde{C}(A) \subseteq \tilde{C}(B)$.

C3') fuzzy idempotence: $\tilde{C}(\tilde{C}(A)) = \tilde{C}(A)$,

hold for each $A, B \in L^S$.

Usually, a compactness property is added. An L -consequence operator is said to be compact if, and only if, it satisfies the condition

(C4') fuzzy compactness: for any $A \in L^S$ and any $x \in S$ there exist a finite subset $G \subset S$ such that $\tilde{C}(A)(x) = \tilde{C}(A \upharpoonright G)(x)$,

where $X \upharpoonright G$ represents the L -subset

$$(X \upharpoonright G)(y) = \begin{cases} X(y) & \text{if } y \in G \\ 0 & \text{otherwise} \end{cases}$$

Proposition 1 Let $\tilde{C}$ be an L -consequence operation on S . Then,

i) The restriction of $\tilde{C}$ to 2^S , C_1 , defined by $a \in C_1(A)$ if and only if $\tilde{C}(A)(a) = 1, \forall a \in S, \forall A \subseteq S$, is a classical consequence operator.

ii) If $\tilde{C}$ satisfies the compactness axiom, C_1 also satisfies the classical compactness axiom.

Remark 2. Given $\alpha \neq 1$, the restriction of $\tilde{C}$ to 2^S , C_α , defined by $a \in C_\alpha(A)$ if and only if $\tilde{C}(A)(a) \geq \alpha, \forall a \in S, \forall A \subseteq S$, is not in general a classical consequence operator, since idempotence is not always followed.

Lemma 1. If C_α is a classical consequence operator for every $\alpha \in L$, then

$$C_\beta(C_\alpha(A)) = \begin{cases} C_\beta(A) & \text{if } \beta \leq \alpha \\ C_\alpha(A) & \text{otherwise} \end{cases}$$

Proof: If $\beta \leq \alpha$ then $C_\alpha(A) \subseteq C_\beta(A)$ and $C_\beta(A) \subseteq C_\beta(C_\alpha(A)) \subseteq C_\beta(C_\beta(A)) = C_\beta(A)$. On the other hand, if $\alpha \leq \beta$ then $[C(C_\alpha(A))]_\beta \subseteq [\tilde{C}(C_\alpha(A))]_\alpha = C_\alpha(C_\alpha(A)) = C_\alpha(A)$, hence $C_\alpha(A) \subseteq C_\beta(C_\alpha(A)) \subseteq C_\alpha(A)$.

Proposition 2. C_α is a classical consequence operator for every $\alpha \in L$ if, and only if,

$$\tilde{C}(C_\alpha(A))(a) = \begin{cases} 1 & \text{if } a \in C_\alpha(A) \\ \tilde{C}(A)(a) & \text{otherwise} \end{cases}$$

Corollary 1. The condition "If A is a classical set then $\tilde{C}(A)$ is a classical set" holds if and only if $C_\alpha = C_1$ for every $\alpha \leq 1$.

Theorem 3. The family of all L -consequence operator on S is a complete lattice under the ordering

$$\tilde{C}_1 \leq \tilde{C}_2 \quad \text{iff} \quad \tilde{C}_1(A) \leq \tilde{C}_2(A), \text{ for each } A \in L^S.$$

Here,

$$\inf \tilde{C}i(A) = \bigcap_i \tilde{C}i(A)$$

$$\sup \tilde{C}i = \inf \{ \tilde{C} : \tilde{C}_i \subseteq \tilde{C}, \text{ for all } i \}.$$

Theorem 4. i) Every L -consequence operator $\tilde{C}$ on S defines the L -closure system

$$\tilde{C}(\tilde{C}) = \{ A \in L^S / \tilde{C}(A) = A \},$$

i.e. the family of all $\tilde{C}$ -closed L -subsets of S . Clearly, it is verified:

- a) $\bigcap \tilde{C}(\tilde{C}) = \tilde{C}(\leq)$, where $\leq(a) = 0$ for each $a \in S$,
- b) $S \in \tilde{C}(\tilde{C})$.

ii) Every L -closure system C on S defines the L -consequence operator

$$\tilde{C}(\tilde{C})(A) = \bigcap \{ T \in \tilde{C} / A \subseteq T \}.$$

Theorem 5. The complete lattice of all L -consequence operations on S and the complete lattice of all L -closure systems on S are dually isomorphic under the correspondences $\tilde{C} \rightarrow \tilde{C}(\tilde{C})$, and $\tilde{C} \rightarrow \tilde{C}(\tilde{C})$.

Remark 3. Analogously to the bivaluated case, the elements of $\tilde{C}(\tilde{C})$ are called the theories of the L -valued logic $\tilde{C}$. Thus, $\tilde{C}(\tilde{C})$ is the L -valued logic whose theories are $\tilde{C}$, and the previous theorem shows the equivalence between the theories of an L -valued logic, and the L -consequence operator defining this logic.

3 L-valued Logics as Indexed Families of Bivaluated Logics

Given an L -consequence operator $\tilde{C}$ on S , we have the associated L -closure system $\tilde{C}(\tilde{C})$ on S , which has associated the family $\mathcal{C}(\tilde{C})_\alpha \forall \alpha \in L$ of classical closure systems on S ; and so, the family $C^\alpha = C(\mathcal{C}(\tilde{C})_\alpha) \forall \alpha \in L$ of classical consequence operators on S associated to everyone of these classical closure system. Thus, we can associated a family of classical consequence operator C^α for every L -fuzzy consequence operator:

Definition 3. For each $\alpha \in L$, every L -consequence operator $\tilde{C}$ on S defines a classical consequence operator C^α , obtained from $\mathcal{C}(\tilde{C})_\alpha$ through the equivalence between classical consequence operator and classical closure systems. Thus,

$$C^\alpha(A) = \bigcap \{ T_\alpha / \tilde{C}(T) = T \text{ and } A \subseteq T_\alpha \}.$$

A more operational expression to $C^\alpha(A)$ is obtained in the following three propositions:

Proposition 3. $C^\alpha(A) = \tilde{C}(S_A^\alpha)_\alpha$, where

$$S_A^\alpha(x) = \begin{cases} \alpha & \text{if } x \in A \\ 0 & \text{otherwise} \end{cases}$$

Proposition 4. If $\tilde{C}$ is defined by means of an L -semantic on S , that is, if there exists a family $S \subseteq L^S$ such that

$$\tilde{C}(A) = \bigcap \{ T \in S / A \subseteq T \},$$

then

$$C^\alpha(A) = \{ x \in S / T(x) \geq \alpha \text{ if } T(a) \geq \alpha \forall a \in A \}.$$

Proposition 5. If $\tilde{C}$ is defined by means of an L -syntax $\langle \mathcal{A}, \mathcal{R} \rangle$ (see [8]), i.e. $\tilde{C} = C_{\mathcal{R}, \mathcal{A}}$, ($\mathcal{R}$ the fuzzy set of rules and $\mathcal{A}$ the fuzzy set of axioms) then

$x \in C^\alpha(A)$ if and only if for every $\beta < \alpha$, there exist a $w \in \mathbf{P}(S, \mathcal{R})$ such that $w^\neg = x$, and $v(w, \mathcal{A}, S_A^\alpha) > \beta$.

Roughly speaking, $x \in C^\alpha(A)$ if, and only if, for every $\beta < \alpha$ there exist a $\mathcal{R}$ -proof of x with validity greater than β under the hypothesis A at level α and the axioms $\mathcal{A}$.

From these general results, we have a pseudo-syntax and a pseudo-semantic for the bivaluated logics associated to the different L -valued logic presents in the literature:

Example 1. Let us consider the Lukasiewicz's logic on $[0, 1]$, i.e., with valuation functions:

$$v(\neg p) = 1 - v(p); \quad v(p \vee q) = \text{Max}(v(p), v(q))$$

$$v(p \wedge q) = \text{Min}(v(p), v(q))$$

$$v(p \leftrightarrow q) = \text{Min}(1, 1 - v(p) + v(q)).$$

The L -consequence operator obtained from that semantic is $C(A)(x) = \inf\{v(x) / v \text{ is a valuation with } v(a) \geq A(a) \forall a \in P\}$, P being the set of propositions of the logic.

Example 2. Let us consider the propositional logic on $[0, 1]$ with the following axioms [6]:

$$A1) p \leftrightarrow p \vee q/1 \quad A4) p \wedge q \leftrightarrow q \wedge p/1$$

$$A2) p \vee q \leftrightarrow q \vee p/1 \quad A5) q \vee p \leftrightarrow q/1/2$$

$$A3) p \wedge q \leftrightarrow p/1$$

and the only rule of inference

$$\frac{P/a, P \leftrightarrow Q/b}{Q, a * b},$$

being $*$ a t -norm.

In this case, when $0 \leq \alpha < 1/2$, $q \in C^\alpha(q \vee p)$ but, when $1/2 < \alpha \leq 1$, $q \notin C^\alpha(q \vee p)$.

This one is an example of a non-trivial fuzzy logic in which the hypothesis "A classical set implies $\tilde{C}(A)$ classical set", falls, since from $\{p/1, q/0\}$ follows q with validity $1/2$ by using $A5$.

A question is now considered: Are the L -consequence operators characterized by its associated family of classical consequence operators, i.e.,

it is $\tilde{C}_1 = \tilde{C}_2$ if and only if $C_1^\alpha = C_2^\alpha$ for every $\alpha \in L$?

The answer is negative in general, since $\tilde{C}_1^\alpha = \tilde{C}_2^\alpha$ for every $\alpha \in L$ if and only if $\mathcal{C}(\tilde{C}_1)_\alpha = \mathcal{C}(\tilde{C}_2)_\alpha$ for every $\alpha \in L$, but it is possible although $\tilde{C}(\tilde{C}_1) \neq \tilde{C}(\tilde{C}_2)$.

Nevertheless, if the index theories maps

$$\tilde{C}(\tilde{C}) \xrightarrow{\text{index}_\alpha} \mathcal{C}(\tilde{C})_\alpha$$

$$T \rightarrow T_\alpha$$

are added, then the families $\{(C^\alpha, \text{index}_\alpha) \mid \alpha \in L\}$ characterize and determine the L -consequence operator $\tilde{C}$.

Equivalence Theorem. Let $\tilde{C}$ be an L -consequence operator. Then the family $\{(C^\alpha, \text{index}_\alpha) \mid \alpha \in L - \{0\}\}$ verifies.

(I) For every $T \in \tilde{C}(\tilde{C})$, if $\alpha \leq \beta$ then $T_\alpha \subseteq T_\beta$,

(II) If $\tilde{\mathcal{F}} \subseteq \tilde{C}(\tilde{C})$ and $T = \bigcap \tilde{\mathcal{F}}$ then $T_\alpha = \bigcap \{T'_\alpha / T' \in \tilde{\mathcal{F}}\}$

(III) For every $T \in \tilde{C}(\tilde{C})$ and every $x \in S$, there exist $\max\{\alpha \in L / x \in T_\alpha\}$.

Conversely, if we have a family $\{C(\alpha), \text{Index}(\alpha) / \alpha \in L - \{0\}\}$ where $C(\alpha)$ is a consequence operator for all $\alpha \in L$ and $\text{index}(\alpha)$ is an indexation over the closure systems associated $i \rightarrow T_i^\alpha$ such that it is verified

(I) If $\alpha \leq \beta$ then $T_i^\beta \subseteq T_i^\alpha$, and

(II) If $J \subseteq I$ and $\bigcap_{j \in J} T_j^\alpha = T_{j_0}^\alpha$, then $\bigcap_{j \in J} T_j^\beta = T_{j_0}^\beta$,

(III) for every $i \in I$ and every $x \in S$, there exists $\max\{\alpha \in L / x \in T_i^\alpha\}$,

then we can define an L -consequence operator $\tilde{C}$ and a mapping $(\tilde{C})$ such that $C^\alpha = C(\alpha)$ and $T_i^\alpha = T_{i_\alpha}$. $\tilde{C}$ is defined as $\tilde{C}(\tilde{C})$ where

$$\tilde{C} = \{T_i / i \in I\},$$

$$T_i(x) = \max\{\alpha \in L / x \in T_i^\alpha\}.$$

Proof: The properties of $\{(C^\alpha, \text{index}_\alpha) \mid \alpha \in L - \{0\}\}$ are evident. Conversely, if $\{C(\alpha), \text{Index}(\alpha) / \alpha \in L - \{0\}\}$ satisfies the condition (I), (II) and (III), then set $\tilde{C} = \tilde{C}(\tilde{C})$ where

$$\tilde{C} = \{T_i / i \in I\},$$

$$T_i(x) = \max\{\alpha \in L / x \in T_i^\alpha\}.$$

$\tilde{C}$ is a L -closure system since if $\tilde{\mathcal{F}} \subseteq \tilde{C}$, then there exists a $J \subseteq I$ such that

$$\tilde{\mathcal{F}} = \{T_j / j \in J\},$$

and for every $\alpha \in L$, $(\bigcap \tilde{\mathcal{F}})_\alpha = \bigcap_{j \in J} T_{j_0}^\alpha = \bigcap_{j \in J} T_j^\alpha = T_{j_0}^\alpha$ by the definition of T_j and (II). Hence $(\bigcap \tilde{\mathcal{F}}) = T_{j_0} \in \tilde{C}$. Finally, it is obvious that $T_{i_\alpha} = T_i^\alpha$ and thus $C(\alpha) = C^\alpha$.

Roughly speaking, an L -consequence operator is equivalent to an L -indexed family of classical consequence operator together a common indexation over its theories such that the indexation assign the same index to intersections for every $\alpha \in L$. Observe that if L is finite, then condition (III) always hold

4 Checking Properties of L-valued Logics

The equivalence theorem is for checking properties of L -valued logics by checking the properties of its associated bivaluated logics.

The first property we study is the compactness:

Theorem 6. *Let $\tilde{C}$ be an L -consequence operator. If $\tilde{C}$ is compact, then C^α is compact for every $\alpha \in L$.*

Proof: Let us suppose that $\tilde{C}$ is compact and show that C^α is compact. Let $A \subseteq S$, then

$$1) C^\alpha(A) = \tilde{C}(S_A^\alpha)_\alpha.$$

From the compactness of $\tilde{C}$ follows that for each $x \in S$ there exist a finite subset $G \subseteq S$ such that $\tilde{C}(A)(x) = \tilde{C}(A \mid G)(x)$. Let $A \subseteq S$, and let G a finite subset of S such that

$$2) \tilde{C}(S_A^\alpha \mid G)(x) = \tilde{C}(S_A^\alpha)(x).$$

Let us check that

$$3) S_A^\alpha \mid G = S_{A \cap G}^\alpha.$$

If $x \in A \cap G$ then $(S_A^\alpha \mid G)(a) = \alpha = S_{A \cap G}^\alpha(a)$, and if $x \notin A \cap G$ then $(S_A^\alpha \mid G)(a) = 0 = S_{A \cap G}^\alpha(a)$.

From 1, 2 and 3 follows $C^\alpha(A) = \tilde{C}(S_A^\alpha)_\alpha$ and $x \in C^\alpha(A)$ if and only if

$$\tilde{C}(S_A^\alpha \mid G)(x) \geq \alpha \text{ iff } \tilde{C}(S_{A \cap G}^\alpha)(x) \geq \alpha \text{ iff } x \in C^\alpha(A \cap G),$$

hence C^α is compact.

Other interesting property is the decidability. The notion of decidability of an L -subset was introduced by Biacino and Gerla [1].

Definition 4. *We will say that an L -consequence operator $\tilde{C}$ is decidable if every theory of the logic is decidable, that is, $\tilde{C}(A)$ is decidable for every $A \in L^S$.*

Theorem 7. *Let $\tilde{C}$ be an L -consequence operator. If $\tilde{C}$ is decidable, then C^α is decidable for every $\alpha \in L$, that is, $C^\alpha(A)$ is decidable for any $A \subseteq S$ and any $\alpha \in L$.*

Proof: If $\tilde{C}$ is decidable then, for every $A \in L^S$, the mappings

$\tilde{C}(A) : N \rightarrow L$ and $\neg\tilde{C}(A) : N \rightarrow L$ are recursive enumerable,

hence there exist two recursive maps $h_1(A), h_2(A) : NxN \rightarrow L'$ such that

$$\tilde{C}(A)(x) = \sup\{h_1(A)(x, n) \mid n \in N\}, \text{ and}$$

$$\neg\tilde{C}(A)(x) = \sup\{h_2(A)(x, n) \mid n \in N\}$$

hence, for every $A \subseteq S$,

$C^\alpha(A) = \{x \mid \tilde{C}(S_A^\alpha)(x) \geq \alpha\} = \{x \mid \exists n \in Nx h_1(S_A^\alpha)(x, n) \geq \alpha\}$ is recursively enumerable and $C^\alpha(A)^c = \{x \mid x \notin C^\alpha(A)\} = \{x \mid \tilde{C}(S_A^\alpha)(x) < \alpha\} = \{x \mid \neg\tilde{C}(S_A^\alpha)(x) \geq \neg\alpha\} = \{x \mid \exists n \in Nx h_2(S_A^\alpha)(x, n) \geq \neg\alpha\}$ is recursively enumerable, therefore $C^\alpha(A)$ is decidable.

Remark 4. *This two results can be used in order to prove that a multivalued logic $\tilde{C}$ has not such a*

property. For example, in order to prove that $\tilde{C}$ is not compact is enough to prove that C^α is not compact for some $\alpha \in L$.

References

- [1] L. BIACINO and G. GERLA, Fuzzy subsets: a constructive approach, *Fuzzy Sets and Systems* 45, 161-168 (1992).
- [2] D.J. BROWN and R.SUSZKO, Abstracts Logics, *Dissertationes Mathematicae* 102, pp. 9-42.(1973)
- [3] J.L.CASTRO and E.TRILLAS, Sobre preórdenes y operadores de consecuencias de Tarski, *Theoria* 11, 419-425. (1989).
- [4] J.L. CASTRO and E. TRILLAS, On the semantic of Implication, *Proceedings of First International Conference on Fuzzy Logic & Neural Networks*, Iizuka (Japan), 719-723.(1990)
- [5] J.L. CASTRO, *Contribución al estudio de modelos lógicos para la Inteligencia Artificial*, Ph.D. Thesis. Universidad de Granada, (1991).
- [6] J.A. GOGUEN, The logic of inexact concepts. *Synthese* 19, 325-373, (1969).
- [7] V. MURALI, Lattice of fuzzy subalgebras and closure systems in I^X . *Fuzzy set and Systems*, 41,101-111 (1991).
- [8] J. PAVELKA, On Fuzzy Logic I., *Zeitschr. f. Math. Logik und Grundlagen d. Math.* Bd.25, 45-52 (1979).
- [9] PAVELKA, J., On Fuzzy Logic II., *Zeitschr. f. Math. Logik und Grundlagen d. Math.* Bd.25, 119-134 (1979).
- [10] PAVELKA, J., On Fuzzy Logic III., *Zeitschr. f. Math. Logik und Grundlagen d. Math.* Bd.25, 447-464 (1979).

SOBRE LA IDENTIDAD DE CONDICIONALES MATERIALES

E. Trillas

Dept. Inteligencia Artificial U.P.M.

Facultad de Informática. 28660 Boadilla del Monte. Madrid.

S. Cubillo

Dept. Matemática Aplicada U.P.M.

Facultad de Informática. Campus de Montegancedo.

28660 Boadilla del Monte. Madrid.

e-mail: scubillo@fi.upm.es

A. Rodríguez

Dept. Dirección y Economía de la Empresa.

Universidad de León.

Resumen

Dados $\mu, \eta : X \rightarrow [0, 1]$, se estudia cuándo se da la igualdad $I_\mu^T = I_\eta^T$, siendo T una t -norma continua, y I_θ^T el preorden elemental:

$$I_\theta^T(y/x) = \sup\{z ; T(\theta(x), z) \leq \theta(y)\}.$$

Palabras clave: Igualdad de T -preórdenes, Condicionales Materiales

1 Introducción

Dado un subconjunto $A \subset X$, $A \neq \emptyset$, llamamos condicional material a la relación en $X \times X$ dada por $\rightarrow_A = (A \times A) \cup (A' \times X)$, que claramente es un preorden. La caracterización de un subconjunto por su condicional material, nos viene dada por el siguiente resultado ([2]).

Teorema 1 Si A y B son subconjuntos no vacíos de X , entonces:

$$A = B \text{ si y sólo si } \rightarrow_A = \rightarrow_B$$

Demostración. Obviamente se da la condición necesaria.

Por otra parte, teniendo en cuenta que $\rightarrow_A = (A \times A) \cup (A' \times X) = (X \times X) - (A \times A')$, se obtiene que si $\rightarrow_A = \rightarrow_B$, es decir, si $(X \times X) - (A \times A') = (X \times X) - (B \times B')$, entonces $A \times A' = B \times B'$, lo que es cierto si y sólo si $A = B$. $\square$

Por otra parte, dada una relación borrosa $I : X \times X \rightarrow [0, 1]$, y una t -norma continua T , es conocido ([3]) que I es un T -preorden si y sólo si

$$I = \inf_{\mu \in \mathcal{F}} I_\mu^T,$$

siendo $\mathcal{F}$ el conjunto de todos los T -estados lógicos de I , es decir $\mathcal{F} = \{\mu : X \rightarrow [0, 1] ; T(\mu(x), I(y/x)) \leq \mu(y) \forall x, y\}$.

Al tomar $\mu = \varphi_A$, la función característica clásica de A , es

$$I_{\varphi_A}^T(y/x) = \varphi_{\rightarrow_A} ;$$

por tanto, puede considerarse que los preórdenes elementales borrosos son una generalización de los condicionales materiales clásicos. Y de ahí la pregunta de si cabe extender el teorema 1 a los preórdenes borrosos.

2 Igualdad de preórdenes elementales

Comenzamos estudiando la equivalencia para la t-norma Min , obteniendo que si los preórdenes I_μ^{Min} , I_η^{Min} son iguales, entonces μ y η se diferencian a lo más en los valores máximos de μx y ηx .

Teorema 2 $[I_\mu^{Min} = I_\eta^{Min}]$ si y sólo si $[\mu x \neq \eta x \Rightarrow \mu x \geq \mu y$ y $\eta x \geq \eta y$ para todo $y]$.

Demostración. Sea $I_\mu^{Min}(y/x) = I_\eta^{Min}(y/x)$, para todo x, y . Sabemos que:

$$\begin{aligned} I_\theta^{Min}(y/x) &= \sup\{z; Min(\theta x, z) \leq \theta y\} \\ &= \begin{cases} 1, & \text{si } \theta x \leq \theta y \\ \theta y, & \text{si } \theta x > \theta y. \end{cases} \end{aligned}$$

Tomemos x tal que $\mu x \neq \eta x$, y supongamos que existe un y verificando $\mu x < \mu y$. Obtendremos $I_\mu^{Min}(y/x) = 1 = I_\eta^{Min}(y/x)$ y, por tanto, $\eta x \leq \eta y$.

Por otra parte, $I_\mu^{Min}(x/y) = \mu x = I_\eta^{Min}(x/y) < 1$, y así $I_\eta^{Min}(x/y) = \eta x = \mu x$, llegando a una contradicción; por tanto, ha de ser $\mu x \geq \mu y$ para todo y . De forma similar, se obtiene que $\eta x \geq \eta y$ para todo y .

Recíprocamente, supongamos que si es $\mu x \neq \eta x$ entonces es $\mu x \geq \mu y$ y $\eta x \geq \eta y$, para todo y . Comprobemos la igualdad $I_\mu^{Min} = I_\eta^{Min}$.

- Si $\mu x = \eta x$ y $\mu y = \eta y$, obviamente $I_\mu^{Min}(y/x) = I_\eta^{Min}(y/x)$.

- Si $\mu x = \eta x$ y $\mu y \neq \eta y$, será $\mu y \geq \mu x$, $\eta y \geq \eta x$, y $I_\mu^{Min}(y/x) = I_\eta^{Min}(y/x) = 1$.

- Si $\mu x \neq \eta x$ y $\mu y = \eta y$, se obtiene $\mu x \geq \mu y$, $\eta x \geq \eta y$, y $I_\mu^{Min}(y/x) = \mu y = \eta y = I_\eta^{Min}(y/x)$.

- Por último, si es $\mu x \neq \eta x$ y $\mu y \neq \eta y$, por ser $\mu x \geq \mu y$, $\eta x \geq \eta y$, y $\mu y \geq \mu x$, $\eta y \geq \eta x$, obviamente es $\mu x = \mu y$, $\eta x = \eta y$, y $I_\mu^{Min}(y/x) = I_\eta^{Min}(y/x)$. $\square$

Nota. Si T es arquimediana con generador aditivo h ([1]), entonces de:

$$I_\mu^T(y/x) = \sup\{z / T(\mu x, z) \leq \mu y\} = \sup\{z / h^{(-1)}(h(\mu x) + h(z)) \leq \mu y\},$$

se sigue que:

- Si $\mu x \leq \mu y$, es $h(\mu x) \geq h(\mu y)$, y $h(\mu x) + h(z) \geq h(\mu y)$ para todo z . Por tanto, $h^{(-1)}(h(\mu x) + h(z)) \leq h^{(-1)}h(\mu y) = \mu y$ para todo z , y $I_\mu^T(y/x) = 1$.

- Si $\mu x > \mu y$, $h(\mu x) < h(\mu y)$ y $h(\mu y) - h(\mu x) \in [0, h(0)]$.

Veamos que $I_\mu^T(y/x) = h^{(-1)}(h(\mu y) - h(\mu x))$.

En efecto, $h^{(-1)}(h(\mu x) + h \circ h^{(-1)}(h(\mu y) - h(\mu x))) = h^{(-1)}(h(\mu x) + h(\mu y) - h(\mu x)) = h^{(-1)}(h(\mu y)) = \mu y$ y para todo $z > h^{(-1)}(h(\mu y) - h(\mu x))$, es $h(z) < h(\mu y) - h(\mu x)$, $h(\mu x) + h(z) < h(\mu y) \leq h(0)$ y $h^{(-1)}(h(\mu x) + h(z)) > h^{(-1)}(h(\mu y)) = \mu y$.

Por tanto, se obtiene

$$I_\mu^T(y/x) = \begin{cases} 1, & \text{si } \mu x \leq \mu y \\ h^{(-1)}(h(\mu y) - h(\mu x)), & \text{si } \mu x > \mu y. \end{cases}$$

El siguiente teorema afirma que, en el caso de las t-normas arquimedias estrictas, los preórdenes asociados a μ y η son iguales si y sólo si μ y η tienen "formas similares".

Teorema 3 Si T es una t-norma arquimediana estricta, con generador aditivo h , es

$[I_\mu^T = I_\eta^T]$ ssi [existe un $k \in R$ tal que para todo x es $h(\mu x) = k + h(\eta x)$].

Demostración. Por ser T estricta, $h^{(-1)} = h^{-1}$. Sea $I_\mu^T = I_\eta^T$, y elijamos x, y cualesquiera.

Primeramente observemos que si $\mu x < \mu y$, es $I_\eta^T(x/y) = I_\mu^T(x/y) = h^{(-1)}(h(\mu x) - h(\mu y)) < 1$; por tanto $\eta x < \eta y$. De la misma forma, si $\eta x < \eta y$, es $\mu x < \mu y$.

Así, si es $\mu x < \mu y$, $I_\mu^T(x/y) = h^{-1}(h(\mu x) - h(\mu y)) = I_\eta^T(x/y) = h^{-1}(h(\eta x) - h(\eta y))$, ssi $h(\mu x) - h(\mu y) = h(\eta x) - h(\eta y)$, ssi $h(\mu x) - h(\eta x) = h(\mu y) - h(\eta y)$.

De la misma forma, si $\mu y < \mu x$, es $h(\mu y) - h(\eta y) = h(\mu x) - h(\eta x)$.

Por último, si $\mu x = \mu y$, es $\eta x = \eta y$, y de nuevo $h(\mu x) - h(\eta x) = h(\mu y) - h(\eta y)$.

Por tanto el valor $h(\mu x) - h(\eta x)$ es una constante k para todo x , y $h(\mu x) = k + h(\eta x)$ para todo x .

Recíprocamente, si existe $k \in R$ tal que para todo x es $h(\mu x) = k + h(\eta x)$, al ser $h(\mu x) - h(\eta x) = k = h(\mu y) - h(\eta y)$, se obtiene $h(\mu x) + h(\eta y) = h(\mu y) + h(\eta x)$; entonces $\mu x \leq \mu y$ ssi $h(\mu x) \geq h(\mu y)$, ssi $h(\eta x) \geq h(\eta y)$, ssi $\eta x \leq \eta y$.

Por tanto, si $\mu x \leq \mu y$, es $\eta x \leq \eta y$ y $I_\mu^T(y/x) = 1 = I_\eta^T(y/x)$; y si $\mu x > \mu y$, es $\eta x > \eta y$, y $I_\mu^T(y/x) = h^{-1}(h(\mu y) - h(\mu x)) = h^{-1}(k + h(\eta y) - k - h(\eta x)) = h^{-1}(h(\eta y) - h(\eta x)) = I_\eta^T(y/x)$. $\square$

Particularizando al caso en que μ y η "tienen puntos", es decir, están normalizados, obtenemos:

Corolario 1 Si T es arquimediana estricta y μ y η son tales que existen x, y con $\mu x = 1, \eta y = 1$, entonces

$$I_\mu^T = I_\eta^T \text{ si y sólo si } \mu = \eta.$$

Demostración. Que la condición es suficiente es obvio.

Por otra parte si $I_\mu^T = I_\eta^T$, existe $k \in R$ tal que para todo z es $h(\mu z) = k + h(\eta z)$.

En particular, $h(\mu x) = 0 = k + h(\eta x) \leq h(\eta x)$, y ha de ser $k \leq 0$.

Y análogamente, $h(\mu y) = k + h(\eta y) = k + 0$, y por tanto $k \geq 0$.

Luego $k = 0$, y así $h(\mu z) = h(\eta z)$, y $\mu z = \eta z$ para todo z . $\square$

Por consiguiente, para el caso en que μ y η "no tienen algunos puntos", se obtiene el

Corolario 2 Si T es arquimediana estricta y μ y η son tales que existen x, y con $\mu x = 0, \eta y = 0$, entonces

$$I_\mu^T = I_\eta^T \text{ si y sólo si } \mu = \eta$$

Demostración. Si $I_\mu^T = I_\eta^T$, existe $k \in R$ con $h(\mu z) = k + h(\eta z)$ para todo z .

Por ser $\mu x = 0$, es $h(0) = k + h(\eta x) \leq k + h(0)$, con lo que $0 \leq k$.

Y por ser $\eta y = 0$, $h(\mu y) = k + h(0) \leq h(0)$, y $k \leq 0$. Llegamos a $k = 0$ y $\mu = \eta$. $\square$

Teorema 4 Si T es una t -norma arquimediana no estricta, con generador aditivo h , entonces:

$[I_\mu^T = I_\eta^T]$ si y sólo si $[\text{existe } k \in R \text{ tal que para todo } x \text{ es } h(\mu x) = k + h(\eta x)]$.

Además, si esto ocurre, para todo x es $k \leq h(0) - h(\eta x)$.

Demostración. Sea $I_\mu^T = I_\eta^T$, y elijamos x, y cualesquiera.

- Si es $\mu x < \mu y$, $h(\mu x) > h(\mu y)$ y $I_\mu^T(x/y) = h^{(-1)}(h(\mu x) - h(\mu y)) = I_\eta^T(x/y) < 1$, luego $\eta x < \eta y$, $h(\eta x) > h(\eta y)$ y $I_\eta^T(x/y) = h^{(-1)}(h(\eta x) - h(\eta y)) = h^{(-1)}(h(\mu x) - h(\mu y))$, lo que implica (por ser $h(\eta x) - h(\eta y) \in [0, h(0)]$ y $h(\mu x) - h(\mu y) \in [0, h(0)]$) que $h(\eta x) - h(\eta y) = h(\mu x) - h(\mu y)$ y $h(\mu x) - h(\eta x) = h(\mu y) - h(\eta y)$.

- De la misma forma, si es $\mu y < \mu x$, entonces $\eta y < \eta x$ y $h(\mu y) - h(\eta y) = h(\mu x) - h(\eta x)$.

- En el caso de ser $\mu x = \mu y$, es $\eta x = \eta y$ y $h(\mu x) - h(\eta x) = h(\mu y) - h(\eta y)$.

Luego para todo x, y se verifica $h(\mu x) - h(\eta x) = h(\mu y) - h(\eta y)$, y existe $k \in R$ tal que para todo x es $h(\mu x) - h(\eta x) = k$ y $h(\mu x) = k + h(\eta x)$.

Recíprocamente, supongamos que existe $k \in R$ tal que para todo x es $k \leq h(0) - h(\eta x)$ y $h(\mu x) = k + h(\eta x)$.

- Si $\mu x \leq \mu y$, $k + h(\eta x) \geq k + h(\eta y)$, $h(\eta x) \geq h(\eta y)$ y $\eta x \leq \eta y$. Luego $I_\mu^T(y/x) = I_\eta^T(y/x) = 1$.

- Si $\mu x > \mu y$, $k + h(\eta x) < k + h(\eta y)$, $h(\eta x) < h(\eta y)$, y $\eta x > \eta y$.

Por tanto, $I_\mu^T(y/x) = h^{(-1)}(h(\mu y) - h(\mu x)) = h^{(-1)}(k + h(\eta y) - k - h(\eta x)) = h^{(-1)}(h(\eta y) - h(\eta x)) = I_\eta^T(y/x)$.

Por último, en este caso, si existiese x con $k > h(0) - h(\eta x)$ se obtendría la contradicción $h(\mu x) = k + h(\eta x) > h(0) - h(\eta x) + h(\eta x) = h(0)$. Por tanto, $k \leq h(0) - h(\eta x)$ para todo x . $\square$

Observemos que de nuevo, si los preórdenes asociados a μ y a η coinciden, μ y η han de tener "formas similares", pero ahora, además, su distancia está acotada por el valor de $h(0) - \sup_x \{h(\eta x)\}$.

Análogamente, particularizando al caso en que μ y η "tienen puntos", y al caso en que sus complementarios "tienen puntos", se obtiene el:

Corolario 3 Si T es arquimediana no estricta y μ y η son tales que existen x, y con $\mu x = 1, \eta y = 1$, o existen x, y con $\mu x = 0, \eta y = 0$, entonces

$$I_\mu^T = I_\eta^T \text{ si y sólo si } \mu = \eta.$$

Demostración. Similar al caso de las t -normas arquimediana estrictas.

Agradecimientos Los autores están en deuda con el profesor Claudi Alsina (Barcelona) y con los dos "referees" anónimos, por sus sugerencias en el desarrollo de este trabajo.

Bibliografía

- [1] Schweizer, B., Sklar, A., *Probabilistic Metric Spaces*, Elsevier North-Holland, New York (1983).
- [2] Trillas, E., Cubillo, S., Rodríguez, A., *An Essay on Name and Extension of Rule-given properties*, IPMU'94, París, 1994.
- [3] Trillas, E., Valverde, L., An inquiry on indistinguishability operators, Skala et. al., Eds., *Aspects of Vagueness* (Reidel, Dordrecht, 1984) 231-256.

Setup Semantics of a System of Fuzzy Deontic Logic

Fco. José Díez Ausín
C.A.L.I.J.
Avda. Alcalde Elósegui 275
20017 San Sebastián, España
e-mail: ylbdiauf@sf.ehu.es

Lorenzo Peña
Instituto de Filosofía del CSIC
Pinar 25
28006 Madrid, España
e-mail: lavinius@cc.csic.es

Resumen

En este trabajo presentamos un sistema de lógica deóntica fuzzy, que es una versión modificada de otro anterior, y esbozamos una semántica para el mismo. Igualmente, discutimos algunas cuestiones importantes en teoría del derecho y lógica deóntica, como son los principios de permisión y de agregación y el problema de las lagunas y su relación con el tercio excluso. Este análisis muestra a las claras la adecuación y pertinencia del sistema aquí propuesto, que pivota sobre dos ideas clave: (i) que las situaciones normativas no pueden ser identificadas con un único código, y (ii) que las calificaciones deónticas no son cuestiones de todo-o-nada, sino que más bien la obligatoriedad, prohibición y permisión presentan grados e infinitos matices.

Palabras clave: lógica deóntica, gradualismo, problemas en teoría del derecho.

1 Introduction

In a previous paper we have developed a system of fuzzy deontic logic based on a nonconservative extension of Anderson & Belnap's system of entailment $\mathbb{E}$. Unlike our previous paper the present one is based, not on $\mathbb{P5}$, but on $\mathbb{P10}$. The chain of systems $\mathbb{P5} \dots \mathbb{P10}$ is constituted by nonconservative extensions of entailmental logic $\mathbb{E}$ proposed by Anderson & Belnap; unlike $\mathbb{P5}$, $\mathbb{P10}$ contains a classical negation and is a conservative extension of classical logic. The present paper offers a semantic for our system of deontic logic.

But first we briefly present the (Hilbert-style) axiomatic system:

Primitive symbols: $\wedge$, $\neg$, $\rightarrow$, $\vdash$, $\Box$, $\Diamond$ etc are used as schematic letters. Notational conventions are à la Church: all two-place functors have the same weight and associate leftwards; a dot after a two-place functor stands for a left parenthesis with its mate as far to the right as possible.

Definitions: $[p \vee q]$ abbr $[N(Np \wedge Nq)]$; $[\forall p]$

abbr $[N(p \rightarrow Np)]$

$\mathbb{P0} = \mathbb{E}$

$\mathbb{P1} = \mathbb{P0} + \text{Factor (i.e. } [p \rightarrow q \rightarrow .p \rightarrow .p \wedge q])$

$\mathbb{P2} = \mathbb{P1} + \text{Linearization } ([p \rightarrow q \vee .q \rightarrow p])$

$\mathbb{P3} = \mathbb{P2} + \text{the Self-Self Principle } ([N(p \rightarrow p \rightarrow N(p \rightarrow p)) \rightarrow .p \rightarrow p])$

$\mathbb{P4} = \mathbb{P3} + \mathbf{IF}$ [Implicational Funnel] $([p \rightarrow q \vee .p \rightarrow q \rightarrow r])$

$\mathbb{P5} = \mathbb{P4} + \text{one of these: Aristotle } ([N(p \rightarrow Np)]),$

Boethius $([p \rightarrow q \rightarrow N(p \rightarrow Nq)])$ or Contradiction (namely, for some particular sentential constant j : $[j \wedge Nj]$)

$\mathbb{P3.5} = \mathbb{P3} + \text{both these principles: } [p \rightarrow q \rightarrow .NHNHp \rightarrow Hq] \text{ and } [Hp \rightarrow q \vee .Np \rightarrow r].$

('H' is read: 'It is altogether true that')

$\mathbb{P8} = \mathbb{P3.5} + \text{these two: } [\alpha] \text{ and } [\alpha \rightarrow p \vee .p \rightarrow q].$

('α' is a sentential constant meaning the conjunction of all truths)

$\mathbb{P10} = \mathbb{P8} + [HN(N\alpha \rightarrow \alpha)]$

Implication, ' $\rightarrow$ ', is a connective such that $\vdash_{P10} p \rightarrow r$ only if $\vdash_{CL} p \supset r$, but not conversely. Specifically deontic axiomatic schemata are these ones: ($\vdash p \supset r$ abbr $\vdash L(p \rightarrow r)$;

$\vdash p \vdash r$ abbr $\vdash p \rightarrow r \wedge r \rightarrow p$; $\vdash p = r$ abbr $\vdash L(p \vdash r)$

$\vdash fp$ is read as 'it is forbidden that p' and $\vdash ap$ is read as 'it is allowed that p'. Both 'a' and 'f' are here primitive operators.

Between any two facts, that-p and that-q, there may be a relation of causation, which we express as ' $\triangleleft$ '. A statement ' $p \triangleleft q$ ' means that

the fact that p causes the fact that q. (In fact what we have in mind is not exactly [mere] causation but a more particular relation of "constraining": your choosing to go to the cinema causes you not to attend the meeting but it does not constrain you to be absent from the meeting.)

Ax1 $p = r \rightarrow fp \vdash r$

Ax2 $p \triangleleft q \wedge p \supset q$ (the causal consequences of a fact are facts)

Ax3 $p \triangleleft q \rightarrow .a Nq \rightarrow fp$ (it is forbidden to do something which constrains somebody not to perform an action he is rightfully entitled to perform)

Ax4 $fp \rightarrow a Np$ (in so much as an action is forbidden, to refrain from doing it is allowed)

2 Semantic Treatment

A normative setup is a lattice of codes. Each code contains both factual statements and two sorts of deontic statements: permissions and bannings. The content of a ban or a permission can be negative or positive.

Upon any such lattice of codes we recognize an operator, $\circ$, such that $C^1 \circ C^2$ is the code containing all formulae ' q ' such that

$\vdash p \rightarrow q$ is contained by C^1 while C^2 contains $\vdash p$. Some additional constraints have to be

imposed upon $\circ$, but we will explore them in a further paper.

In each setup there are three special sets of codes: (1) the minimum code, which represents utter falsity; (2) a topmost ranking code; and (3) a special code, α , such that all designated codes rank at least as high as α . If a code is designated, all codes ranking at least as high as it are designated, too.

For a setup S a valuation $\mathbf{v}$ is a function carrying sentences into sets of codes such that for all sentences ' p ', ' q ':

$\mathbf{v}(p) =: \{z \in S: \vdash p \in z\}$ (if ' p ' is an atomic sentence)

$\mathbf{v}(p \wedge q) =: \{z \in S: \forall x \in S: \vdash p \in x \circ z \ \& \ \vdash q \in x \circ z\}$

$\mathbf{v}(p \vee q) =: \{x \circ z \in S: \vdash p \in z \text{ or } \vdash q \in z\}$

$\mathbf{v}(Np) =: \{z \circ x \in S: \vdash p \rightarrow q \in z \ \& \ \vdash Nq \in x\}$

$\mathbf{v}(p \rightarrow q) =: \{x \in S: x \in S \ \& \ z \in S: \vdash p \in z \ \& \ \vdash q \in x \circ z\}$

$\mathbf{v}(Hp) =: \{x \in S: \forall z, u \in S: \text{if } x = z \circ u, \vdash q \rightarrow p \in z \text{ for every formula } \vdash q\}$

$\mathbf{v}(\alpha) =: \text{the set of all designated codes}$

$\mathbf{v}(fp) =: \{x \circ z \in S: \vdash p \triangleleft q \in x \ \& \ \vdash a Nq \in z\}$

$\mathbf{v}(fp \wedge fq) =: \{x \in S: \forall z \in S: \vdash fp \in z \circ x \ \& \ \vdash fq \in z \circ x\}$

$\mathbf{v}(ap) =: \{z \in S: \forall x \in S: \vdash p \triangleleft Nq \in x \text{ only if } \vdash fq \in x \circ z\}$

A sentence ' p ' is said to hold in a setup S

for a valuation $\mathbf{v}$ iff $\mathbf{v}(p)$ is a superset of the set of designated codes of S.

We define negation of deontic sentences in a recursive way:

$\mathbf{v}(Nfp) =: \mathbf{v}(a Np)$

$\mathbf{v}(Np) =: \mathbf{v}(f Np)$

3 Related Issues in Theory of Law

First, a comment seems now required about **PP** or the **permission principle** to the effect that whatever is not forbidden is allowed (the so-called Alchourrón-Bulygin problem). It has been contended that even if a code remains silent on whether or not an action, A, is

allowed, and also on whether a different action, B, is allowed, we cannot draw the conclusion that people have a legal right to do A and that they also have a right to do B (in accordance with the legal system in operation), since a person's doing A may thwart (or prevent or hamper) another person's doing B; and there is a juridical principle to the effect that nobody has the right of staying in the way of another person's exercising a right.

We solve the problem by, on the one hand, refusing to identify a juridical setup with a single code, and, on the other hand, admitting degrees of duty, prohibition and permission. In many cases it may well happen that A is allowed, and so is B, but yet A is also, to some extent, forbidden — in so much as it prevents another person performing action B — at the same time that, in turn, B is also to some extent prohibited for a similar reason.

Let us take an example: the curator is bound to close the building neither before 6:30 nor after 7:00. The visitors are bound to leave the building before 7:00 (or in other words: they are required to be outside the building after 7:00). Have they the right to remain within until 6:59? Has the curator the right to close at 6:45 if there still are visitors within the building? In virtue of **PP** we can gather that the visitors can remain until 7 p.m.; if the curator closes the doors when they are in the building, they are prevented from leaving; since they are banned from remaining after 7:00, they are bound to leave, and hence they also have the right to leave before 7:00; the doors being closed prevents them from leaving; so it is forbidden to close the doors when the visitors are still in the building before 7:00. But in virtue of the permission principle we also draw the conclusion that the curator has indeed the right to close the doors at any moment after 6:30 — which in turn leads us, by an identical reasoning, to the conclusion that the visitors are in duty bound to leave at 6:30.

Such a "paradox" is just a contradiction, and within our approach it can be treated as a merely partial contradiction, which may turn out to be true (up to a point). Both rights may well exist, even though one of them ranks higher (or perhaps neither ranks higher — it all depends on a broad context within which the normative setup has to be interpreted). But even if the

curator has more right to close the doors than have the visitors to remain, this does not entirely nullify the latter right. Rights, duties, prohibitions are not all-or-nothing matters. That is one of the most important lessons we can learn from fuzzy deontic logic.

Within our approach a course of action can be both allowed, up to a point, and yet also forbidden to some extent. But our system also encompasses a strong assertion operator, 'H', read as 'It is entirely the case that', such that a person's being entirely allowed to do A completely rules out his being forbidden to do A.

It seems to us that, whatever the difficulties surrounding it, **PP** plays a major role in our common juridical thought, and is in fact the implicit base for many of our dearest claims against encroachment. More often than not we — rightly — feel entitled to a course of action, not in virtue of the law explicitly recognizing it, but because of our [implicit] right to do whatever is not expressly prohibited.

To be sure the principle gives rise to widely known dilemmas. For no right exists unless all other people are bound to respect it. (In fact our semantic treatment goes further than that, by laying down that nothing is forbidden except transgression of other people's rights, at the same time that no action is allowed except in so much as all are forbidden to interfere with the action.) Such is the **principle of ensuant obligation**, **PEO** for short. In other words, what makes a right a right for somebody to do some action is nothing else but the duty for everybody else not to stand in the way of his doing that action. In fact the right's existence is probably reducible to the ensuant obligation binding all other people. (We have a case of mutual reducibility.)

Now, whenever we encounter a legal gap in virtue of which two different, but mutually incompatible, courses of action — each of them open to a different person — are such that neither is against the law, both can be regarded as allowed, and therefore, in virtue of **PEO**, both are forbidden!

Our proposal entails that in such cases what we face is a situation wherein both courses of action are to some extent allowed and yet also to some extent prohibited — in so far as each of them stands in the way of the alternative course

of action. The legal setup in operation is constituted by a large number of different codes, each of which would fill in the gaps in a particular way. As a whole, the setup neither uniquely allows one course of action neither uniquely allows the opposite course; accordingly, the setup also fails to uniquely forbid either.

Depending on the ranking of those different codes within the legal setup, one of the two courses of action may be more lawful than the other, and so one of the claimants may have a higher entitlement — which nevertheless does not entirely obliterate the alternative course of action.

Let us take an example we owe to Prof. Iturralde, that of a couple, Mary and Joseph, who have prepared frozen embryos meant to be used in a future implanting — and before the bill on assisted reproduction was passed. Meanwhile they quarrel. Is Joseph entitled to have the embryos destroyed? Or rather has Mary the right to go proceed with the implantation regardless of the divorce? The law is (was) silent on the matter.

Admitting degrees of entitlement seems to us the best way out of the difficulty. Our view is that in such a case both claims are [to some extent] right and both are [up to a point] wrong. Neither course of action is uniquely right. So each of them is to some extent right and also to some extent contrary to the legal setup in operation — which cannot be reduced to a particular code, but includes a broad range of alternative completions thereof.

It seems to us that too strong difficulties surround the two classical approaches, one of which entirely rejects the existence of any legal rights unless and until they are explicitly recognized by the legal *corpus*, while the other unqualifiedly accepts as an unrestrictedly rightful entitlement whatever is not expressly forbidden by the law. The latter approach can hardly cope with the possibility of mutually incompatible courses of action by two different persons, neither of which is prohibited by law. The former approach seems to us to lead to an unacceptable conclusion, namely that the juridical order leaves us unprotected whenever the explicit wording of the legal texts is vague or general.

What is more, the legal-gaps approach — that according to which such courses of action as have not been explicitly forbidden are not *eo ipso* implicitly allowed — seems to us not just fraught with unpalatable consequences but in some sense hard to understand, at least if we identify our having a certain right with other people being obliged not to stand in the way of our exercising such a right. For the approach amounts to this: (1) in such situations, neither course of action is the exercise of a legal right; hence, (2) nobody transgresses the law by preventing or thwarting such a course of action; therefore, (3) any such preventive move cannot be a [legal] ground for charging its performer. But (3) entails that the preventive move is not unlawful, i.e. it is within the pale of the law, and hence that: (4) it is rightful to perform such a preventive move.

In order to avoid the conclusion ((4)), the contender has to claim that some actions, on account of which their performers cannot be charged, are not lawful either. Now, if we accept the reductive approach to rights — for you to have a right to do A is nothing else but for all other people to stand under the obligation not to hinder your doing A —, then having such a right entails that, to the extent that you have indeed the right, you cannot be charged on account of having done A (otherwise the threat of arraignment would obviously be a deterrent, a hindrance against your doing A). Conversely, if you are free from prosecution on account of doing A, then, according to the reductive view, you have indeed the right to do A, unless other people are legally entitled to prevent you doing A. Yet it would be an extremely awkward and irksome situation for an action to be immune from public prosecution but open to private obstruction. To unreservedly allow such situations would mean to withdraw the law's protection.

Coming back to our example about frozen embryos, we think that analogical considerations are to the point for the judges to figure out what of the two rival claims is more legal — in other words which of the codes resulting from attempts to complete the current [explicit] set of laws in operation are more in accordance with “the spirit of the law”, or with the purpose of the legal system, or with natural law. We guess that Mary's right to have the

embryos implanted on her is higher, if being deprived of such an implantation would cause her more grief or distress than the implantation itself would cause Joseph.

We must face an objection to our view: if, in the situations being considered, the two courses of action are to some extent legal and therefore if they are, both, to some extent unlawful — each of them in so far as it obstructs or impedes the other, and also legal, course of action — then their performance can be a basis for legal prosecution, which goes against the principle *nullum crimen sine lege*. We reply that many actions which are not explicitly forbidden by statute law are implicitly prohibited by the legal system as a whole, and hence can constitute grounds for prosecution. In the last resort, natural law is to be fallen back on as a sufficient basis for prosecution, in case more explicit norms fail to protect human rights — each legal system being conceived as a particular implementation of natural law. Thus, if our guess concerning Mary & Joseph's case is right, should Joseph smash the embryos against Mary's will, he'd be liable to pay due reparation. No legal system worth thinking about can fail to implicitly incorporate guidelines to cope with such issues, at least through analogical reasoning. (We are not denying that there are cases in which the legal system fails to give even implicit prevalence to either right; even then it often happens that **in some respects** one of them ranks higher while the other ranks higher in other respects, neither ranking **all in all** higher.)

An outcome of our approach is the maintenance of the principle of excluded middle, **PEM**. Since in all such cases neither course of action is wholly forbidden, both are [to some extent] allowed. Hence each of them is either allowed or forbidden. Other nonclassical approaches to deontic logic waive **PEM**. There are several strategies for the abandonment of **PEM**. One is to maintain an external **PEM**, [$p \vee Np$], rejecting, though, the internal principle, namely that each action is either allowed or else forbidden; in order to make such a strategy workable you only need to relinquish the equivalence between **not-forbidden** and **allowed**. A more daring strategy consists in sacrificing **PEM** altogether. Our present considerations and our logical proposal cast

doubt on any justification for such a strategy on the basis of deontic and juridical reasoning. Our view is that scrapping **PEM** is a desperate course which frequently stems from a failure to recognize degrees of legality — and more generally degrees of duty and moral entitlement.

To be sure, the classical construal of **PEM** is — in virtue of mistaking **to be true** for **to be wholly true** — that each situation is either entirely the case or else not the case at all. Concerning things legal, that means that either you are wholly entitled to do A or else fully forbidden to do A. Such a perspective seems to us unattractive and unrealistic. In fact we firmly believe that the admission of degrees of legality is the most promising and significant aspect of our current proposal. We grant that its acceptance calls for a reshaping of much juridical practice. But such a reshaping is anyway mandatory on the basis of current research in an astonishingly broad range of fields — concerning human and nonhuman reproduction, relations between members of our own species and other animals and so on. Leaving the **all or nothing** approach of rights seems to us an overdue step in our legal evolution.

One of the most serious difficulties our current proposal has to face is entailed by the **principle of aggregation, PA**, namely that to the extent that there is an obligation to do A and there is an obligation to do B, there is one to do both A and B. Many classical and nonclassical approaches to deontic logic have preferred to relinquish **PA**. The existence of normative conflicts seems a reasonably sufficient ground for doing so. Not even the acceptance of degrees of dutifulness and that of legal [partial] contradictions seems to offer a helpful way out, since we are often obliged to do A in a high degree, to also do B in a high degree, while it is [utterly] impossible to do A-and-B.

Alternatively Kant's principle — what is mandatory is possible — could be given up: we could contemplate situations wherein a wholly impossible course of action would be nevertheless obligatory. Such a strategy seems to us unsavoury, to say the least.

For the time being at least, our proposal amounts to this: the degree to which we are obliged to do A is one thing; the degree of our

doing A under such an obligation is quite another. Suppose you are bound to do A and also bound to do B, but you cannot do A-and-B. To be precise, suppose that, to the extent that you do A, you fail to do B — doing A is included in not-doing B. Now, suppose you are fairly bound to do A and also fairly bound to do B, but that it is completely impossible to fairly-do-A and fairly-do-B. Were your being fairly obliged to do A the same as your being obliged to fairly do A, **PA** would entail an overcontradiction, a situation wherein an action would be mandatory — and hence possible — and yet utterly impossible. On the other hand if the degree applies to the obligation only, not to the action it renders mandatory, then no overcontradiction results unless A-and-B (with no degree precisification added) is itself overcontradictory.

Let's come to specifics. Suppose you are a physician who can save a person's life, and who can save another person's life, but not both. If saving A's life and B's life is an utterly impossible task, then **PA** puts you under such an unbearable burden. But we can guess that, while — under prevailing circumstances — you cannot fully save A's life if you save B's life, at least in many cases you can to some extent save both. You can give to each of them a dose of the available remedy which can delay their death.

Be it what it may, our current view is that in most cases the alternatively mandatory courses of action are to some extent concretely [partially] compatible, if only at the price of performing each of them in a low degree in order to allow for the other courses of action to also be materialized to some extent. Quite often a half-performance of conflicting duties can be the least bad solution. Anyway our gradualistic or fuzzy logical approach makes such a prospect worth discussing, whereas classical approaches rule it out altogether.

We must admit that what matters for our proposal is not **PA** so much as its corollary, the principle of alternative rights, **PAR**, namely that, to the extent that you are entitled to do either A or B, you are entitled to do A or else you are entitled to do B. Should **PA** be relinquished, the links between duty and permission ought to be also rescinded or loosened. On the other hand, within our logical framework with two conditionals (an

implication, ' $\rightarrow$ ', and a mere conditional, ' $\supset$ ') and two negations (a strong or classical negation, ' $\neg$ ' and a Łukasiewiczian negation, ' N '), we could envisage a number of alternatives. One of them would be to admit **PAR** only for the conditional, ' $\supset$ ', not for implication, ' $\rightarrow$ ': then, since contraposition does not apply to ' $\supset$ ' formulae, we could eschew **PA** without sacrificing **PAR**. Such an option's drawback is that we seem to need the strong (implicational) **PAR**: not only is it true that, if you are entitled to A or B, then either you are entitled to A or else you are entitled to B; not only that, a stronger (implicational) link seems required; otherwise, you could be entitled to A in a very small degree, and to B not at all, while being entitled in a high degree to A-or-B! Such a legal right would be a mockery of justice. Furthermore, while the straightforward **PA** would then become unavailable, a "perverse", if more refined, version would still be with us, to the effect that, if it is completely forbidden not to do A at all and also completely forbidden not to do B at all, then it is completely forbidden to wholly refrain from doing A-and-B. Proof:

- (2) $a(\neg p \vee \neg q) \supset .a \neg p \vee a \neg q$
- (3) $a \neg(p \wedge q) \supset .a \neg p \vee a \neg q$
- (4) $\neg(a \neg p \vee a \neg q) \supset \neg a \neg(p \wedge q)$
- (5) $\neg a \neg p \wedge \neg a \neg q \supset \neg a \neg(p \wedge q)$
- (6) $HNa \neg p \wedge HNa \neg q \supset HNa \neg(p \wedge q)$
- (7) $Hf \neg p \wedge Hf \neg q \supset Hf \neg(p \wedge q)$

So, despite a lingering qualm we feel concerning **PA**, we think that, upon reflection, we had better accept it, as we have in fact done in our present approach.

References

- [1] DIEZ AUSIN, F.J. y PEÑA, L., "Un sistema de lógica deóntica sin principio de cierre", en BUSTOS, E. de, ECHEVERRÍA, J., PEREZ SEDENO, E. y SANCHEZ BALMASADA, M.I., eds., *Actas del I Congreso de la Sociedad de Lógica, Metodología y Filosofía de la Ciencia en España*, Madrid, UNED, 1993, pp. 44-47.
- [2] PEÑA, L., *Rudimentos de lógica matemática*, Madrid, CSIC, 1991.

Generadores y similitudes duales

D. Boixader

J. Jacas

E.T.S. d'Arquitectura del Vallés
Univ. Politècnica de Catalunya
Sant Cugat del Vallés.
e-mail: boixader@ea.upc.es

E.T.S. d'Arquitectura de Barcelona
Univ. Politècnica de Catalunya
Diagonal 649. 08028 Barcelona.
e-mail: jacas@ea.upc.es

Resumen ¹

La principal idea de este trabajo es la introducción de una dualidad entre los elementos de un conjunto X y sus subconjuntos difusos. Esta dualidad induce una interesante interpretación de los elementos como subconjuntos difusos y viceversa. Como consecuencia natural de este planteamiento, el conjunto de los conjuntos difusos sobre $[0, 1]^X$ puede ser considerado como una extensión clásica de X . En el caso que X esté dotado de un operador de T -indistinguibilidad E y considerando sólo el conjunto H de los generadores de E , el operador E puede ser definido de forma natural sobre la mencionada extensión. Por otro lado, una interpretación dual del teorema de representación, conduce a un teorema similar pero partiendo de un conjunto de puntos como generadores. Finalmente, se muestra también que las T -similitudes sobre $[0, 1]^H$ son una herramienta útil para construir una fundamentación coherente para el estudio de la inferencia por semejanza.

Palabras clave: operador de T -indistinguibilidad, Teorema de Representación, generador, Regla Composicional de Inferencia, Dualidad.

Introducción.

Uno de los puntos más conflictivos de la Teoría de Conjuntos Difusos radica en el hecho de que el punto de partida es un conjunto clásico X con elementos clásicos.

En este trabajo, a través del concepto de dualidad entre los elementos de X y el de sus partes difusas $[0, 1]^X$, ambos, puntos y conjuntos difusos pueden ser considerados como conjuntos difusos, aunque definidos sobre conjuntos diferentes. Por tanto, cuando se definen conjuntos difusos sobre un conjunto clásico X , esta construcción transmite, en cierto modo, la borrosidad a los elementos de X .

Después de una sección preliminar, la introducción de esta idea será el principal objetivo de la sección 2. En la sección 3, se muestra que el conjunto de

los difusos sobre $[0, 1]^X$, puede ser considerado como una extensión de X y consecuentemente, se estudia la compatibilidad de un operador de T -indistinguibilidad sobre el conjunto extendido y obtenemos el Teorema de Representación dual para relaciones de similitud sobre $[0, 1]^X$ tomando puntos como generadores.

En la sección 4, introducimos el concepto de similitud propia y mostramos que la mayor de ellas, denominada similitud natural, resulta ser una herramienta útil para el razonamiento aproximado via la RCI generalizada, dado que esta verifica la propiedad de que la similitud natural entre dos hipótesis es siempre mayor o igual que la similitud de las correspondientes tesis. Se finaliza el trabajo con unas conclusiones.

1 Preliminares.

En esta sección recogemos algunos conceptos referentes a operadores de T -indistinguibilidad ampliamente utilizados en diferentes contextos [2,6,7].

Sea T una t -norma continua, $\hat{T}$ representará su **quasi-inversa** [9] definida por

$$\hat{T}(x|y) = \sup\{\alpha \in [0, 1] \mid T(x, \alpha) \leq y\}$$

Definición 1.1 Dada una t -norma T , [1,8] un operador de T -indistinguibilidad E sobre un conjunto X es una relación difusa en X que satisface

$$1.1.1. E(x, x) = 1, \forall x \in X \text{ (Reflexividad),}$$

$$1.1.2. E(x, y) = E(y, x), \forall x, y \in X \text{ (Simetría),}$$

$$1.1.3. T(E(x, y), E(y, z)) \leq E(x, z), \forall x, y, z \in X \text{ (T-transitividad).}$$

En Valverde [9] se demuestra que cualquier operador de T -indistinguibilidad puede ser construido partiendo de una familia de conjuntos difusos.

Lema 1.2 Dado un subconjunto difuso h de un conjunto X , la relación difusa E_h definida por

$$E_h(x, y) = \hat{T}(\text{Max}(h(x), h(y)) \mid \text{Min}(h(x), h(y))),$$

¹Financiado parcialmente por el proyecto número PB91-0334-C03 de la DGICYT

es un operador de T -indistinguibilidad en X

En particular, $\forall x, y \in [0, 1]$, $E_T(x, y) = T(\text{Max}(x, y) | \text{Min}(x, y))$, es una T -indistinguibilidad en el intervalo unidad y en consecuencia, podemos utilizar la siguiente notación $E_h := E_T \circ (h \times h)$

Teorema de Representación 1.3 [3,6,9] Una relación difusa en un conjunto X es un operador de T -indistinguibilidad si y sólo si existe una familia $\{h_i\}_{i \in I}$ de subconjuntos difusos de X tal que

$$E = \inf_{i \in I} E_{h_i} = \inf_{i \in I} E_{T \circ (h_i \times h_i)}.$$

El Teorema 1.3 sugiere la definición siguiente:

Definición 1.4 Dado un operador de T -indistinguibilidad E en X , un generador [5] de E es un conjunto difuso de X que pertenece a una familia generadora de E en el sentido del teorema precedente.

Denotando por H_E el conjunto de generadores de E , está claro que $h \in H_E$ si y sólo si $E_h \geq E$

Definición 1.5 El mínimo de los cardinales de las familias generadoras de un operador de T -indistinguibilidad E se llama su **dimensión** y una familia con esta cardinalidad se denomina una **base** de E [3,4]

El operador propuesto en la siguiente definición es una herramienta útil para el estudio de la estructura del conjunto H_E

Definición 1.6 Dada una t -norma T y un operador de T -indistinguibilidad E sobre un conjunto X , ϕ_E denotará el operador de $[0, 1]^X$ en $[0, 1]^X$ definido por

$$\phi_E(h)(x) = \sup_{y \in X} \{T(E(x, y), h(y))\}.$$

Con esta definición, un subconjunto difuso h de X es un **T -vector propio** [3] de E si y sólo si $\phi_E(h) = h$.

En Jacas & Recasens [6] se demuestra la equivalencia entre generadores y T -vectores propios.

Teorema 1.7 h es un generador de un operador de T -indistinguibilidad E , si y sólo si h es un T -vector propio de E .

Por lo tanto, H_E puede interpretarse como el conjunto de generadores de E o como el conjunto de sus T -vectores propios.

2 Dualidad entre elementos y subconjuntos difusos.

Observación El concepto de dualidad a que se refiere este artículo no guarda relación con la dualidad en el sentido de las Ternas de De Morgan

Dados un conjunto X y una familia $H \subseteq [0, 1]^X$ de subconjuntos difusos es posible, bajo ciertas condiciones, interpretar los elementos de X ($x \in X$) como subconjuntos difusos de H ($x^* \in [0, 1]^H$). Tal interpretación se realiza mediante la identificación ψ_H que se establece en la Proposición 2.2.

Definición 2.1 Dados X y $H \subseteq [0, 1]^X$, se dice que H **separa elementos en X** , si para todo par $x_1, x_2 \in X$, existe $h \in H$ tal que $h(x_1) \neq h(x_2)$.

Proposición 2.2 Dados X , $H \subseteq [0, 1]^X$ y $[0, 1]^H = \{\alpha : H \rightarrow [0, 1]\}$, se considera la aplicación ψ_H definida por

$$\psi_H(x) : X \rightarrow [0, 1]^H$$

y

$$\begin{aligned} \psi_H(x) : H &\rightarrow [0, 1] \\ h &\mapsto \psi_H(x)(h) = h(x) \end{aligned}$$

Con estas hipótesis ψ_H es inyectiva si, y sólo si H separa elementos en X

Notación. Si no existe ambigüedad con respecto a H , notaremos $\psi_H(x) = x^*$, quedando así $\psi_H(x)(h) = x^*(h) = h(x)$.

Así pues, podemos interpretar $[0, 1]^H$ como una extensión de X en sentido conjuntista clásico.

3 Dualidad y T -indistinguibilidades.

En esta sección se extienden de forma natural los resultados de la sección anterior al caso en que X esté dotado de un operador de T -indistinguibilidad

Proposición 3.1 Sea (X, E, T) una T -indistinguibilidad, $H_E \subseteq [0, 1]^X$ el conjunto de sus generadores y $\psi_{H_E} : X \rightarrow [0, 1]^{H_E}$ definida en 2.1. En tal situación, ψ_{H_E} es inyectiva si y sólo si E es T -indistinguibilidad separadora.

En lo sucesivo se tomarán en consideración únicamente T -indistinguibilidades separadoras.

Representación puntual de una T-indistinguibilidad.

Una lectura dual del Teorema de Representación sugiere que un subconjunto $X_0 \subseteq X$ permite generar una T-indistinguibilidad sobre $[0, 1]^X$ y, por restricción, sobre cualquier subconjunto $S \subseteq [0, 1]^X$.

Proposición 3.2 Dados $X_0 \subseteq X$ y $H \subseteq [0, 1]^X$ la relación difusa $F : H \times H \rightarrow [0, 1]$ definida por:

$$F(h_1, h_2) = \inf_{x \in X_0} E_T(h_1(x), h_2(x)) \quad h_1, h_2 \in H,$$

es una relación de T-indistinguibilidad.

Si E es una T-indistinguibilidad, y $H_E \subseteq [0, 1]^X$ es el conjunto de sus generadores, la proposición 3.1 asegura que la relación difusa $[E^*(\alpha, \beta) = \inf_{h \in H_E} E_T(\alpha(h), \beta(h))]$ para todo $\alpha, \beta : H_E \rightarrow [0, 1]$ es una relación de T-indistinguibilidad. El siguiente teorema expresa la relación entre E y E^* .

Teorema 3.3 Si E es T-indistinguibilidad separadora, la aplicación

$$\begin{aligned} \psi_{H_E} : X &\longrightarrow [0, 1]^{H_E} \\ x &\longmapsto \psi_{H_E}(x) = x^* \end{aligned}$$

es un morfismo inyectivo de T-indistinguibilidades entre (X, E, T) y $([0, 1]^{H_E}, E^*, T)$.

Definición 3.4 Sea E una T-indistinguibilidad separadora sobre un conjunto E . Una representación puntual para E es un par $((H, F), \varphi)$ en donde.

- 1) $H \subseteq [0, 1]^Y$, para un cierto conjunto Y
- 2) F es T-indistinguibilidad sobre H definida por

$$F(h_1, h_2) = \inf_{y \in Y} E_T(h_1(y), h_2(y))$$

- 3) $\varphi : X \rightarrow H$ es un isomorfismo entre T-indistinguibilidades.

Corolario 3.5 Toda T-indistinguibilidad (X, E, T) separadora admite representación puntual.

Definición 3.6 Sean $H \subseteq [0, 1]^X$, F una T-indistinguibilidad separadora sobre H y $\mathfrak{H}_F \subseteq [0, 1]^H$ el conjunto de los generadores de F . Se dirá que F es propia si $\psi_H(X) \subseteq \mathfrak{H}_F$.

Mediante la identificación a través de ψ_H , decir que una T-indistinguibilidad en H es propia significa que los puntos de X se encuentran entre sus generadores.

Las T-indistinguibilidades propias vienen caracterizadas por el siguiente teorema.

Teorema 3.7 Sea F una T-indistinguibilidad sobre $H \subseteq [0, 1]^X$. Las siguientes proposiciones son equivalentes:

a) F es propia.

b) $F \leq \bar{E}$, donde $\bar{E}(h_1, h_2) = \inf_{x \in X} E_T(h_1(x), h_2(x))$.

c) Para todo $x \in X$, $h \in H$, $\sup_{g \in H} T(g(x), F(h, g)) = h(x)$.

Definición 3.8 La relación de T-indistinguibilidad definida en $H \subseteq [0, 1]^X$ mediante:

$$\bar{E}(h_1, h_2) = \inf_{x \in X} E_T(h_1(x), h_2(x)).$$

se llamará T-indistinguibilidad natural en H .

En particular, como consecuencia del Teorema 3.5, la T-indistinguibilidad $\bar{E}$ es la mayor entre todas las posibles T-indistinguibilidades propias definidas sobre H .

4 La T-indistinguibilidad natural $\bar{E}$ y la RCI.

El operador de clausura difusa ϕ_E asociado a una T-indistinguibilidad (X, E, T) se ha definido en 1.6 como:

$$\begin{aligned} \phi_E : [0, 1]^X &\longrightarrow H_E \\ g &\longmapsto \phi_E(g) \end{aligned}$$

$$\begin{aligned} \phi_E(g) : X &\rightarrow [0, 1] \\ x &\mapsto \phi_E(g)(x) = \sup_{u \in X} T(g(u), E(u, x)) \end{aligned}$$

Si μ_x denota la función característica de x , ($\mu_x(u) = 0$ si $u \neq x$, $\mu_x(x) = 1$), se tiene:

Proposición 4.1 a) $E(x, y) = \bar{E}(\phi_E(\mu_x), \phi_E(\mu_y))$ para todos $x, y \in X$

b) $\bar{E}(\mu, \nu) \leq \bar{E}(\phi_E(\mu), \phi_E(\nu))$ para todos $\mu, \nu \in [0, 1]^X$

Dados $A, A' : U \rightarrow [0, 1]$ y $B : V \rightarrow [0, 1]$ sea $B' : V \rightarrow [0, 1]$ obtenido a través de la Regla Composicional de Inferencia generalizada (CRI) mediante la fórmula

$$B'(v) = \sup_{u \in U} T(A'(u), I(A(u)|B(v)))$$

(en donde se ha tomado T como función generadora de modus ponens). Notaremos $B' = RCI_{A,B}(A')$, asociada a la R implicación $I(A(u)|B(v))$

Proposición 4.2 Si $B', B'' : V \rightarrow [0, 1]$ satisfacen $B' = RCI_{A,B}(A')$ y $B'' = RCI_{A,B}(A'')$, se tiene que $\overline{E}_V(B', B'') \geq \overline{E}_U(A', A'')$, donde $\overline{E}_U$ y $\overline{E}_V$ denotan las T -indistinguibilidades naturales en U y V respectivamente.

Así pues, la T -indistinguibilidad natural $\overline{E}$ es un instrumento útil para modelizar procesos de razonamiento aproximado basados en RCI, pues satisface que el grado de T -indistinguibilidad entre las tesis obtenidas a partir de RCI es siempre mayor o igual que el grado de T -indistinguibilidad entre las correspondientes hipótesis.

Conclusiones

Hemos mostrado que cuando definimos una T -indistinguibilidad E sobre $H \subset [0, 1]^X$ (H separadora), la dualidad entre puntos en conjuntos difusos muestra claramente la situación

- E es una similitud propia $E \leq \overline{E}$ (siendo $\overline{E}$ la similitud natural), o
- La similitud E no tiene en cuenta la información dada por algunos puntos de X (¡incluso todos ellos!) vía los valores que los elementos $h \in H$ toman sobre ellos.

Es también interesante señalar que la similitud natural $\overline{E}$ es la mayor de entre las propias y así mismo, es la similitud que supone la menor cantidad de información.

Por otra parte, las similitudes naturales parecen ser una herramienta útil para el razonamiento aproximado dado que la RCI es contractiva bajo su aplicación

Bibliografía

- [1] ALSINA, C., TRILLAS, E. AND VALVERDE, L. On some logical connectives for fuzzy set theory, *J. Math. Anal. Appl.* 93 (1983) 15-26.
- [2] GODO, L., JACAS, J. AND VALVERDE, L. Fuzzy Values in Fuzzy Logic. *International Journal of Intelligent Systems*, 6 (1991) 199-212.
- [3] JACAS, J. On the generators of a T -indistinguishability operator. *Stochastica*, 12 (1988) 49-63.
- [4] JACAS, J. Similarity relations. The calculation of minimal generating families, *Fuzzy Sets and Systems*, 35 (1990) 151-162.
- [5] JACAS, J. AND RECASENS, J. Eigenvectors and generators of fuzzy relations, *Proc. of FUZZ-IEEE'92*, San Diego (1992) 687-694.
- [6] JACAS, J. AND VALVERDE, L. On Fuzzy Relations, metrics and cluster Analysis in: J.L. Verdegay and M. Delgado Eds. *Approximate Reasoning Tools for Artificial Intelligence*, ISR 96 (Verlag TUV, Rheinland, 1990) 21-38.
- [7] KLAUONN, F. AND KRUSE, R. Equality relations as a basis of Fuzzy Control, *Proc. of Fuzz-IEEE'92*, San Diego (1992).
- [8] SCHWEIZER, B. AND SKLAR, A. *Probabilistic Metric Spaces*, (North-Holland, Amsterdam 1983).
- [9] VALVERDE, L. On the structure of F -indistinguishability operators, *Fuzzy Sets and Systems*, 17 (1985) 313-328.

Teoría

Valuaciones y conjuntos convexos de probabilidad*

Andrés Cano, María Teresa Lamata

Departamento de Ciencias de la Computación e I.A.

Universidad de Granada

18071 - Granada - España

e-mail: acu,mtl@robinson.ugr.es

Resumen

Presentamos una forma de representación de los problemas de decisión, que son los Sistemas Basados en Valuaciones. Se describe también una representación gráfica de estos sistemas que son la Redes de Valuaciones. Introducimos los Conjuntos Convexos de Probabilidad como forma alternativa para representar la incertidumbre sobre las variables aleatorias en nuestro problema de decisión, y damos una solución para este tipo de problemas cuando tenemos un problema de decisión canónica y la incertidumbre sobre el estado de la naturaleza se da mediante un conjunto convexo de probabilidades. Finalmente se da un algoritmo para resolver problemas de decisión generales cuando la incertidumbre se expresa con distribuciones de probabilidad.

Palabras clave:

Problema de decisión. Valuación. Conjunto convexo de probabilidad. Incertidumbre.

1 Introducción

Los árboles de decisión son una herramienta flexible, y con una gran utilidad de representación gráfica para los problemas de decisión multietapa. A pesar de estas ventajas, presentan ciertos inconvenientes. Primero, crecen rápidamente en muchos problemas, debido a que representan acciones y eventos. Además las probabilidades de los eventos puede que no estén disponibles en la forma que esta metodología necesita. Por ello se hace necesario un preprocesamiento usando las leyes de la probabilidad. Tercero, los árboles de decisión necesitan distribuciones de probabilidad condicional para cada variable aleatoria. Esto necesita muchas veces del uso de la división, lo cual puede ser una operación innecesaria.

*Este trabajo ha sido financiado por la DGICYT bajo el proyecto PB92-0959

Por ello vamos a estudiar un método alternativo a los árboles para la representación de un problema de decisión. Este método es los *Sistemas Basados en Valuaciones* (VBS). Los VBS tienen una representación gráfica que son las *Redes de Valuaciones*, que hacen hincapie en las características cualitativas de los problemas de decisión, permitiendo la representación de los problemas sin ningún preprocesamiento.

En los VBS representamos el conocimiento con funciones llamadas *valuaciones*. En estos sistemas para hacer inferencias se hace uso de las operaciones de *combinación* y de *marginalización*. La inferencia significa marginalizar todas las variables de la valuación conjunta, que es el resultado de combinar todas las valuaciones.

El trabajo lo planteamos en tres partes: Primero estudiamos la estructura y operaciones posibles en las redes de valuaciones. Segundo planteamos el problema de los conjuntos convexos de probabilidad como resultado de dos ó más informaciones distintas, y por último proponemos un posible algoritmo de resolución.

Para comprender bien como funcionan las distintas representaciones de un problema de decisión vamos a ver un ejemplo tomado de Raiffa. Un buscador de petróleo debe decidir si perforar (d) o no perforar ($\sim d$) en un terreno para sacar petróleo. El no sabe si el pozo está seco (dr), húmedo (we) o empapado (so). En la tabla 1 tenemos la matriz de recompensas de este problema, así como las probabilidades de cada posible estado de la naturaleza, $P(O)$.

Ganancia	d	$\sim d$	$P(O)$
dr	-70,000	0	0.5
we	50,000	0	0.3
so	200,000	0	0.2

Table 1: Matriz de recompensas

El buscador de petróleo podría hacer un sondeo que le ayudaría a determinar la estructura geológica del lugar. El test le revelaría si el terreno no tiene

ninguna estructura (*ns*), una estructura abierta (*os*) o una estructura cerrada (*cs*) que sería lo más esperanzador. En la tabla 2 mostramos las probabilidades condicionadas de los resultados del test a la cantidad de petróleo, $P(R|O)$.

$P(R O)$	<i>ns</i>	<i>os</i>	<i>cs</i>
<i>dr</i>	0.6	0.3	0.1
<i>we</i>	0.3	0.4	0.3
<i>so</i>	0.1	0.4	0.5

Table 2: Probabilidades condicionadas

En figura 1 podemos ver el árbol de decisión para este problema. Para calcular las probabilidades en el árbol hemos tenido que hacer un preprocesamiento de las que teníamos disponibles. El *valor esperado* es de 22,500. La *estrategia óptima* es hacer un test de sondeo, y no perforar si el test revela ninguna estructura (*ns*) y perforar si el test revela una estructura abierta(*os*) o cerrada(*cs*).

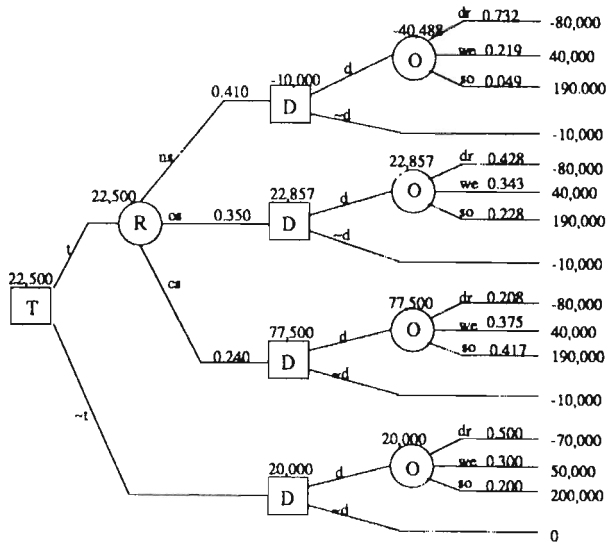


Figure 1: Árbol de decisión

2 Sistemas basados en valuaciones

En los sistemas basados en valuaciones los nodos de decisión se representan como una variable. Los valores de ella son las distintas acciones que puede seguir el decisor. El conjunto de valores de la variable de decisión D se llama el frame de D y lo representamos por $\mathcal{W}_D$. En una red de valuaciones las decisiones las representamos con nodos rectangulares.

Usamos $\mathcal{W}_R$ para representar el frame de la variable aleatoria R . Estas variables se representan en la red con nodos circulares.

Si $\mathcal{X}_D$ es el conjunto de todas las variables decisión y $\mathcal{X}_R$ el conjunto de todas las variables aleatorias entonces $\mathcal{X} = \mathcal{X}_R \cup \mathcal{X}_D$.

Si h es un subconjunto de $\mathcal{X}$, entonces $\mathcal{W}_h$ es el producto cartesiano de los frames de las variables que hay en h . Por ello llamamos a $\mathcal{W}_h$, el frame de h . Los elementos de $\mathcal{W}_h$ los llamamos *configuraciones* de h . El frame para el conjunto vacío $\emptyset$ tiene una única configuración, y usamos el símbolo $\diamond$ para nombrar esta configuración: $\mathcal{W}_\emptyset = \{\diamond\}$.

Si $h \subseteq \mathcal{X}$ entonces una *valuación recompensa* π para h es una función de $\mathcal{W}_h$ en el conjunto de los números reales $\mathbb{R}$. Las valuaciones recompensa las representamos con un rombo en las redes de valuaciones. Dibujaremos arcos no dirigidos entre el nodo de la valuación y todos los nodos correspondientes a sus variables.

Si $h \subseteq \mathcal{X}$ que tiene al menos una variable decisión, entonces un *potencial* ρ para h es una función de $\mathcal{W}_h$ en el intervalo $[0, 1]$. Los potenciales representan probabilidades y en las redes de valuaciones los representamos por triángulos. Dibujamos también arcos no dirigidos entre el potencial y sus variables. Se utilizan los *potenciales condicionales* cuando el potencial representa una probabilidad condicional de una de sus variables aleatoria dadas las demás. En este caso ponemos un arco dirigido entre el potencial y la variable para la que se define el potencial, teniendo ρ que satisfacer:

$$\sum \{\rho(c, r) | r \in \mathcal{W}_R\} = 1 \quad \forall c \in \mathcal{W}_{h - \{R\}} \quad (1)$$

Podemos ver las valuaciones recompensa y los potenciales del problema del buscador de petróleo en las tablas 3 y 4 respectivamente. En la tabla 4 sólo aparecen las probabilidades distintas de cero. En esta misma tabla usamos nr cuando no hay ningún resultado en el test R , debido a que no lo hemos realizado.

$\mathcal{W}_{\{D, O\}}$	Π	$\mathcal{W}_{\{T\}}$	κ
<i>d dr</i>	-70,000	<i>t</i>	-10,000
<i>d we</i>	50,000	<i>~t</i>	0
<i>d so</i>	200,000		
<i>~d dr</i>	0		
<i>~d we</i>	0		
<i>~d so</i>	0		

Table 3: Valuaciones de recompensa pr. buscador petróleo

Denotaremos por $\mathcal{H}_D$ al conjunto de los subconjuntos de $\mathcal{X}$ para los que las valuaciones recompensa existen en el VBS. Supondremos que cada variable decisión de $\mathcal{X}_D$ está incluida en algún elemento de $\mathcal{H}_D$. O sea se cumple:

$$\cup \mathcal{H}_D \supseteq \mathcal{X}_D \quad (2)$$

Denotaremos por $\mathcal{H}_R$ el conjunto de subconjuntos de $\mathcal{X}$ para el que existen los potenciales. Supondremos que cada variable aleatoria está incluida en algún elemento de $\mathcal{H}_R$. O sea se cumple:

$$\cup \mathcal{H}_R \supseteq \mathcal{X}_R \quad (3)$$

Algunas decisiones deben hacerse antes de que ciertas variables aleatorias sean instanciadas (conocido su

$W_{\{O\}}$	ρ	$W_{\{T,R,O\}}$	μ
<i>dr</i>	0.5	<i>t ns dr</i>	0.6
<i>we</i>	0.3	<i>t ns we</i>	0.3
<i>so</i>	0.2	<i>t ns so</i>	0.1
		<i>t os dr</i>	0.3
		<i>t os we</i>	0.4
		<i>t os so</i>	0.4
		<i>t cs dr</i>	0.1
		<i>t cs we</i>	0.3
		<i>t cs so</i>	0.5
		<i>~ t nr dr</i>	1.0
		<i>~ t nr we</i>	1.0
		<i>~ t nr so</i>	1.0

Table 4: Potenciales pr. buscador petróleo

valor) y ciertas variables aleatorias no se instanciarán hasta que ciertas decisiones no sean tomadas. Por ello, es necesario el uso de *restricciones de precedencia*. También podemos tener restricciones de precedencia entre dos nodos de decisión o entre dos variables aleatorias. Usaremos la relación de precedencia $\rightarrow$ para indicar restricciones de precedencia. Las condiciones que impondremos a esta relación son las siguientes:

- La clausura transitiva de la relación, $>$, debe ser un orden parcial de $\mathcal{X}$.
- El orden parcial $>$ debe cumplir que para cualesquiera R y D se debe dar que $D > R$ ó $R > D$.
- Si R es una variable aleatoria sobre la que hay un potencial condicional dados $h - R$, y D es una variable de decisión en h , entonces requerimos que $D > R$.
- Si D es una variable de decisión y existe un potencial para h tal que $D \in h$, entonces requerimos que $D > R$ para cualquier variable aleatoria $R \in h$.

Podemos ver la red de valuaciones para el problema del buscador de petróleo en la figura 2.

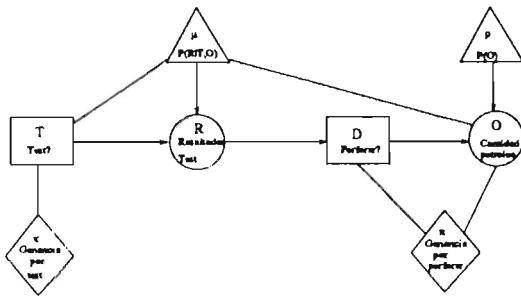


Figure 2: Red de valuaciones

Si g y h son subconjuntos de variables, $h \subseteq g$, y $\mathbf{x}$ es una configuración de g , entonces $\mathbf{x}^{1h}$ es la *proyección* de $\mathbf{x}$ en h . Esta proyección es una configuración de h .

Para resolver sistemas basados en valuaciones usaremos dos operaciones llamadas *combinación* y *marginalización*.

Combinación. La definición de combinación de valuaciones depende del tipo de valuaciones que combinemos. Usaremos la letra π con subíndices para indicar valuaciones de recompensa, y la letra ρ con subíndices para indicar potenciales. Las posibles combinaciones son:

Combinación de valuaciones	Resultado
$(\pi_i \otimes \pi_j)(\mathbf{x}) = \pi_i(\mathbf{x}^{1h}) + \pi_j(\mathbf{x}^{1g})$	V. recompensa
$(\pi_i \otimes \rho_j)(\mathbf{x}) = \pi_i(\mathbf{x}^{1h})\rho_j(\mathbf{x}^{1g})$	V. recompensa
$\rho_j \otimes \pi_i = \pi_i \otimes \rho_j$	V. recompensa
$(\rho_i \otimes \rho_j)(\mathbf{x}) = \rho_i(\mathbf{x}^{1h})\rho_j(\mathbf{x}^{1g})$	potencial

La combinación de valuaciones es conmutativa pero no asociativa en general. Sería asociativa si sólo tuvieramos una valuación recompensa y varias potenciales, ya que la combinación de potenciales si es asociativa. Por ello una forma de combinar todas las valuaciones es combinar todas las valuaciones recompensa obteniendo una única valuación recompensa, y luego combinar con las potenciales.

Marginalización. La definición de la marginalización depende del tipo de variables que sean eliminadas. Si la variable a ser eliminada es aleatoria, la marginalización se lleva a cabo sumando la valuación sobre el frame de la variable a eliminar. En este caso, la valuación puede ser una valuación recompensa o una potencial. Si la variable a eliminar es una variable de decisión, la marginalización se lleva a cabo por maximización (o minimización) sobre el frame de la variable. En este caso, la valuación debe ser una valuación de recompensa.

Puede comprobarse que el orden en que se marginalizan dos variables de decisión en una valuación es irrelevante. Lo mismo ocurre para dos variables aleatorias. Sin embargo el resultado puede ser diferente si las variables son una aleatoria y la otra de decisión. Por ello nosotros haremos la marginalización de la variable Y antes que la X si $X > Y$. Cuando tengamos que marginalizar varias variables en una valuación lo haremos eliminando las variables de una en una, teniendo en cuenta el orden parcial que nos da la relación $>$. Puede haber más de una secuencia posible, pero el resultado es el mismo.

Cada vez que eliminamos una variable de decisión de una valuación recompensa usando maximización, almacenaremos la tabla de valores óptimos de las variables de decisión donde se han alcanzado los máximos. Esta tabla es una función que llamaremos *solución* para la variable decisión, y que definiremos como sigue. Sea h un subconjunto de variables tal que $D \in h$, y ρ es una valuación para h . Una función $\Psi_D : \mathcal{W}_{h-\{D\}} \rightarrow \mathcal{W}_D$ es una solución para D con respecto a π si.

$$\pi^{1(h-\{D\})}(c) = \pi(c, \Psi_D(c)) \quad (4)$$

para todos los $c \in \mathcal{W}_h - \{D\}$.

La tarea principal en un problema de decisión es encontrar una *estrategia*. Una estrategia es una elección de una acción para cada variable decisión D a través de una función de configuraciones de aquellas variables aleatorias R tales que $R \succ D$. Una estrategia es una colección de funciones. Cada una de estas funciones permite conocer la acción que se tomará en una variable decisión conocidas todas las variables aleatorias que pueden influir sobre ella. A estas variables aleatorias las llamamos $Pr(D)$ y se definen como $Pr(D) = \{R \in \mathcal{X}_R | R \succ D\}$.

Llamaremos *problema de decisión canónico* a aquel que tiene una única variable de decisión D , una variable aleatoria R , una valuación de recompensas π para D, R , un potencial ρ para R , y una relación de precedencia $\rightarrow$ definida como $D \rightarrow R$. Este problema se interpreta de la siguiente manera. Los elementos de $\mathcal{W}_D$ son acciones, y los elementos de $\mathcal{W}_R$ son estados de la naturaleza. La potencial ρ es una distribución de probabilidad para R (una probabilidad para cada estado de la naturaleza). La valuación recompensa π es una función condicional de recompensa, si el decisor elige la acción d y se da el estado de la naturaleza r , entonces la recompensa para el decisor es $\pi(d, r)$. La relación de precedencia $\rightarrow$ impone que el verdadero estado de la naturaleza se conoce por el decisor después de que ha elegido una acción.

La recompensa asociada con una acción d es:

$$\sum \{(\pi \otimes \rho)(d, r) | r \in \mathcal{W}_R\} = (\pi \otimes \rho)^{1\{D\}}(d)$$

La máxima recompensa (sea d^* la recompensa máxima) es:

$$\text{MAX}\{(\pi \otimes \rho)^{1\{D\}}(d) | d \in \mathcal{W}_D\} = ((\pi \otimes \rho)^{1\{D\}})^{1\emptyset}(\phi)$$

Una acción d^* será óptima sii:

$$(\pi \otimes \rho)^{1\{D\}}(d^*) = (\pi \otimes \rho)^{1\emptyset}(\phi)$$

3 Conjuntos convexos de probabilidades

El uso de probabilidades simplemente valuadas no nos permite la representación de todos los tipos de incertidumbre que encierra el razonamiento humano. Por ello se introducen los conjuntos convexos de probabilidades. Los convexos son una representación más general que los intervalos de probabilidad, los cuales no son cerrados para la propagación en los casos más simples. Podemos utilizar conjuntos convexos de probabilidad para representar la incertidumbre sobre las variables aleatorias en los problemas de decisión.

Supongamos que tenemos una población, Ω , y una variable, X , que toma valores en un conjunto finito, $U = \{u_1, \dots, u_n\}$. Una información sobre esta variable puede ser una distribución de probabilidad p , es decir, una función:

$$p: U \rightarrow [0, 1]$$

que verifica:

$$\sum_{i=1}^n p(u_i) = 1$$

Cuando el conocimiento acerca de estas frecuencias no es tan exacto, podemos tener un conjunto, H^X , de distribuciones de probabilidad posibles. Este conjunto, H^X , puede definir para cada subconjunto, $A \subseteq U$, dos valores extremos de probabilidades (un intervalo):

$$P_*(A) = \inf_{p \in H^X} \sum_{a \in A} p(a) = \inf_{p \in H^X} P(A)$$

$$P^*(A) = \sup_{p \in H^X} \sum_{a \in A} p(a) = \sup_{p \in H^X} P(A)$$

siendo P la medida de probabilidad asociada a la distribución p . Este par de funciones no son equivalentes al conjunto original, H^X , pues diferentes H^X pueden definir el mismo par de funciones. Se puede asociar un conjunto H a un par de funciones, (P_*, P^*) :

$$H(P_*, P^*) = \{p | P_*(A) \leq P(A) \leq P^*(A)\}$$

Consideraremos que H^X es siempre un conjunto convexo con número finito de puntos extremos, $p_1, \dots, p_k$. Si H^X no es un conjunto convexo, consideraremos su cápsula convexa:

$$\text{CH}(H^X) = \left\{ \sum_i \alpha_i p_i | p_i \in H^X, \alpha_i \in [0, 1], \sum_i \alpha_i = 1 \right\}$$

Podemos definir una valuación potencial usando un conjunto convexo de probabilidades de la misma forma que la definíamos para una distribución de probabilidad normal. Ahora añadiremos una variable transparente por cada variable aleatoria de nuestro problema. Esta variable transparente será la que indique la tendencia del decisor. Es una variable que nos indica que distribución de probabilidad (punto extremo) estamos considerando. De esta forma la manera en que trabajamos utilizando conjuntos convexos de probabilidades para los potenciales es la misma que cuando utilizábamos distribuciones de probabilidad normales. La combinación y marginalización utilizando convexos es igual que utilizando distribuciones de probabilidad. Hay ciertos detalles que hay que especificar como resolveremos. Para ello veamos un ejemplo.

Supongamos el problema de decisión canónica de las tablas 5 y 6. Tras combinar la valuación de recompensa y la potencial obtenemos la valuación de recompensa de la tabla 7. La columna *suma* nos muestra el resultado tras sumar en cada fila, que es equivalente a eliminar la variable aleatoria. Cuando no usábamos conjuntos convexos de probabilidades la decisión óptima se cogía tomando el máximo de los valores de esta columna. Ahora debemos modificar un poco este criterio. Según la distribución de probabilidad t_1 la estrategia óptima es d_1 con una recompensa máxima esperada de 14.1, y según t_2 la estrategia

óptima es d_2 con una recompensa máxima esperada de 12.3. Debemos proporcionar una forma de combinar las informaciones que dan los distintos puntos del conjunto convexo que equivale a eliminar la variable transparente asociada a la variable aleatoria que acaba de borrarse mediante la operación de suma.

Para borrar la variable transparente consideraremos la valuación obtenida tras eliminar la variable aleatoria asociada mediante suma. Tenemos varias posibilidades para eliminar la variable transparente. Una es sumar la información que dan todos los puntos extremos y hacer la media. Otra forma es buscar el máximo y mínimo en cada decisión y hacer la media. Podríamos también tener en cuenta el optimismo o pesimismo del decisor. Hay más posibles formas de combinar las distintas distribuciones de probabilidad, y habría que estudiar cual es más conveniente según cada caso. En el ejemplo tras sumar sobre la variable aleatoria eliminaremos la variable transparente tomando el máximo y mínimo y haciendo la media. Obtenemos los valores 12.9, 12.85, 4.65 para las decisiones d_1, d_2, d_3 respectivamente. Esto nos da una valuación recompensa formada por una variable decisión. Marginalizando esta variable decisión obtenemos que la decisión óptima es d_1 con un valor esperado de 12.9.

Por tanto todo funciona igual con conjuntos convexos de probabilidades, salvo que ahora hemos añadido una variable transparente por cada variable aleatoria de nuestro problema, y que al marginalizar una variable aleatoria debemos eliminar también la variable transparente asociada usando alguno de los criterios mencionados anteriormente.

En un problema de decisión general tendremos informaciones condicionales e informaciones a priori expresadas mediante conjuntos convexos. Cada vez que borremos una variable aleatoria tendremos que borrar la variable transparente asociada.

Π	e_1	e_2	e_3	e_4
d_1	15	7	20	8
d_2	12	15	20	5
d_3	5	-3	5	10

Table 5: Matriz de recompensas

ρ	$P(e_1)$	$P(e_2)$	$P(e_3)$	$P(e_4)$
t_1	0.2	0.1	0.4	0.3
t_2	0.4	0.3	0.1	0.2

Table 6: Convexo de probabilidades

4 Un algoritmo para resolver VBS

Describiremos un algoritmo para resolver VBS usando cálculos locales. Usando el método indicado más arriba para calcular la estrategia óptima y el

$\Pi \otimes \rho$	e_1	e_2	e_3	e_4	suma	óptimo
d_1, t_1	3	0.7	8	2.4	14.1	d_1
d_2, t_1	2.4	1.5	8	1.5	13.4	
d_3, t_1	1	-0.3	2	3	5.7	
d_1, t_2	6	2.1	2	1.6	11.7	d_2
d_2, t_2	4.8	4.5	2	1	12.3	
d_3, t_2	2	-0.9	0.5	2	3.6	

Table 7: Combinación valuaciones

máximo valor esperado requiere trabajar con combinaciones en el espacio $\mathcal{W}_X$. Esto puede ser computacionalmente imposible al trabajar con muchas variables.

En este método se van borrando sucesivamente variables del VBS. La secuencia en que las variables son borradas debe respetar las restricciones de precedencia en el sentido que si $X > Y$, entonces Y debe ser borrada antes que X . Al ser $>$ un orden parcial, el problema puede permitir varias secuencias de borrado, obteniendo la misma respuesta. Sin embargo diferentes secuencias de borrado pueden producir diferentes costos en cuanto a cálculo.

Cuando borramos una variable, tenemos que hacer una operación de *fusión* sobre las valuaciones que contienen a esa variable. Definamos ahora la operación de fusión ayudándonos de la operación de división.

División. Sea α un potencial para g , y $h \subseteq g$. Entonces definimos α/α^{1h} como un potencial para g definido como:

$$(\alpha/\alpha^{1h})(r) = \alpha(r)/\alpha^{1h}(r^{1h}) \quad (5)$$

para todo $r \in \mathcal{W}_R$.

Fusión. Esta operación depende del tipo de variable que va a ser eliminada y de la naturaleza de las valuaciones que contienen a la variable.

Caso 1. Sea D una variable decisión que queremos borrar, y ninguna de las potenciales contiene a D . En este caso con la fusión todas las valuaciones que contienen a D se combinan, y la valuación recompensa resultado se marginaliza para eliminar D . Las valuaciones recompensa que no contienen a D , y todos los potenciales, permanecen inalterados.

Caso 2. Sea R una variable aleatoria que queremos borrar, y ninguna de las valuaciones recompensa contiene a R . En este caso con la fusión el conjunto de valuaciones se cambia como sigue. Las valuaciones recompensa permanece inalteradas. Las potenciales que no contienen a R permanecen inalteradas. Las potenciales que contienen a R se combinan y la potencial resultado se marginaliza para borrar R .

Caso 3. Sea R una variable aleatoria que queremos borrar y todas las valuaciones recompensa contienen a R . En este caso el conjunto de valuaciones se cambia como sigue. Todas las valuaciones recompensa y las potenciales que dependen de R se combinan. La valuación recompensa resultado se marginaliza para borrar R . Las potenciales que no contienen a R permanecen

inalteradas.

Caso 4. Sea R una variable aleatoria que queremos borrar. Supongamos que hay una valuación recompensa que contiene a R , y hay una valuación recompensa que no contiene a R . En este caso con la fusión, el conjunto de valuaciones se cambia como sigue. Las valuaciones recompensa que no contienen a R , y las potenciales que no contienen a R quedan inalteradas. Se crea un nuevo potencial $\rho^{1(g-\{R\})}$ siendo ρ la combinación de todas las potenciales que contienen a R . Combinamos todos los potenciales que contienen a R , y lo dividimos por la potencial que fue creada anteriormente. La potencial resultado se combina con las valuaciones recompensa que contienen a R , y marginalizamos la valuación recompensa resultado para borrar R .

Teorema 1 Supongamos que Δ es un problema de decisión bien definido. Supongamos que $X_1, X_2, \dots, X_k$ es una secuencia de variables de $\mathcal{X} = \mathcal{X}_D \cup \mathcal{X}_R$ tal que con respecto al orden parcial $>$, X_1 es un elemento minimal de $\mathcal{X} - \{X_1\}$, etc. Entonces para obtener el máximo valor esperado podemos ir borrando todas las variables en el orden $X_k, \dots, X_2, X_1$.

En el caso 4 para marginalizar es preciso usar la división. Una forma de evitarlo es estar seguro que tenemos una única valuación recompensa, ya que en ese caso se aplican los casos 2 ó 3. Si tenemos más de una valuación recompensa podemos combinarlas para obtener una única valuación recompensa y entonces comenzar el algoritmo. Otra forma de evitar la división es tener una factorización de la distribución de probabilidad conjunta de tal forma que al sumar sobre una variable en una potencial se obtiene una valuación en que todos los elementos son 1 (*valuación vacua*). Una forma de conseguir esto es tener sólo potenciales condicionales para cada variable aleatoria tal que las variables sobre las que se condiciona siempre preceden a la variable aleatoria. La división debe intentarse para tomar ventaja de una factorización aditiva de la valuación conjunta de recompensa. Si tenemos potenciales arbitrarios, entonces la división puede ser inevitable.

Es necesario encontrar una buena secuencia de borrado para que el número de operaciones sea menor. Pero el problema de encontrar una secuencia de borrado óptima es NP-duro. Hay varias heurísticas para encontrar buenas secuencias de borrado. Una de ellas es la llamada *one-step-look-ahead*, en la que la variable que se borrará es aquella que conduce a combinaciones sobre los frames más pequeños.

5 Conclusiones

En el trabajo hemos estudiado los Sistemas Basados en Valuaciones y su representación gráfica. Hemos visto la forma en que se resuelven los problemas de decisión generales utilizando esta representación. En estos sistemas hemos presentado la posibilidad de utilizar conjuntos convexos de probabilidad como alternativa para representar la incertidumbre en los problemas de decisión. Se ha visto un ejemplo de como se

puede resolver este problema en el caso de un problema de decisión canónico en el que sólo tenemos una variable aleatoria y por tanto un único potencial. De forma similar esperamos poder generalizar esta solución para un problema de decisión general.

Bibliografía

- [1] CAMPOS L.M. DE, HUETE J., MORAL S.. Probability intervals: a tool for uncertain reasoning. Submitted to: *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems* (1993).
- [2] CANO J.E., MORAL S., VERDEGAY-LÓPEZ J.F.. Combination of Upper and Lower Probabilities. In: *Proceedings of the 7th Conference on Uncertainty in A.I.* D'Ambrosio, Smets, Bonissone, eds. (1991), 61-68.
- [3] CANO J.E., MORAL S., VERDEGAY-LÓPEZ J.F.. Propagation of convex sets of probabilities in directed acyclic graphs. In: *Uncertainty in Intelligent Systems* (B. Bouchon-Meunier, L. Valverde, R.R. Yager, eds.) Elsevier (Amsterdam) (1993), 15-26.
- [4] CANO J.E., DELGADO M., MORAL S.. An axiomatic framework for the propagation of uncertainty in directed acyclic networks. In: *International Journal of Approximate Reasoning* 9 (1993), 253-280.
- [5] CANO J.E.. Propagación de probabilidades inferiores y superiores en grafos. DECSAI. Universidad de Granada. (Tesis Doctoral) (1992).
- [6] CHIN H.L., COOPER G.F.. Bayesian Belief Network Inference Using Simulation. *Uncertainty in Artificial Intelligence* 3. (Kanal L.N., Lemmer J.F., eds.) North-Holland (1989).
- [7] MORAL S.. Calculating uncertainty intervals from conditional convex sets of probabilities. In: *Uncertainty in Artificial Intelligence. Proceedings of the 8th Conference* (D. Dubois, M.P. Wellman, B.D'Ambrosio, Ph. Smets, eds.) Morgan & Kaufmann, San Mateo (1992), 199-206.
- [8] MORAL S., CAMPOS L.M. DE. Updating uncertain information. In: *Uncertainty in Knowledge Bases* (B. Bouchon-Meunier, R.R. Yager, L.A. Zadeh, eds.) Springer Verlag, Lectures Notes in Computer Science N. 521 (1991) 58-67.
- [9] NEAPOLITAN R.E.. Probabilistic Reasoning in Expert Systems. Theory and Algorithm. Wiley-Interscience, New York (1990).
- [10] SHENOY P.P.. Valuation-Based Systems for Bayesian Decision Analysis, In: *Operations Research*; Vol. 40, No. 3 (1992).
- [11] SMITH J.A.. Decision Analysis. A Bayesian Approach. Chapman & Hall (1988).

Números intuicionistas fuzzy

P. Burillo López

U. Pública de Navarra
Campus de Arrosadía
31006 Pamplona, España

H. Bustince Sola

U. Pública de Navarra
Campus de Arrosadía
31006 Pamplona, España
e-mail bustince@si.upna.es

V. Mohedano Salillas

U. Pública de Navarra
Campus de Arrosadía
31006 Pamplona, España
vmohedano@si.upna.es

Resumen

En este trabajo proponemos una definición de número intuicionista fuzzy que recupera la ya conocida definición de número fuzzy. Ofrecemos un método de construcción de números intuicionistas fuzzy a partir de números fuzzy. Presentamos diferentes formas de obtención de números fuzzy a partir de números intuicionistas fuzzy, analizando las propiedades de los diferentes números así obtenidos, por último estudiamos perturbaciones de números intuicionistas fuzzy.

Palabras clave: Número intuicionista fuzzy; número fuzzy; operador de Atanassov.

1 Introducción

El desarrollo de la Teoría de Conjuntos Fuzzy ha sido espectacular en las tres últimas décadas. No obstante, existen problemas que para un mejor análisis requieren una filosofía similar a la fuzzy en la que no se tenga en cuenta tan solo, el grado de pertenencia de un elemento al conjunto, sino también su grado de no pertenencia al mismo, entendiendo éste como un valor menor o igual que el complementario de la pertenencia. En este sentido, K. Atanassov en 1986, introdujo el concepto de conjunto intuicionista fuzzy.

Un conjunto intuicionista fuzzy sobre un referencial U es una expresión E de la

siguiente forma:

$$E = \{ \langle x, \mu_E(x), \nu_E(x) \rangle \mid x \in U \}, \text{ donde}$$

$$\mu_E : U \longrightarrow [0, 1]$$

$$\nu_E : U \longrightarrow [0, 1]$$

siempre que $0 \leq \mu_E(x) + \nu_E(x) \leq 1$ para todo x de U , siendo μ_E y ν_E las funciones de pertenencia y de no pertenencia del elemento x al conjunto E respectivamente.

Un estudio completo de tales conjuntos se puede encontrar en ([1], [3] y [4]).

En el trabajo que nos ocupa el conjunto referencial U será $\mathbf{R}$ y por tanto, siempre que hablemos de números intuicionistas fuzzy daremos por entendido que estamos hablando de números intuicionistas fuzzy sobre el conjunto de los números reales.

Se llama *índice intuicionista* del elemento x en el conjunto E a:

$$\pi_E(x) = 1 - \mu_E(x) - \nu_E(x).$$

Existen muchas formas de obtener conjuntos fuzzy a partir de conjuntos intuicionistas fuzzy, una de ellas fue establecida por K. Atanassov en 1989, mediante el siguiente operador:

Sea E un conjunto intuicionista fuzzy y $\alpha \in [0, 1]$

$$D_\alpha(E) =$$

$$= \{ \langle x, (1 - \alpha)\mu_E(x) + \alpha(1 - \nu_E(x)) \rangle \mid x \in U \}.$$

Las propiedades más importantes de este operador se encuentran analizadas en ([2], [4]).

Utilizaremos también el operador $F_{\alpha,\beta}$ definido en ([4]) el cual permite pasar de un conjunto E intuicionista fuzzy a otro $F_{\alpha,\beta}(E)$ también intuicionista, de la forma siguiente:

$$F_{\alpha,\beta}(E) = \{ \langle x, \mu_E(x) + \alpha\pi_E(x), \nu_E(x) + \beta\pi_E(x) \rangle \mid x \in U \}$$

con $\alpha, \beta \in [0, 1]$ tales que $0 \leq \alpha + \beta \leq 1$.

Merece la pena destacar que para $\alpha = 1 - \beta$, este operador coincide con el de Atanassov.

En este trabajo nos proponemos estudiar los números intuicionistas fuzzy. Empezamos dando su definición, analizamos también un método de construcción de estos números a partir de números fuzzy y nos planteamos el proceso contrario, es decir, obtener números fuzzy a partir de números intuicionistas fuzzy por medio del operador de Atanassov, por último, estudiamos la acción del operador $F_{\alpha,\beta}$ sobre los números intuicionistas fuzzy. Entre las posibles definiciones conocidas de número fuzzy elegimos la expuesta en ([5]). Un *número fuzzy* X es un subconjunto fuzzy de la recta real *normal* ($\exists x_0 \in \mathbb{R} \mid \mu_X(x_0) = 1$) y *convexo* ($\mu_X(t) \geq \min\{\mu_X(s), \mu_X(r)\}$ para $r \leq t \leq s$ números reales).

Denotaremos con

$$C_X = \{x \in \mathbb{R} \mid \mu_X(x) = 1\}.$$

2 Números intuicionistas

Definición 1 Un conjunto intuicionista A , es *normal* si existe algún x_0 de $\mathbb{R}$ tal que

$$\mu_A(x_0) = 1.$$

Naturalmente la función de no pertenencia en esos valores vale cero.

Definición 2 Un conjunto intuicionista es *convexo* para la pertenencia si

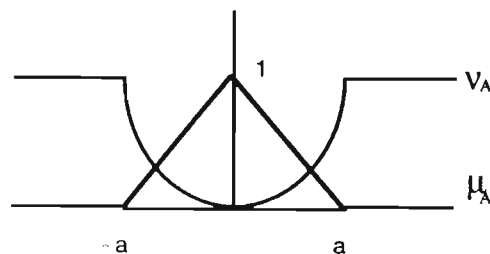
$\mu_A(t) \geq \min\{\mu_A(s), \mu_A(r)\}$ para $s \leq t \leq r$ siendo $s, t, r \in \mathbb{R}$.

Definición 3 Un conjunto intuicionista es *cóncavo* para la no pertenencia si

$\nu_A(t) \leq \max\{\nu_A(s), \nu_A(r)\}$ para $s \leq t \leq r$ siendo $s, t, r \in \mathbb{R}$.

Definición 4 Un *número intuicionista fuzzy* es un subconjunto intuicionista de la recta real, *normal*, *convexo* para la pertenencia y *cóncavo* para la no pertenencia.

Un número intuicionista fuzzy es por ejemplo:



donde

$$\mu_A(x) = \begin{cases} 0 & \text{si } x \leq -a \\ \frac{x}{a} + 1 & \text{si } -a \leq x \leq 0 \\ -\frac{x}{a} + 1 & \text{si } 0 \leq x \leq a \\ 0 & \text{si } x \geq a \end{cases}$$

$$\nu_A(x) = \begin{cases} 1 & \text{si } x \leq -a \\ \frac{1}{a^2}x^2 & \text{si } -a \leq x \leq a \\ 1 & \text{si } x \geq a \end{cases}$$

3 Construcción de números intuicionistas

En este apartado exponemos un método de construcción de números intuicionistas fuzzy a partir de un número fuzzy y estudiamos la forma de recuperar el número

fuzzy utilizado en la construcción por medio del operador de K. Atanassov. Obsérvese que con esta definición, todo número difuso X es un número intuicionista.

Teorema 1 Sea $X = \{ \langle x, \mu_X(x) \rangle \mid x \in \mathbf{R} \}$ un número fuzzy. El conjunto

$$A = \{ \langle x, \mu_A(x), \nu_A(x) \rangle \mid x \in \mathbf{R} \}$$

es un número intuicionista fuzzy, donde

$$\begin{cases} \mu_A(x) = \mu_X(x) \\ \nu_A : \mathbf{R} \longrightarrow [0, 1] \\ x \longrightarrow \nu_A(x) \leq 1 - \mu_A(x) \end{cases}$$

para todo x de $\mathbf{R}$ y tal que $\nu_A(x)$ sea monótona decreciente para los $x \leq \text{Inf}(C_X)$ y monótona creciente para los $x \geq \text{Sup}(C_X)$.

Demostración. Es claro que

$$\begin{aligned} 0 &\leq \mu_A(x) \leq 1 \\ 0 &\leq \nu_A(x) \leq 1 \\ 0 &\leq \mu_A(x) + \nu_A(x) \leq 1 \end{aligned}$$

para todo $x \in \mathbf{R}$.

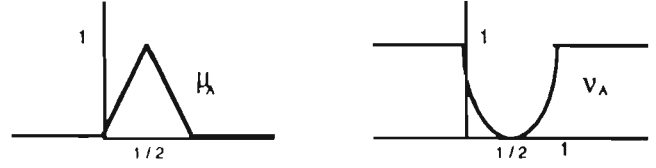
Además μ_A es normal y convexa puesto que μ_X lo es. La concavidad de ν_A queda garantizada por la condición de monotonía impuesta. $\square$

Obsérvese que este teorema permite construir a partir de un número fuzzy infinitas familias de números intuicionistas, bastaría por ejemplo, definir ν_A de la siguiente forma

$$\nu_A(x) = p(1 - \mu_A(x))$$

con $p \in [0, 1]$ para todo $x \in \mathbf{R}$, teniendo de esta forma una familia infinita de números.

A continuación presentamos un ejemplo de esta construcción:



siendo el número fuzzy

$$X = \{ \langle x, \mu_X(x) \rangle \mid x \in \mathbf{R} \}$$

el dado por

$$\mu_A(x) = \begin{cases} 0 & \text{si } x \leq 0 \\ 2x & \text{si } 0 \leq x \leq \frac{1}{2} \\ 2 - 2x & \text{si } \frac{1}{2} \leq x \leq 1 \\ 0 & \text{si } x \geq 1. \end{cases}$$

Tomando

$$\nu_A(x) = \begin{cases} 1 & \text{si } x \leq 0 \\ 4(x - \frac{1}{2})^2 & \text{si } 0 \leq x \leq 1 \\ 1 & \text{si } x \geq 1 \end{cases}$$

que satisface las condiciones exigidas en el Teorema, resulta que el conjunto

$$A = \{ \langle x, \mu_A(x), \nu_A(x) \rangle \mid x \in \mathbf{R} \}$$

es un número intuicionista fuzzy.

Teorema 2 Si A es un número intuicionista fuzzy construido mediante el procedimiento del Teorema anterior a partir del número fuzzy X , entonces

$$D_0(A) = X.$$

Demostración. Basta recordar que

$$\mu_{D_\alpha(A)}(x) = \mu_X(x) + \alpha \pi_A(x) \quad \square.$$

4 Operadores aplicados a un número intuicionista

Este apartado consta de dos subsecciones. En la primera proponemos un Teorema fundamental de caracterización para números intuicionistas. Se estudia además, las condiciones que ha de verificar un número intuicionista fuzzy A , de manera que al aplicarle el operador de K. Atanassov, $D_\alpha(A)$

para cualquier valor de $\alpha \in [0, 1]$, su resultado sea un número fuzzy. En la segunda proponemos un método de perturbación de números intuicionistas fuzzy mediante el operador $F_{\alpha, \beta}$.

4.1 Operador D_α sobre un número intuicionista

El siguiente resultado nos permite caracterizar los números intuicionistas fuzzy por medio del operador de K. Atanassov.

Teorema 3 *A es un número intuicionista fuzzy si y solamente si $D_0(A)$ y $D_1(A)$ son números fuzzy.*

Demostración. $\Rightarrow$) Como $\mu_{D_0(A)}(x) = \mu_A(x)$ se tiene que $D_0(A)$ es número fuzzy. Por otra parte $\mu_{D_1(A)}(x) = 1 - \nu_A(x)$ es normal puesto que para los x_0 tales que $\mu_A(x_0) = 1$, el valor de $\nu_A(x_0)$ es cero por la condición de intuicionismo.

La convexidad de $\mu_{D_1(A)}$ queda garantizada por la concavidad de ν_A y la dualidad de la t-norma mínimo.

$\Leftarrow$) Sean X e Y dos números fuzzy tales que $0 \leq \mu_Y(x) - \mu_X(x) \leq 1$ para todo $x \in \mathbb{R}$. Sea el número A dado por

$$\mu_A(x) = \mu_X(x)$$

$$\nu_A(x) = 1 - \mu_Y(x)$$

evidentemente A es un número intuicionista fuzzy y por tanto

$$\mu_{D_0(A)}(x) = \mu_A(x) = \mu_X(x)$$

$$\mu_{D_1(A)}(x) = 1 - \nu_A(x) = \mu_Y(x) \quad \square$$

El siguiente teorema establece las condiciones bajo las cuales podemos asegurar que a partir de un número intuicionista fuzzy obtenemos siempre un número fuzzy por medio del operador D_α .

Teorema 4 *Sea A un número intuicionista fuzzy, verificando la siguiente condición*

$$\begin{aligned} & \text{Signo}(\mu_A(r) - \mu_A(s)) \neq \\ & \neq \text{Signo}(\nu_A(r) - \nu_A(s)) \end{aligned} \quad (1)$$

para toda pareja de números reales r y s . En estas condiciones, $D_\alpha(A)$ es un número fuzzy para todo α de $[0, 1]$.

Demostración. La condición de normalidad de D_α (para todo $\alpha \in [0, 1]$) es evidente por la normalidad de A y la definición de D_α .

La convexidad de la pertenencia y la concavidad de la no pertenencia de A nos permite asegurar que

$$\begin{aligned} (1 - \alpha)\mu_A(t) & \geq \\ & \geq \min\{(1 - \alpha)\mu_A(s), (1 - \alpha)\mu_A(r)\} \\ \alpha(1 - \nu_A(t)) & \geq \\ & \geq \min\{\alpha(1 - \nu_A(s)), \alpha(1 - \nu_A(r))\} \end{aligned}$$

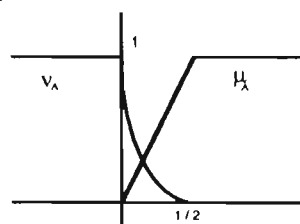
con $s \leq t \leq r$ y para todo $\alpha \in [0, 1]$. Sumando estas dos expresiones y teniendo en cuenta la condición (1) queda probado que

$$\mu_{D_\alpha(A)}(t) \geq \min\{\mu_{D_\alpha(A)}(s), \mu_{D_\alpha(A)}(r)\}$$

y por tanto la convexidad de A . $\square$

Es importante resaltar que la condición (1) es verificada por cualquier número fuzzy y no así por todo número intuicionista fuzzy.

El siguiente ejemplo cumple las hipótesis del Teorema

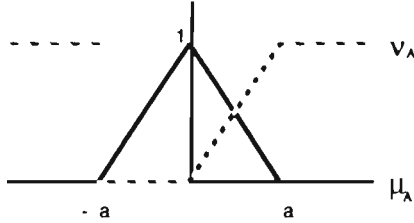


siendo

$$\mu_A(x) = \begin{cases} 0 & \text{si } x \leq 0 \\ 2x & \text{si } 0 \leq x \leq \frac{1}{2} \\ 1 & \text{si } \frac{1}{2} < x \leq 1 \\ 1 & \text{si } x > \frac{1}{2} \end{cases}$$

$$\nu_A(x) = \begin{cases} 1 & \text{si } x < 0 \\ 4(x - \frac{1}{2})^2 & \text{si } 0 \leq x \leq \frac{1}{2} \\ 0 & \text{si } x > \frac{1}{2} \end{cases}$$

Notemos que la condición (1) es excesivamente restrictiva ya que existen números intuicionistas fuzzy que no cumpliéndola verifican la tesis del teorema. Por ejemplo:



siendo

$$\mu_A(x) = \begin{cases} 0 & \text{si } x \leq -a \\ 1 + \frac{1}{a}x & \text{si } -a \leq x \leq 0 \\ 1 - \frac{1}{a}x & \text{si } 0 \leq x \leq a \\ 0 & \text{si } x \geq a. \end{cases}$$

$$\nu_A(x) = \begin{cases} 1 & \text{si } x < -a \\ 0 & \text{si } -a \leq x \leq 0 \\ \frac{1}{a}x & \text{si } 0 \leq x \leq a \\ 1 & \text{si } x \geq a. \end{cases}$$

evidentemente A es un número intuicionista fuzzy que no satisface la condición (1) para todo $r, s \in \mathbf{R}$ y sin embargo, D_α es un número fuzzy para todo $\alpha \in [0, 1]$.

4.2 Perturbación de números intuicionistas fuzzy

En este apartado estudiamos las condiciones que debe cumplir un número intuicionista fuzzy de forma que al perturbarlo con el operador $F_{\alpha, \beta}$ sigamos obteniendo otro número intuicionista fuzzy.

Teorema 5 Sea A un número intuicionista fuzzy, verificando la siguiente condición

$$\begin{aligned} \text{Signo}(\mu_A(r) - \mu_A(s)) &\neq \\ &\neq \text{Signo}(\nu_A(r) - \nu_A(s)) \end{aligned}$$

para toda pareja de números reales r y s . En estas condiciones, $F_{\alpha, \beta}(A)$ es un número intuicionista fuzzy para todo α, β de $[0, 1]$ con $0 \leq \alpha + \beta \leq 1$.

Demostración. La prueba de la normalidad y de la convexidad para la pertenencia es similar a la realizada en la demostración del teorema 4. Probaremos la convexidad de la no pertenencia. La convexidad de la pertenencia y la concavidad de la no pertenencia de A nos permite asegurar que

$$\begin{aligned} (1 - \beta)\nu_A(t) &\leq \\ &\leq \max\{(1 - \beta)\nu_A(s), (1 - \beta)\nu_A(r)\} \\ \beta(1 - \mu_A(t)) &\leq \\ &\leq \max\{\beta(1 - \mu_A(s)), \beta(1 - \mu_A(r))\} \end{aligned}$$

con $s \leq t \leq r$ y para todo $\alpha, \beta \in [0, 1]$ tal que $0 \leq \alpha + \beta \leq 1$. Sumando estas dos expresiones y teniendo en cuenta la condición (1) queda probada la concavidad de la no pertenencia. $\square$

5 Conclusiones

Los conceptos introducidos en este artículo permiten, por una parte, un desarrollo ulterior para construir la aritmética intuicionista difusa como una generalización de la aritmética de intervalos así como introducir conceptos de cálculo intuicionista difuso. Por otra parte pueden construir una herramienta útil para el desarrollo del razonamiento aproximado en general y en particular para el tratamiento de problemas de indistinguibilidad y de estados lógicos en relación con reglas de inferencia de lógicas intuicionistas, problemas de toma de decisión con información no precisa por comparación de números intuicionistas difusos (redes neuronales artificiales) y formalizaciones de la Teoría de la Comunicabilidad. La introducción de la aritmética intuicionista difusa puede manifestarse útil en la construcción de modelos difusos en Lógica y Control.

En este trabajo hemos presentado una definición de número intuicionista fuzzy concreta. No obstante se podrían estudiar otras posibles definiciones que proporcionarían resultados diferentes respecto de la acción de los operadores sobre dichos números.

Por último decir que nos hemos limitado a sentar las bases de estos números, su aritmética y otras propiedades serán objeto de posteriores investigaciones.

Referencias

- [1] K. ATANASSOV Intuitionistic Fuzzy Sets. *Fuzzy Sets and Systems*, **20**: 37-45 (1986).
- [2] K. ATANASSOV More on Intuitionistic Fuzzy Sets. *Fuzzy Sets and Systems*, **33**: 37-45 (1986).
- [3] P. BURILLO Y H. BUSTINCE Estructuras algebraicas en conjuntos intuicionistas fuzzy, II Congreso Espanol de Lógica y Tecnología Fuzzy, Boadilla del Monte, Madrid, (1992).
- [4] H. BUSTINCE Conjuntos Intuicionistas e Intervalo-valorados Difusos: Propiedades y Construcción. Relaciones Intuicionistas Fuzzy. *Memoria de Tesis Doctoral*, Universidad Pública de Navarra, (1994).
- [5] L. A. ZADEH Fuzzy Sets. *Inform. and Control*, **8**: 338-353 (1965).

Negaciones Fractales *

G. Mayor
Dep. Matemàtiques i Informàtica
Universitat de les Illes Balears
07012 Palma de Mallorca, España
e-mail: DMIGMF0@ps.uib.es

T. Calvo
Dep. de Matemàtiques i Informàtica
Universitat de les Illes Balears
07012 Palma de Mallorca, España
e-mail: DMITCS0@ps.uib.es

Resumen

A partir del concepto de atractor de una familia de afinidades contractivas del espacio métrico ordinario R^2 , se estudia la propiedad de fractalidad de la función de De Rham y de otras funciones singulares que se derivan de ésta. En particular, aparecen como fractales las negaciones fuertes llamadas k-negaciones.

Palabras clave: Espacio métrico completo, aplicación contractiva, conjunto compacto, afinidad de R^2 , atractor, conjunto fractal, función de De Rham, negación fuerte, k-negación.

1. Introducción

El objeto de este trabajo es un estudio sobre la fractalidad de las negaciones fuertes. Empezamos por describir el ámbito en el que "viven" los fractales de R^2 , y a continuación se define el concepto de fractal asociado a una

familia de transformaciones de dicho espacio. A tal fin es necesario presentar algunos preliminares

2. Preliminares

Definición 1. Sea (X,d) un espacio métrico. Una transformación $h: X \rightarrow X$ es contractiva si existe una constante r , $0 \leq r < 1$, tal que se verifica $d(h(x), h(y)) \leq r d(x,y)$ para todo $x,y \in X$.

Teorema 1. Sea X un espacio métrico completo, y sea h una transformación contractiva de X . Se verifica que h posee un único punto fijo $s \in X$ ($h(s) = s$). Por otra parte, para cualquier $x \in X$, la sucesión $(h^n(x))$ converge a s .

Consideremos, ahora, el plano euclídeo (espacio métrico completo (R^2, d) siendo d la distancia euclídea). Indicaremos mediante $K(R^2)$ el conjunto cuyos elementos son los subconjuntos compactos (no vacíos) de R^2 . Si $x \in X$ y $A \in K(R^2)$ se define la distancia entre x

y A mediante: $d(x,A) = \min \{ d(x,y) ; y \in A \}$.
 Si $A, B \in K(\mathbb{R}^2)$, se define la distancia entre A y B por: $d(A,B) = \max \{ d(x,B) ; x \in A \}$.
 Finalmente, para A y B como antes, se define la distancia (Hausdorff) entre A y B : $d_H(A,B) = \max \{ d(A,B), d(B,A) \}$.

Observemos que si $F: [0,1] \rightarrow [0,1]$ es una función continua, entonces F es un compacto de $\mathbb{R}^2$ ($F \in K(\mathbb{R}^2)$).

Puede probarse el siguiente resultado ([1])

Teorema 2. $(K(\mathbb{R}^2), d_H)$ es un espacio métrico completo.

Una transformación $w: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ de la forma

$$w \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} + \begin{pmatrix} e \\ f \end{pmatrix}$$

con a, b, c, d, e, f números reales, es una afinidad. Puede demostrarse fácilmente que una afinidad w es contractiva si, y sólo si, verifica $ab+cd=0$, $a^2+c^2 < 1$ y $b^2+d^2 < 1$. Estamos aquí interesados en afinidades del tipo "cambio de escala" y que sean contractivas, o sea en transformaciones de la forma

$$w \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} a & 0 \\ 0 & d \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} + \begin{pmatrix} e \\ f \end{pmatrix} \quad 0 < a, d < 1$$

Si consideramos un conjunto finito $\{w_1, \dots, w_n\}$ de afinidades contractivas, entonces $W: K(\mathbb{R}^2) \rightarrow K(\mathbb{R}^2)$ definida por

$$W(X) = \bigcup_{k=1}^n w_k(X)$$

es contractiva (respecto de d_H). Al único

punto fijo de W lo llamamos el atractor, A , de $\{w_1, \dots, w_n\}$ y es $A = \lim_{k \rightarrow \infty} W^k(X)$ para cualquier X de $K(\mathbb{R}^2)$. Diremos, también, que A es el fractal definido por el conjunto $\{w_1, \dots, w_n\}$.

Definición 2. Una función $N: [0,1] \rightarrow [0,1]$ es una negación fuerte si es involutiva ($N^2 = id$) y decreciente.

Definición 3. Sea k , $0 < k < 1$. Una negación fuerte N se llama una k -negación si verifica $N(kx) = (1-k)N(x) + k$ para todo x de $[0,1]$.

Para más detalles sobre funciones de negación y k -negaciones puede consultarse ([2] a [4] y [6]).

3. La función de De Rham

Dada una constante k , $0 < k < 1$, la única función $R: [0,1] \rightarrow [0,1]$ solución del sistema de ecuaciones funcionales siguiente:

$$\left. \begin{aligned} R\left(\frac{x}{2}\right) &= k R(x) \\ R\left(\frac{x+1}{2}\right) &= (1-k) R(x) + k \end{aligned} \right\} 0 \leq x \leq 1 \quad (DR[k])$$

viene dada por

$$R(x) = \begin{cases} 0, & \text{si } x = 0 \\ \sum_{n=1}^{\infty} \left(\frac{1-k}{k}\right)^{n-1} k^{\alpha_n}, & \text{si } x = \sum_{n=1}^{\infty} \frac{1}{2^{\alpha_n}} \\ \text{con } \alpha_1 < \alpha_2 < \dots \text{ enteros positivos} \end{cases}$$

Esta función es estrictamente creciente, con $R(1)=1$, y continua. Si $k=1/2$ entonces $R(x)=x$. Si $k \neq 1/2$, entonces su derivada es cero casi por

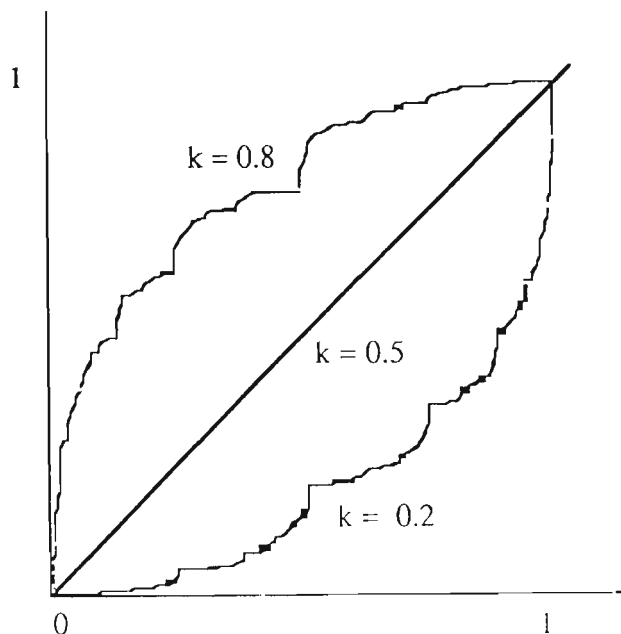
todo en $[0,1]$. La función R se llamará la función de De Rham de parámetro k (o k -función de De Rham). Para más detalles sobre esta función puede consultarse ([5]).

Teorema 3. Dado k , $0 < k < 1$, la función de De Rham de parámetro k , R , es el atractor de $\{w_1, w_2\}$ donde w_1, w_2 están definidas por:

$$w_1 \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 1/2 & 0 \\ 0 & k \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$$

$$w_2 \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 1/2 & 0 \\ 0 & 1-k \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} + \begin{pmatrix} 1/2 \\ k \end{pmatrix}$$

Demostración. Evidentemente $R \in K(R^2)$. Por otra parte, w_1 y w_2 son afinidades contractivas que verifican, como puede comprobarse fácilmente, $w_1(R) \cup w_2(R) = R$; por tanto R es el atractor de $\{w_1, w_2\}$. Vale la pena observar que $w_1(R) \cap w_2(R) = \{(1/2, k)\}$. En la siguiente figura aparece la representación gráfica del fractal R para distintos valores de k



4. Las funciones $N_{kk'}$. k -negaciones.

Hace algún tiempo, y en relación con unos estudios sobre dualidad de una clase de operaciones binarias en el intervalo real $[0,1]$, utilizamos unas funciones dependientes de dos parámetros (k y k') $N_{kk'}$, que ahora reconsideramos con el fin de descubrir, al igual que para las funciones de De Rham, su carácter fractal.

Dados k, k' con $0 < k, k' < 1$, la única solución $N: [0,1] \rightarrow [0,1]$ del sistema de ecuaciones funcionales

$$\begin{cases} N(kx) = (1-k')N(x) + k' \\ N((1-k)x + k) = k'N(x) \end{cases} \quad \left. \begin{array}{l} 0 \leq x \leq 1 \\ (\Sigma(k, k')) \end{array} \right\}$$

es $N(x) = R'(1-R^{-1}(x))$, siendo R y R' las funciones de De Rham de parámetros k y k' respectivamente. Para detalles y demostraciones puede consultarse ([3]).

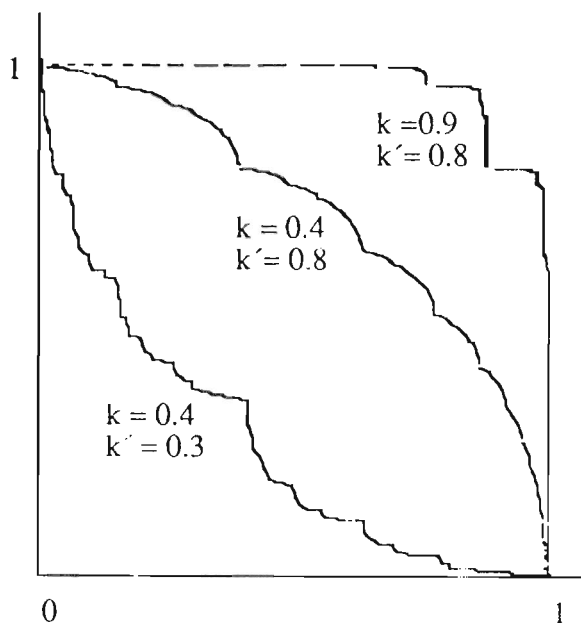
Teorema 4. Dados $k, k': 0 < k, k' < 1$, sea $N_{kk'}$ la solución del sistema $\Sigma(k, k')$ (que indicaremos simplemente por N), es decir, $N(x) = R'(1-R^{-1}(x))$. Entonces N es el atractor de $\{w_1, w_2\}$ siendo

$$w_1 \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} k & 0 \\ 0 & 1-k' \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} + \begin{pmatrix} 0 \\ k' \end{pmatrix}$$

$$w_2 \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 1-k & 0 \\ 0 & k' \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} + \begin{pmatrix} k \\ 0 \end{pmatrix}$$

Demostración. Para empezar N es un compacto de R^2 . Las transformaciones w_1, w_2 son,

evidentemente, afinidades contractivas y verifican $w_1(N) \cup w_2(N) = N$. Por tanto N es el atractor de $\{w_1, w_2\}$. Es fácil comprobar, por otra parte, que $w_1(N) \cap w_2(N) = \{(k, k')\}$. En la Fig. 2 se muestran distintos gráficos del fractal N para varios valores de k y k' .



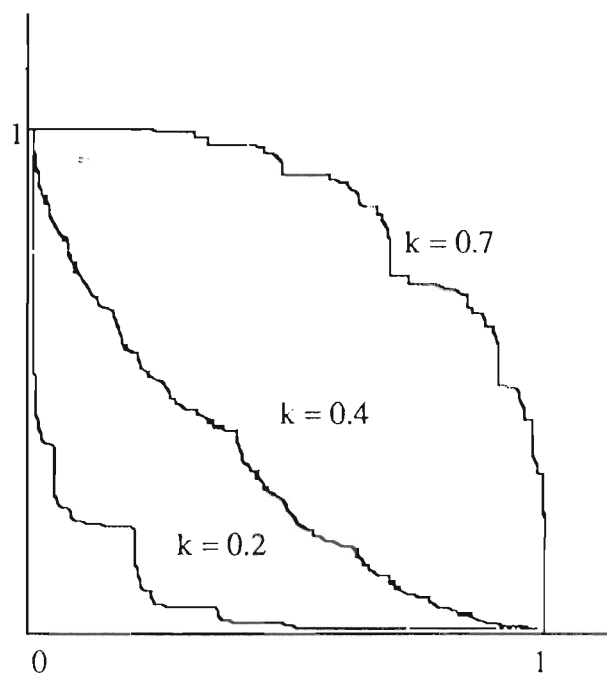
$N_{kk'}$ (Fig. 2)

Corolario 1. *Las k -negaciones son fractales.*

Demostración. En el caso $k=k'$ se obtiene como solución del sistema $\sum (k, k)$ la función $N(x) = R(1-R^{-1}(x))$ que resulta ser involutiva y por tanto es una negación (fuerte). Esta es la única negación fuerte que verifica la igualdad

$$N(kx) = (1-k)N(x) + k \text{ para todo } x \text{ de } [0,1].$$

Algunos gráficos de k -negaciones aparecen en la Fig 3.



N_k (Fig. 3)

Agradecimientos

A los profesores M^a. Luisa Santamaría y Gabriel Oliver del Departamento de Matemáticas e Informática de la U.I.B., por su valiosa ayuda en la elaboración de los gráficos que acompañan el desarrollo de este artículo.

Bibliografía

- [1] Barnsley, M.F. *Fractal Everywhere*, Academic Press Inc., New York, 1988.
- [2] Fraile, A., Mayor, G. y Monserrat, M., *k-negacions*. Actas del IV Congrés Català de Lògica. Barcelona (1985), 67-68.

- [3] Mayor, G y Torrens, J. *Duality for a class of binary operations on $[0,1]$* . Fuzzy Sets and Systems 47 (1992), 77-80.
- [4] Mayor, G y Torrens, J. *De Rham systems and the solution of a class of functional equations*. Aequationes Mathematicae 47 (1994), 43-49.
- [5] Takács, L. *An increasing continuous singular function*. Math. Monthly (1978), 35-37.
- [6] Trillas, E., *Sobre funciones de negación en la teoría de conjuntos difusos*. Stochastica 3 (1979), 1, 47-60.

Estudio del retículo de conceptos L-Fuzzy

A. Burusco Juandeaburre
Dep. de Matemática e Informática.
Universidad Pública de Navarra.
Campus de Arrosadía.
31006 Pamplona. España.
E-mail: Burusco@si.UPNA.es

R. Fuentes González.
Dep. de Matemática e Informática.
Universidad Pública de Navarra.
Campus de Arrosadía.
31006 Pamplona. España.

Resumen

La teoría de conceptos $L - Fuzzy$ que desarrollamos permite establecer clasificaciones de los objetos y atributos pertenecientes a un contexto partiendo de relaciones $L - Fuzzy$. Esta teoría generaliza la teoría formal de conceptos iniciada por R. Wille ([6]).

En este trabajo partimos de las definiciones de conceptos $L - Fuzzy$ que amplían las ya dadas en la teoría formal de conceptos y estudiamos la estructura reticular del conjunto de conceptos $L - Fuzzy$ indicando un método constructivo para calcular dicho retículo. Por último aplicamos este método constructivo a un ejemplo concreto sobre el que se han aplicado otras técnicas.

Palabras clave:

Conceptos $L - Fuzzy$, jerarquías de conceptos, conjuntos $L - Fuzzy$, retículos.

1 Introducción

Tomando como estructura $(L, \leq, ', \top)$, sistema algebraico determinado por:

- Un retículo completo $(L, \leq)$
- Una complementación $'$ en $L(\{1\})$
- Una t-conorma $\top$ en $L(\{1\})$

dábamos en [2] las definiciones básicas que soportan la teoría de conceptos $L - Fuzzy$. Estas definiciones junto con otros resultados que aparecen expuestos en el siguiente apartado nos permitirán estudiar y representar gráficamente la estructura reticular del conjunto de los conceptos $L - Fuzzy$.

2 Operadores derivación en estructuras $(L, \leq, ', \top)$

Sean L^X y L^Y las clases de conjuntos $L - Fuzzy$ asociadas a X e Y respectivamente y sea $R \in L^{X \times Y}$ una relación $L - Fuzzy$.

Dado $A \in L^X$, le asociamos el $L - Fuzzy$ A_1 de L^Y tal que

$$A_1(y) = \inf_{x \in X} (A'(x) \top R(x, y)) \quad (1)$$

Así mismo, a $B \in L^Y$ le asociamos $B_2 \in L^X$

$$B_2(x) = \inf_{y \in Y} (B'(y) \top R(x, y)) \quad (2)$$

Definición 1 A los operadores anteriores A_1 y B_2 los llamaremos operadores derivación ponderados por una negación.

Las definiciones (1) y (2) recuperan, en el caso en que $L = \{0, 1\} = \mathbf{2}$, los conceptos de operadores derivación utilizados en la teoría formal de conceptos de R. Wille ([6]).

Proposición 1 Los operadores derivación definidos en (1) y (2) cumplen:

- $A \leq B \implies A_1 \geq B_1, \forall A, B \in L^X.$
- $C \leq D \implies C_2 \geq D_2, \forall C, D \in L^Y.$
- $\bigvee_{i \in I} (A_i)_1 \leq \left(\bigwedge_{i \in I} A_i \right)_1, \forall A_i \in L^X, i \in I$
- $\bigvee_{i \in I} (B_i)_2 \leq \left(\bigwedge_{i \in I} B_i \right)_2, \forall B_i \in L^Y, i \in I.$

Si $\top$ es una t -conorma semicontinua superiormente en L , entonces

$$\alpha \top \bigwedge_{i \in I} (\beta_i) = \bigwedge_{i \in I} (\alpha \top \beta_i), \quad \forall \alpha, \beta_i \in L$$

y se cumple:

$$\text{iii)} \left(\bigvee_{i \in I} A_i \right)_1 = \bigwedge_{i \in I} (A_i)_1, \quad \forall A_i \in L^X, i \in I.$$

$$\text{iii')} \left(\bigvee_{i \in I} B_i \right)_2 = \bigwedge_{i \in I} (B_i)_2, \quad \forall B_i \in L^Y, i \in I.$$

Demostración i) Sean $A, B \in L^X, R \in L^{X \times Y}$.

Si $A \leq B$, entonces $\tilde{A}(x) \leq \tilde{B}(x) \quad \forall x \in X \Rightarrow$

$$\tilde{A}'(x) \geq \tilde{B}'(x) \quad \forall x \in X \Rightarrow$$

$$\tilde{A}'(x) \top \tilde{R}(x, y) \geq \tilde{B}'(x) \top \tilde{R}(x, y) \quad \forall x \in X, \forall y \in Y.$$

Tomando ínfimos:

$$\inf_{x \in X} (\tilde{A}'(x) \top \tilde{R}(x, y)) \geq \inf_{x \in X} (\tilde{B}'(x) \top \tilde{R}(x, y))$$

$$\Rightarrow \tilde{A}_1(y) \geq \tilde{B}_1(y) \quad \forall y \in Y \Rightarrow \tilde{A}_1 \geq \tilde{B}_1.$$

ii) $\forall A_i \in L^X, i \in I$ se cumple que $\tilde{A}_i \geq \bigwedge_{i \in I} A_i$

Aplicando el operador derivación

$$(\tilde{A}_i)_1 \leq \left(\bigwedge_{i \in I} A_i \right)_1, \quad \forall i \in I;$$

luego $\left(\bigwedge_{i \in I} A_i \right)_1$ será una cota superior de la familia

$\{(\tilde{A}_i)_1, i \in I\}$. Como el supremo es la menor de las cotas superiores:

$$\bigvee_{i \in I} (\tilde{A}_i)_1 \leq \left(\bigwedge_{i \in I} A_i \right)_1.$$

iii) Teniendo en cuenta que L es retículo completo y $\top$ semicontinua superiormente, $\forall y \in Y$

$$\begin{aligned} \left(\bigvee_{i \in I} A_i \right)_1(y) &= \inf_{x \in X} \left\{ \left(\bigvee_{i \in I} A_i \right)'(x) \top \tilde{R}(x, y) \right\} \\ &= \inf_{x \in X} \left\{ \left(\bigwedge_{i \in I} A_i' \right)(x) \top \tilde{R}(x, y) \right\} \\ &= \inf_{x \in X} \left\{ \left(\inf_{i \in I} A_i'(x) \right) \top \tilde{R}(x, y) \right\} \\ &= \inf_{x \in X} \left\{ \inf_{i \in I} \left(A_i'(x) \top \tilde{R}(x, y) \right) \right\} \\ &= \inf_{i \in I} \left\{ \inf_{x \in X} \left(A_i'(x) \top \tilde{R}(x, y) \right) \right\} \\ &= \inf_{i \in I} \left\{ (\tilde{A}_i)_1(y) \right\} = \bigwedge_{i \in I} (\tilde{A}_i)_1(y). \end{aligned}$$

Las demostraciones de i'), ii') y iii') son análogas a las anteriores. ■

Si escribimos A_{12} y B_{21} para representar los conjuntos $L - Fuzzy$ $(\tilde{A}_1)_2$ y $(\tilde{B}_2)_1$ respectivamente, podemos definir los operadores φ y ψ :

$$\varphi : L^X \rightarrow L^X \text{ t.q. } \varphi(\tilde{A}) = A_{12} \quad (3)$$

$$\psi : L^Y \rightarrow L^Y \text{ t.q. } \psi(\tilde{B}) = B_{21} \quad (4)$$

de los que nos serviremos para dar posteriormente nuestra definición de conceptos $L - Fuzzy$.

Definición 2 Llamaremos a los operadores φ y ψ dados en (3) y (4) constructores de conceptos $L - Fuzzy$.

Estos operadores preservan el orden como se demuestra en [2].

3 Retículo de conceptos L -Fuzzy

En este punto estamos en condiciones de introducir las definiciones básicas que soportan este trabajo.

Definición 3 Llamaremos contexto $L - Fuzzy$ a la terna $(X, Y, \tilde{R})$ con X e Y conjuntos de objetos y atributos respectivamente, y $\tilde{R} \in L^{X \times Y}$ una relación $L - Fuzzy$.

Definición 4 Sea $\varphi : L^X \rightarrow L^X$ t.q. $\varphi(\tilde{A}) = A_{12}$ definida anteriormente, y el conjunto $fix(\varphi) = \{\tilde{A} \in L^X / \tilde{A} = \varphi(\tilde{A})\}$.

Si $\tilde{M} \in fix(\varphi)$, al par $(\tilde{M}, \tilde{M}_1)$ lo llamaremos concepto $L - fuzzy$ del contexto $L - Fuzzy$ $(X, Y, \tilde{R})$.

Obsérvese que esta definición recupera como caso particular la dada en la teoría de conceptos en $L = \{0, 1\}$.

En nuestro caso L^X y L^Y son retículos completos, y los operadores constructores de conceptos φ y ψ definidos en (3) y (4) preservan el orden, por tanto, por el teorema de Tarski de puntos fijos ([5]), $\Omega = (\text{fix}(\varphi), \leq, 0_\Omega, 1_\Omega, \vee_\Omega, \wedge_\Omega)$ y $\Sigma = (\text{fix}(\psi), \geq, 0_\Sigma, 1_\Sigma, \vee_\Sigma, \wedge_\Sigma)$ son retículos completos.

Para calcular los elementos máximo y mínimo de dichos retículos, así como el supremo y el ínfimo de una familia recurriremos al trabajo desarrollado por P. Cousot y R. Cousot ([3]), que nos proporciona una versión constructiva del teorema de Tarski ([5]). En ese trabajo se define $\text{luis}(f)(A)$, con f una función monótona, como el límite de una secuencia iterativa superior estacionaria de f que comienza en $\underline{A}$ y $\text{llis}(f)(B)$ como el límite de una secuencia iterativa inferior estacionaria de f que comienza en $\underline{B}$.

A partir de estos límites se calcula.

$$0_\Omega = \text{luis}(\varphi)(0) \quad 0_\Sigma = \text{luis}(\psi)(0) \quad (5)$$

$$1_\Omega = \text{llis}(\varphi)(1) \quad 1_\Sigma = \text{llis}(\psi)(1) \quad (6)$$

$$\forall \{A_i, i \in I\} \in \Omega$$

$$\begin{aligned} \bigvee_{\Omega} A_i &= \text{luis}(\varphi) \left(\bigvee A_i \right) \\ \bigwedge_{\Omega} A_i &= \text{llis}(\varphi) \left(\bigwedge A_i \right) \end{aligned} \quad (7)$$

$$\forall \{B_i, i \in I\} \in \Sigma$$

$$\begin{aligned} \bigvee_{\Sigma} B_i &= \text{luis}(\psi) \left(\bigvee B_i \right) \\ \bigwedge_{\Sigma} B_i &= \text{llis}(\psi) \left(\bigwedge B_i \right) \end{aligned} \quad (8)$$

Sean $\mathcal{L}_1 = \{(\underline{A}, \underline{A}_1) \text{ t.q. } \underline{A} \in \text{fix}(\varphi)\}$ y $\mathcal{L}_2 = \{(\underline{B}_2, \underline{B}) \text{ t.q. } \underline{B} \in \text{fix}(\psi)\}$.

Proposición 2 Los conjuntos $\mathcal{L}_1$ y $\mathcal{L}_2$ coinciden.

Demostración Veamos que $\mathcal{L}_1 \subseteq \mathcal{L}_2$:

$\forall (\underline{A}, \underline{A}_1) \in \mathcal{L}_1$, $\underline{A} \in \text{fix}(\varphi)$ y $\underline{A}_1 \in \text{fix}(\psi)$ ya que $\psi(\underline{A}_1) = (\underline{A}_1)_{21} = (\underline{A})_{12} = (\varphi(\underline{A}))_1 = \underline{A}_1$. Además: $(\underline{A}_1)_2 = \underline{A} = \underline{A}_{12}$, luego $(\underline{A}, \underline{A}_1) \in \mathcal{L}_2$.

Análogamente se demuestra que $\mathcal{L}_2 \subseteq \mathcal{L}_1$ y por tanto $\mathcal{L}_1 = \mathcal{L}_2$. ■

Llamaremos $\mathcal{L}$ al conjunto anterior. Obsérvese que $\mathcal{L} \subseteq \text{fix}(\varphi) \times \text{fix}(\psi)$.

Definición 5 Definimos la relación de orden $\leq^*$ en $\mathcal{L}$ de la siguiente forma:

$$(\underline{A}, \underline{B}) \leq^* (\underline{C}, \underline{D}) \text{ si } \underline{A} \leq \underline{C} \quad (9)$$

$$\forall (\underline{A}, \underline{B}), (\underline{C}, \underline{D}) \in \mathcal{L}$$

Proposición 3 Dados $(\underline{A}, \underline{B}), (\underline{C}, \underline{D}) \in \mathcal{L}$, es equivalente decir que $\underline{A} \leq \underline{C}$ a $\underline{B} \geq \underline{D}$.

Demostración Es evidente por la proposición 1 i), teniendo en cuenta que $\underline{A}_1 = \underline{B}$ y $\underline{C}_1 = \underline{D}$. ■

Nótese que la relación de orden $\leq^*$ viene inducida por las relaciones de orden $\leq$ en Ω y su opuesta $\geq$ en Σ .

Teorema 1 $(\mathcal{L}, \leq^*)$ es un retículo completo.

Demostración. Sea una familia de conceptos

$$\mathcal{F} = \{(\underline{A}_i, (\underline{A}_i)_1), \underline{A}_i \in \Omega, i \in I\} \subseteq \mathcal{L}.$$

Veamos que se puede calcular su supremo en $\mathcal{L}$:

Como $\bigvee_{\Omega} \underline{A}_i$ existe por ser Ω retículo completo,

$$\left(\bigvee_{\Omega} \underline{A}_i, \left(\bigvee_{\Omega} \underline{A}_i \right)_1 \right) \text{ será cota superior ya que } \bigvee_{\Omega} \underline{A}_i \geq \underline{A}_i \quad \forall i \in I \left(y \left(\bigvee_{\Omega} \underline{A}_i \right)_1 \leq (\underline{A}_i)_1 \right).$$

Además se comprueba que es la menor cota superior: si $(\underline{A}, \underline{A}_1) \in \mathcal{L}$, es otra cota superior. $\underline{A} \geq \underline{A}_i, \forall i \in I$, luego $\underline{A} \geq \bigvee_{\Omega} \underline{A}_i$

$$\left(y \text{ respectivamente } \underline{A}_1 \leq \left(\bigvee_{\Omega} \underline{A}_i \right)_1 \right).$$

El elemento mínimo será $\left(0_\Omega, \left(0_\Omega \right)_1 \right)$ donde 0_Ω será el elemento mínimo del retículo Ω . ■

Definición 6 Al retículo $(\mathcal{L}, \leq^*)$ lo llamaremos retículo de conceptos $L = \text{Fuzzy}$ del contexto $L = \text{Fuzzy}(X, Y, R)$.

Proposición 4 Los elementos mínimo y máximo del retículo $(\mathcal{L}, \leq^*)$ son respectivamente:

$$0_{\mathcal{L}} = (0_{\Omega}, 1_{\Sigma}) \text{ y } 1_{\mathcal{L}} = (1_{\Omega}, 0_{\Sigma})$$

Demostración. Puesto que $\forall A \in \Omega$, se cumple que

$$0_{\Omega} \leq A \leq 1_{\Omega}, \text{ entonces}$$

$$(0_{\Omega})_1 \geq A_1 \geq (1_{\Omega})_1, \forall A_1 \in \Sigma, \text{ y como}$$

$$\Sigma = \{A_1 \text{ t.q. } A \in \Omega\}, \text{ se deduce que}$$

$$(0_{\Omega})_1 = 1_{\Sigma} \text{ y } (1_{\Omega})_1 = 0_{\Sigma},$$

luego el menor elemento de $\mathcal{L}$ con el orden $\leq^*$ definido en (9) vendrá dado por la expresión

$$0_{\mathcal{L}} = \left(0_{\Omega}, (0_{\Omega})_1\right) = (0_{\Omega}, 1_{\Sigma}), \text{ y el mayor}$$

$$1_{\mathcal{L}} = \left(1_{\Omega}, (1_{\Omega})_1\right) = (1_{\Omega}, 0_{\Sigma}). \blacksquare$$

Teorema 2 Si utilizamos una *t*-conorma triangular $\top$ semicontinua superiormente para las definiciones de las funciones φ y ψ , dada una familia

$$\mathcal{F} = \{(A_i, (A_i)_1), A_i \in \Omega\}$$

$$= \{((B_i)_2, B_i), B_i \in \Sigma\} \subseteq \mathcal{L}$$

se puede expresar el supremo y el ínfimo de $\mathcal{F}$ de la siguiente forma:

$$\bigvee_{\mathcal{L}} (A_i, (A_i)_1) = \left(\bigvee_{\Omega} A_i, \bigwedge_{\Sigma} (A_i)_1 \right)$$

$$\bigwedge_{\mathcal{L}} ((B_i)_2, B_i) = \left(\bigwedge_{\Omega} (B_i)_2, \bigvee_{\Sigma} B_i \right)$$

donde los supremos e ínfimos en Σ y Ω se calcularán de acuerdo con las expresiones (7) y (8).

Demostración. Para la expresión del supremo basta tener en cuenta la demostración del teorema 1 y la proposición 1; y para el ínfimo se procede de la misma manera. $\blacksquare$

En general, no se puede ampliar al caso $L = Fuzzy$ la construcción del supremo de

una familia de conceptos que nos daba el teorema fundamental de la teoría de conceptos ordinarios de Wille ([6]). Únicamente es cierta la desigualdad:

Proposición 5 Sea una familia de conceptos $\mathcal{F} = \{(A_i, (A_i)_1), A_i \in \Omega, i \in I\}$, se verifica

$$\{((B_i)_2, B_i), B_i \in \Sigma, i \in I\} \subseteq \mathcal{L}$$

$$\bigvee_{\mathcal{L}} (A_i, (A_i)_1) \geq (\varphi(\bigvee A_i), (\varphi(\bigvee A_i))_1)$$

Demostración. Como $\Omega \subseteq L^X$, podemos asegurar que $\bigvee_{\Omega} A_i \geq \bigvee A_i, \forall A_i \in \Omega$. Tomando φ en ambos miembros y teniendo en cuenta que $\bigvee_{\Omega} A_i$ es un elemento de Ω y por lo tanto punto fijo de φ , obtenemos la desigualdad:

$\bigvee_{\Omega} A_i \geq \varphi(\bigvee A_i), \forall A_i \in \Omega$, luego podemos concluir que

$$\bigvee_{\mathcal{L}} (A_i, (A_i)_1) \geq (\varphi(\bigvee A_i), (\varphi(\bigvee A_i))_1)$$

$\forall A_i \in \Omega, i \in I. \blacksquare$

4 Método de construcción del retículo de conceptos

Si L, X e Y son finitos con cardinales k, m y n respectivamente, siempre podemos tomar los k^m posibles subconjuntos $L = Fuzzy$ de L^X y ver si son puntos fijos de φ , en otras palabras, calcular

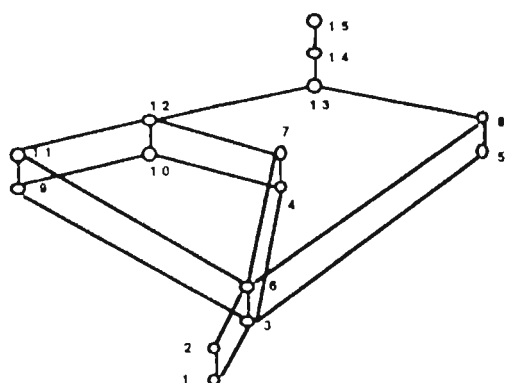
$$\mathcal{M} = \{A \in L^X \text{ t.q. } A = \varphi(A)\}$$

y de esa forma $\forall A \in \mathcal{M}$ tendremos el concepto (A, A_1) y podremos estar seguros de que estamos calculando todo el retículo.

Este método que sólo es aplicable en el caso finito, es el más eficaz si lo que nos interesa es calcular el retículo de conceptos completo.

No obstante, en el caso de que tengamos volúmenes grandes de datos es bastante lento.

Podemos representar gráficamente este retículo de conceptos L - Fuzzy de la siguiente forma:



En general, para interpretar un concepto L -Fuzzy, nos fijaremos en los objetos y atributos cuyos grados de pertenencia destaquen frente al resto. Así, por ejemplo, si nos fijamos en el concepto L - Fuzzy 4 podremos decir que los alumnos 1,10,11 y 13 son a los que más les gustan las ecuaciones diferenciales. Además, si lo comparamos con el concepto L - Fuzzy 10 vemos que de estos alumnos, los que tienen preferencia también por la lógica son el número 1 y el 13.

Estos resultados son deducibles através de una observación minuciosa de la tabla; la ventaja es que nuestra teoría propone un algoritmo para calcularlos.

Por otra parte, si nos fijamos en los conceptos 7,8,10,11 y 12 estamos viendo combinaciones de asignaturas que prefieren los alumnos y su aceptación por parte de estos. Así, observamos en los conceptos 8 y 11 que las parejas de asignaturas que tienen más aceptación son: álgebra-análisis funcional y lógica-álgebra. Por otra parte, del concepto 12 se deduce que la combinación lógica-álgebra-ec.diferenciales es la que más se ajusta a los intereses de los alumnos.

De alguna manera, los conceptos L - Fuzzy nos permiten en este caso clasificar y agrupar

las asignaturas. Esto se visualiza de forma más clara en el gráfico donde se observan las jerarquías que hay entre los conceptos de modo que si, por ejemplo, interesa conocer la opinión del alumno sobre la lógica, sola o junto con otras materias, bastará con analizar la parte izquierda del gráfico que es donde se localizan los conceptos que corresponden a esa materia.

3. Conclusiones

La teoría de conceptos L - Fuzzy nos permite, por una parte, ver de forma más clara y concisa algunos resultados que se podrían visualizar directamente de la tabla, y por otra, extraer conclusiones más elaboradas, siempre através de un proceso algoritmizable.

Referencias

- [1] BURILLO, P., FUENTES, R. Perturbations on L-Fuzzy Sets. *Aportaciones matemáticas en memoria del profesor Onieva*. Universidad de Cantabria. Santander (1991), 43-54.
- [2] BURUSCO, A., FUENTES, R. Aplicación de técnicas L-Fuzzy a la representación formal del conocimiento conceptual. III Congreso Español sobre Tecnologías y Lógica Fuzzy. Santiago de Compostela.(1993).
- [3] COUSOT, P. AND COUSOT, R. Constructive versions of Tarski's fixed point theorems, *Pacific Journal of Mathematics*, 82, (1979), 43-57.
- [4] GANTER, B., STAHL, J. AND WILLE, R. Conceptual measurement and many-valued contexts, *Classification as a Tool of Research*, edited by W.Gaul and M.Schader, North Holland. Amsterdam (1986), 169-176.
- [5] TARSKI, A. A lattice theoretical fixpoint theorem and its applications. *Pacific J.Math.*, 5, (1955), 285-309.
- [6] WILLE, R. Restructuring lattice theory: an approach based on hierarchies of concepts, *Rival I.(ed.).Ordered Sets*.Reidel,Dordrecht-Boston (1982).445-470.

Acerca de la Calculabilidad de las Funciones Difusas.

R. Morales Bueno¹ J.L. Pérez de la Cruz¹ B. Clares² R. Conejo¹

(¹)Dpto. Lenguajes y Ciencias de la
Computación
Facultad de Informática Universidad de
Málaga.
29071-Málaga.
e-mail: morales@ctima.uma.es

(²)Dpto. Lenguajes y Sistemas
Informáticos
Facultad de Ciencias. Universidad de
Granada.
18071-Granada.

Resumen

En el presente trabajo estudiamos diversas clases de funciones difusas (entendidas como subconjuntos difusos de un conjunto producto) en relación con los cálculos realizados por una máquina de Turing difusa y relacionamos estas clases con las anteriormente definidas en la literatura.

Palabras clave: fuzzy computability, Turing W-machine, recursive W-map.

1. Introducción y notación

La calculabilidad borrosa presenta ciertos rasgos propios en relación con el caso nítido. Por ejemplo, el concepto de función recursiva o calculable, definido unívocamente para las funciones nítidas, puede dar lugar a diversos conceptos no equivalentes. En este trabajo estudiamos las consecuencias de considerar la diferencia entre funciones (nítidas) parciales y totales a la hora de definir el concepto de conjunto difuso, y por tanto los de algoritmo y función difusos. Para ello empleamos como herramienta la máquina de Turing difusa o W-máquina, definida en la sección 2. En la sección 3 distinguimos varias formas en las que una W-máquina puede calcular una función difusa: polítropamente, monótopamente y monótonamente, y para cada una de ellas establecemos la clase de funciones calculables.

En lo que sigue consideraremos que $(W, \oplus, \otimes, \leq)$ es un semianillo ordenado finito cuyo supremo es 1 y cuyo ínfimo es 0. Sea una sucesión infinita de elementos de W , $\{w_i\}$. Consideremos la sucesión dada por $s_0=0$, $s_{i+1}=s_i+w_i$. Es evidente que $\{s_i\}$ converge en un número finito de pasos.

Un subconjunto difuso (W-conjunto) de A es una función (nítida) total de A en W . Notaremos los W-conjuntos de X mediante letras griegas minúsculas subindicadas por X . Así, por ejemplo, μ_A, λ_A serán subconjuntos difusos de A . Se define como soporte de μ_A , $sp(\mu_A)$, al subconjunto de A formado por los $u \in A$ tales que $\mu_A(u) \neq 0$. La restricción trivial $\mu_{sp}: sp(A) \rightarrow W$ de μ_A es el W-conjunto de $sp(\mu_A)$ dado para todo $u \in sp(A)$ por $\mu_{sp}(u) = \mu_A(u)$. Sea μ_A un W-conjunto de A , $A \subseteq B$. Se define la extensión trivial de μ_A en B , μ_B , como el W-conjunto de B μ'_B dado por

$$\begin{aligned} \mu'_B(u) &= \mu_A(u) & \text{si } u \in A \\ \mu'_B(u) &= 0 & \text{en otro caso.} \end{aligned}$$

Si A y B son conjuntos (nítidos) tales que $A \subseteq B$, $A \neq B$ y μ_A es un subconjunto difuso de A , entonces ni μ_A ni su restricción trivial μ_{sp} son subconjuntos difusos de B -ya que no son funciones totales de B en W - pero su prolongación trivial μ'_B sí lo es.

Representaremos la restricción trivial de un W-conjunto μ_A mediante la notación

En general, $S(d)$ puede ser vacío, unitario o constar de un número cualquiera finito de descripciones elementales. Sea $\mathcal{E} = \{(d, d') \in \mathcal{D}^2 \mid d' \in S(d)\}$. Ahora podemos definir una W -función parcial $\pi_{\mathcal{E}}$ que describe como transitar en un paso de una descripción elemental a otra.

Definición 3.- Sea $T = (V, S, \psi_E, s_0)$ una W -MT. Definimos una W -función parcial $\pi_{\mathcal{E}}$ de $\mathcal{D}$ en $\mathcal{D}$ como sigue:

$$\pi_{\mathcal{E}}(d, d') = \begin{cases} p(s, c, c', s') & \text{si } d = usc w, d' = us' c' w \\ & \text{si } d = uscc' w, d' = ucs' c' w \\ p(s, c, R, s') & \text{o } d = usc \text{ y } d' = ucs' \# \\ & \text{si } d = uc' scw, d' = us' c' cw \\ p(s, c, L, s') & \text{o } d = scw \text{ y } d' = s' \# cw \end{cases}$$

donde $\#$ es el símbolo blanco, $u, w \in V^*$, $s, s' \in S$ y $c, c' \in V$. La función de transición considerada habitualmente en la literatura ([1], [2], [3] [4]) es precisamente la extensión trivial de $\pi_{\mathcal{E}}$ a $\mathcal{D}^2$.

La situación de una W -MT vendrá descrita por un W -conjunto de descripciones elementales, al que llamaremos descripción.

Definición 4.- Sea $T = (V, S, \psi_E, s_0)$ una W -MT, $\mathcal{D}$ su conjunto de descripciones elementales y $\mathcal{F}$ un subconjunto finito (nítido) $\mathcal{F} \subseteq \mathcal{D}$. Notaremos por $\mu_{\mathcal{F}}$ los W -conjuntos de $\mathcal{F}$ y los denominaremos $\mathcal{F}$ -descripciones.

Diremos que $\mu_{\mathcal{F}}$ es una $\mathcal{F}$ -descripción inicial si $\mathcal{F} = \{s_0 u\}$, $\mu_{sp} = \{1/s_0 u\}$. Diremos que $\mu_{\mathcal{F}}$ es una $\mathcal{F}$ -descripción final si, para cada $d \in \mathcal{F}$, d es una descripción elemental final.

Para cada descripción podemos definir unívocamente un subconjunto (nítido) de V^* y un W -conjunto del mismo, que pueden considerarse ya calculados por la W -MT en esa descripción.

Definición 5.- Sea $\mathcal{F} = \{d_1, \dots, d_r\}$ y $\mu_{\mathcal{F}}$ una $\mathcal{F}$ -descripción. Definimos el conjunto de salida

de $\mathcal{F}$, $O(\mathcal{F})$, como el subconjunto (nítido) de V^* formado por las cadenas asociadas a alguna configuración de $\mathcal{F}$. Definimos el conjunto asociado a $\mu_{\mathcal{F}}$, $\mu'_{O(\mathcal{F})}$, como el W -conjunto de $O(\mathcal{F})$ dado para todo $u \in O(\mathcal{F})$ por

$$\mu'_{O(\mathcal{F})}(u) = \bigcup_{i=1}^r \beta_i$$

donde $\beta_i = \alpha_i$ si u es la cadena asociada a la descripción elemental final d_i y 0 en otro caso.

Veamos ahora cómo de una descripción $\mu_{\mathcal{F}1}$ se transita en un paso a otra descripción $\mu'_{\mathcal{F}2}$. Nótese que en cada paso sólo se pueden alcanzar un número finito de nuevas descripciones elementales, por lo que de un conjunto finito $\mathcal{F}1$ siempre transitamos a otro también finito $\mathcal{F}2$.

Definición 6.- Sea T una W -MT. Sean $\mathcal{F}1$ un conjunto unitario $\mathcal{F}1 = \{d\}$, $\mu_{\mathcal{F}1}$ una $\mathcal{F}1$ -descripción con $\mu_{sp} = \{\alpha/d\}$ y $\mu'_{\mathcal{F}2}$ una $\mathcal{F}2$ -descripción. Diremos que T partiendo de $\mu_{\mathcal{F}1}$ alcanza $\mu'_{\mathcal{F}2}$ en un paso, y lo notaremos $\mu_{\mathcal{F}1} \Rightarrow_T \mu'_{\mathcal{F}2}$, si:

- Si d es una descripción elemental final, entonces $\mu_{\mathcal{F}1} = \mu'_{\mathcal{F}2}$. En otro caso,
- $Im(\pi_{\mathcal{E}}, \mathcal{F}1) = \mathcal{F}2$.
- para todo $d' \in \mathcal{F}2$, $\mu'_{\mathcal{F}2}(d') = \alpha \otimes \pi_{\mathcal{E}}(d, d_i)$.

Definición 7.- Sea T una W -MT. Sean $\mathcal{F}1 = \{d_1, \dots, d_r\}$, $\mu_{\mathcal{F}1}$ una $\mathcal{F}1$ -descripción y $\mu'_{\mathcal{F}2}$ una $\mathcal{F}2$ -descripción. Diremos que T partiendo de $\mu_{\mathcal{F}1}$ alcanza $\mu'_{\mathcal{F}2}$ en un paso, y lo notaremos $\mu_{\mathcal{F}1} \Rightarrow_T \mu'_{\mathcal{F}2}$, si:

- $\mu_{\mathcal{F}1}$ no es una descripción final.
- Para cada $d_i \in \mathcal{F}1$, $\mathcal{F}2_i = Im(\pi_{\mathcal{E}}, \{d_i\})$ y $\{\mu_{\mathcal{F}1}(d_i) / d_i\} \Rightarrow_T \mu^i_{\mathcal{F}2_i}$
- $\mathcal{F}2 = \bigcup_{i=1}^r Im(\pi_{\mathcal{E}}, \mathcal{F}2_i)$
- $\mu'_{\mathcal{F}2} = \bigcup_{i=1}^r \mu^i_{\mathcal{F}2_i}$

Definición 8.- Sea T una W -MT. Consideremos las $\mathcal{F}_i$ -descripciones $\mu_{\mathcal{F}1}$ alcanzadas en i pasos

En general, $S(d)$ puede ser vacío, unitario o constar de un número cualquiera finito de descripciones elementales. Sea $\mathcal{E} = \{(d, d') \in \mathcal{D}^2 \mid d' \in S(d)\}$. Ahora podemos definir una W-función parcial $\pi_{\mathcal{E}}$ que describe como transitar en un paso de una descripción elemental a otra.

Definición 3.- Sea $T = (V, S, \psi_E, s_0)$ una W-MT. Definimos una W-función parcial $\pi_{\mathcal{E}}$ de $\mathcal{D}$ en $\mathcal{D}$ como sigue:

$$\pi_{\mathcal{E}}(d, d') = \begin{cases} p(s, c, c', s') & \text{si } d = uscw, d' = us' c' w \\ & \text{si } d = uscc' w, d' = ucs' c' w \\ p(s, c, R, s') & \text{o } d = usc \text{ y } d' = ucs' \# \\ p(s, c, L, s') & \text{si } d = uc' scw, d' = us' c' cw \\ & \text{o } d = scw \text{ y } d' = s' \# cw \end{cases}$$

donde $\#$ es el símbolo blanco, $u, w \in V^*$, $s, s' \in S$ y $c, c' \in V$. La función de transición considerada habitualmente en la literatura ([1], [2], [3] [4]) es precisamente la extensión trivial de $\pi_{\mathcal{E}}$ a $\mathcal{D}^2$.

La situación de una W-MT vendrá descrita por un W-conjunto de descripciones elementales, al que llamaremos descripción.

Definición 4.- Sea $T = (V, S, \psi_E, s_0)$ una W-MT, $\mathcal{D}$ su conjunto de descripciones elementales y $\mathcal{F}$ un subconjunto finito (nítido) $\mathcal{F} \subseteq \mathcal{D}$. Notaremos por $\mu_{\mathcal{F}}$ los W-conjuntos de $\mathcal{F}$ y los denominaremos $\mathcal{F}$ -descripciones.

Diremos que $\mu_{\mathcal{F}}$ es una $\mathcal{F}$ -descripción inicial si $\mathcal{F} = \{s_0 u\}$, $\mu_{sp} = \{1/s_0 u\}$. Diremos que $\mu_{\mathcal{F}}$ es una $\mathcal{F}$ -descripción final si, para cada $d \in \mathcal{F}$, d es una descripción elemental final.

Para cada descripción podemos definir unívocamente un subconjunto (nítido) de V^* y un W-conjunto del mismo, que pueden considerarse ya calculados por la W-MT en esa descripción.

Definición 5.- Sea $\mathcal{F} = \{d_1, \dots, d_r\}$ y $\mu_{\mathcal{F}}$ una $\mathcal{F}$ -descripción. Definimos el conjunto de salida

de $\mathcal{F}$, $O(\mathcal{F})$, como el subconjunto (nítido) de V^* formado por las cadenas asociadas a alguna configuración de $\mathcal{F}$. Definimos el conjunto asociado a $\mu_{\mathcal{F}}$, $\mu'_{O(\mathcal{F})}$, como el W-conjunto de $O(\mathcal{F})$ dado para todo $u \in O(\mathcal{F})$ por

$$\mu'_{O(\mathcal{F})}(u) = \bigcup_{i=1}^r \beta_i$$

donde $\beta_i = \alpha_i$ si u es la cadena asociada a la descripción elemental final d_i y 0 en otro caso.

Veamos ahora cómo de una descripción $\mu_{\mathcal{F}_1}$ se transita en un paso a otra descripción $\mu'_{\mathcal{F}_2}$. Nótese que en cada paso sólo se pueden alcanzar un número finito de nuevas descripciones elementales, por lo que de un conjunto finito $\mathcal{F}_1$ siempre transitamos a otro también finito $\mathcal{F}_2$.

Definición 6.- Sea T una W-MT. Sean $\mathcal{F}_1$ un conjunto unitario $\mathcal{F}_1 = \{d\}$, $\mu_{\mathcal{F}_1}$ una $\mathcal{F}_1$ -descripción con $\mu_{sp} = \{\alpha/d\}$ y $\mu'_{\mathcal{F}_2}$ una $\mathcal{F}_2$ -descripción. Diremos que T partiendo de $\mu_{\mathcal{F}_1}$ alcanza $\mu'_{\mathcal{F}_2}$ en un paso, y lo notaremos $\mu_{\mathcal{F}_1} \Rightarrow_T \mu'_{\mathcal{F}_2}$, si:

- Si d es una descripción elemental final, entonces $\mu_{\mathcal{F}_1} = \mu'_{\mathcal{F}_2}$. En otro caso,
- $Im(\pi_{\mathcal{E}}, \mathcal{F}_1) = \mathcal{F}_2$.
- para todo $d' \in \mathcal{F}_2$, $\mu'_{\mathcal{F}_2}(d') = \alpha \otimes \pi_{\mathcal{E}}(d, d_i)$.

Definición 7.- Sea T una W-MT. Sean $\mathcal{F}_1 = \{d_1, \dots, d_r\}$, $\mu_{\mathcal{F}_1}$ una $\mathcal{F}_1$ -descripción y $\mu'_{\mathcal{F}_2}$ una $\mathcal{F}_2$ -descripción. Diremos que T partiendo de $\mu_{\mathcal{F}_1}$ alcanza $\mu'_{\mathcal{F}_2}$ en un paso, y lo notaremos $\mu_{\mathcal{F}_1} \Rightarrow_T \mu'_{\mathcal{F}_2}$, si:

- $\mu_{\mathcal{F}_1}$ no es una descripción final.
- Para cada $d_i \in \mathcal{F}_1$, $\mathcal{F}_{2i} = Im(\pi_{\mathcal{E}}, \{d_i\})$ y $\{\mu_{\mathcal{F}_1}(d_i) / d_i\} \Rightarrow_T \mu^i_{\mathcal{F}_{2i}}$
- $\mathcal{F}_2 = \bigcup_{i=1}^r Im(\pi_{\mathcal{E}}, \mathcal{F}_{2i})$
- $\mu'_{\mathcal{F}_2} = \bigcup_{i=1}^r \mu^i_{\mathcal{F}_{2i}}$

Definición 8.- Sea T una W-MT. Consideremos las $\mathcal{F}_i$ -descripciones $\mu_{\mathcal{F}_i}$ alcanzadas en i pasos

desde la descripción inicial $\{I/s_0u\}$. Sean $O(\mathcal{F}_i)$ los correspondientes conjuntos de salida. Nótese que, en general, para cada i se tendrán un $\mathcal{F}_i$ y un $O(\mathcal{F}_i)$ distintos. La sucesión generada por T con entrada u es la sucesión de pares de la forma $(O(\mathcal{F}_i), \mu'_{O(\mathcal{F}_i)})$ donde cada $\mu'_{O(\mathcal{F}_i)}$ es el W -conjunto de $O(\mathcal{F}_i)$ asociado a la descripción $\mu_{\mathcal{F}_i}$.

Nótese que: i) la sucesión $\{O(\mathcal{F}_i)\}$ es monótona no decreciente con respecto a la inclusión; ii) siempre existe su supremo, dado por la unión numerable $\cup O(\mathcal{F}_i)$, que puede ser $\emptyset$; iii) sea un $i \in \mathbb{N}$ y sea un $u \in O(\mathcal{F}_i)$. Consideremos los números $j \geq i$. Entonces, la sucesión de elementos de W dada por $\{\mu'_{O(\mathcal{F}_j)}(u)\}$ es monótona no decreciente y, por lo observado en la sección 1, converge en un número finito de pasos.

Sea $u \in \cup O(\mathcal{F}_i)$. Llamamos primer índice de u , $I(u)$, al $i \in \mathbb{N}$ tal que $u \in O(\mathcal{F}_i)$. y para todo $j < i$, $v \notin O(\mathcal{F}_j)$. Diremos entonces que $W(u) = \mu'_{O(\mathcal{F}_i)}(u)$ es el primer grado de u .

Sean dos W -máquinas T_1, T_2 . En [4] se define su composición, denotada por $T_1 \rightarrow T_2$. Análogamente se define la composición paralela de dos W -máquinas T_1, T_2 , denotada por $[T_1|T_2]$ que es la W -máquina obtenida añadiendo un nuevo estado inicial desde el cual se transita con grado 1 tanto al estado inicial de T_1 como al estado inicial de T_2 . En [1] se prueba (prop. 6.1) que existe una W -máquina universal U que con entrada la codificación de una W -máquina T y la codificación de una entrada u produce cada una de las salidas que produciría T con su grado correspondiente. Esa W -máquina U , convenientemente modificada -llamémosla U' - con entrada la codificación de una W -máquina T , la codificación de una entrada u y un $n \in \mathbb{N}$ obtiene el resultado antes citado pero relativo a la realización de exactamente n pasos. Por tanto, en nuestra notación, U' calcula $(O(\mathcal{F}_n), \mu'_{O(\mathcal{F}_n)})$. A partir de U' es fácil construir otra W -máquina universal U'' que realiza cíclicamente estos dos

procesos: i) generación del siguiente número natural j ; ii) actuación de U'' con tercer argumento el número j .

3. Calculabilidad mediante W -máquinas de Turing

Definición 9.- Sean T una W -MT y μ_S un W -conjunto de un $S \subseteq V^*$. Sea $\{(O_i, \mu^i_{O_i})\}$ la sucesión generada por T con entrada u . Se dice que T calcula polítropamente μ_S con entrada u cuando

$$i) S = \cup O_i$$

y para todo $v \in S$ existe k tal que

$$ii) \mu^k_{O_k}(v) = \mu_S(v).$$

$$iii) \text{ para todo } n > k, \mu^n_{O_n}(v) = \mu_S(v).$$

Sea ϕ_D una W -función de U^* en V^* . Si para todo $u \in U^*$ T calcula polítropamente $\text{Im}(\phi_D, u)$, se dice que T calcula polítropamente ϕ_D y que ϕ_D es polítropamente calculable. Nótese que toda W -MT con entrada u calcula polítropamente algún $\lambda_R(u)$.

Sea $A \subseteq V^*$. La máquina generadora de A con grado α , $G(A, \alpha)$, es la máquina que con cualquier entrada u genera la misma sucesión $\{(O_i, \mu^i_{O_i})\}$ tal que O_i está formado por los i primeros elementos de A según el orden lexicográfico, y $\mu^i_{O_i}$ asigna a todos los elementos de O_i el mismo grado α . Es fácil ver que $G(A, \alpha)$ existe para cualquier conjunto decidable A .

Proposición 1. La W -función total $\phi_{U^* \times V^*}$ es W -calculable (en el sentido de Santos [4]) sii es polítropamente calculable.

Prueba: si ϕ es total y polítropamente calculable, entonces existe T tal que para todo $u \in U^*$ T calcula -con entrada u - explícitamente $\text{Im}(\phi, u) = V^*$, es decir, calcula con entrada u todo $v \in V^*$ en algún número finito de pasos, y realiza este cálculo con su grado $\phi(u, v)$

(posiblemente nulo), luego es W -calculable. Por otra parte, si ϕ es W -calculable, hay una máquina T que calcula (en el sentido de Santos) la función ϕ . Sea $T' = G(V^*, 0)$. Es fácil ver que la máquina $T'' = [T'T']$ calcula polítropamente ϕ .

Proposición 2. *La W -función parcial ϕ_D es polítropamente calculable sii su extensión trivial $\phi'_{U \times V^*}$ es polítropamente calculable.*

Prueba: Sea ϕ_D parcial y polítropamente calculable. Sea T la W -máquina que lo hace, sea $T' = G(V^*, 0)$ y sea como antes $T'' = [T'T']$. T'' calcula la extensión $\phi'_{U \times V^*}$, luego ésta es polítropamente calculable.

En [1] queda demostrado que una W -función total es W -calculable sii todos sus α -cortes son recursivamente enumerables. Por tanto, podemos afirmar que una W -función parcial ϕ_D es polítropamente calculable sii todos los α -cortes de ϕ_D son recursivamente enumerables.

Definición 10.- Sean T una W -MT y μ_S un W -conjunto de un $S \subseteq V^*$. Sea $\{(O_i, \mu^i_{O_i})\}$ la sucesión generada por T con entrada u . Se dice que T calcula monótropeamente μ_S con entrada u cuando

$$i) S = \cup O_i$$

y para todo $v \in S$, existe k tal que:

$$ii) \mu^k_{O_k}(v) = \mu_S(v).$$

$$iii) \text{ para todo } n < k, v \notin O_n.$$

$$iv) \text{ para todo } m > k, \mu^m_{O_m}(v) = \mu_S(v).$$

Sea ϕ_D una W -función de U^* en V^* . Si para todo $u \in U^*$ T calcula monótropeamente $\text{Im}(\phi_D, u)$, se dice que T calcula monótropeamente ϕ_D y que ϕ_D es monótropeamente calculable. Nótese que no toda W -MT calcula monótropeamente algún $\lambda_R(u)$ con entrada u . Obviamente, si T calcula monótropeamente μ_S con entrada u ,

también lo calcula polítropamente, y si ϕ_D es monótropeamente calculable, también es polítropamente calculable.

Proposición 3. *La W -función total $\phi_{U \times V^*}$ es spr sii su restricción trivial ϕ_{sp} es monótropeamente calculable.*

Prueba: (esbozo) sea $\phi_{U \times V^*}$ spr. Entonces existe T que para todo $(u, v) \in \text{sp}(\phi_{U \times V^*})$ calcula el grado. Construyamos una W -máquina Y que con entrada u realiza los siguientes procesos: i) generar el siguiente elemento de V^* en orden lexicográfico ii) realizar en paralelo y con grado 1 estas dos acciones: a) volver al paso i); b) activar la máquina T y, cuando esta pare, parar con grado el escrito en la cinta. Es evidente que la W -máquina Y calcula monótropeamente la restricción trivial ϕ_{sp} .

Sea ahora una ϕ_D monótropeamente calculable y sea T la W -máquina que la calcula. Consideremos la W -máquina universal U'' con entrada la codificación de T y una cadena u . Construyamos Z nítida con entrada $u \# v$ a partir de U'' de la siguiente manera: para cada $j \in \mathbb{N}$, Z comprueba si su entrada v aparece en la salida de U'' . Si es así, Z se para y escribe el grado correspondiente en la cinta. Z calcula el grado para todo par $(u, v) \in D$, luego ϕ_D es spr.

Proposición 4. *Sea ϕ_D una W -función parcial monótropeamente calculable. La condición necesaria y suficiente para que su extensión trivial $\phi'_{U \times V^*}$ sea monótropeamente calculable es que el conjunto $D' = U^* \times V^* - D$ sea decidable.*

Prueba (esbozo): Supongamos ϕ_D monótropeamente calculable. Entonces existe una W -máquina T que con entrada u produce en un número finito de pasos cada uno de los v tales que $(u, v) \in D$. Por otra parte, si D' es decidable, existe una máquina T' que con entrada (u, v) para siempre, escribiendo S sii $(u, v) \in D'$ y N en

otro caso; T' es fácilmente modificable de forma que escriba v con grado 0 sii $(u,v) \in D'$. Combinando esta W -máquina y la anterior, obtenemos una W -máquina que calcula monótonamente $\phi'_{U \times V^*}$. Por otra parte, si $\phi'_{U \times V^*}$ es monótonamente calculable, es obvio que el conjunto D' es decidable.

Definición 11.- Sea μ_S un W -conjunto de un $S \subseteq V^*$. Sea $\{(O_i, \mu_{A_i}^i)\}$ la sucesión generada por T con entrada u . Se dice que T calcula monótonamente μ_S con entrada u cuando se cumplen las siguientes condiciones:

i) T calcula monótonamente μ_S con entrada u .

ii) Si $u < v$, existen k y l tales que $\mu_{O_k}^k(u) = \mu_S(u)$, $\mu_{O_l}^l(v) = \mu_S(v)$ y $n < m$.

Sea ϕ_D una W -función de U^* en V^* . Si existe una W -MT tal que para todo $u \in U^*$ T calcula $Im(\phi_D, u)$ monótonamente, se dice que ϕ_D es monótonamente calculable.

Proposición 5 . La W -función total $\phi_{U \times V^*}$ es recursiva sii es monótonamente calculable.

Prueba: (esbozo) Sea T la W -máquina que calcula monótonamente $\phi_{U \times V^*}$. La máquina nítida que demuestra que la función es recursiva se construye de la siguiente forma: partimos de la entrada $u\#v$. Aplicamos ahora la máquina que tras estas entradas escribe la codificación de T . Consideremos ahora la W -máquina universal U'' con entrada la codificación de T y u . Para cada ciclo de funcionamiento de U'' comparamos sus salidas con v . En el momento en que haya emparejamiento, producimos esa salida con el grado correspondiente.

Recíprocamente, supongamos que $\phi_{U \times V^*}$ es recursiva; por tanto existe T que con entrada $u\#v$ produce su grado. Entonces podemos construir una W -máquina T' que con entrada u escriba junto a u cada una de las cadenas v en orden lexicográfico y active T . Cuando T se pare, T' producirá la salida v con grado el escrito en la cinta y se detendrá.

Proposición 6. La W -función total $\phi_{U \times V^*}$ es monótonamente calculable sii es monótonamente calculable.

Prueba: La implicación en un sentido es trivial. Supongamos por otra parte que $\phi_{U \times V^*}$ es monótonamente calculable. Como sabemos que cada salida v se produce en un número finito de pasos, podemos activar la W -máquina U'' hasta el número de pasos que nos permita alcanzar una v concreta. Las salidas v anteriores en orden lexicográfico pueden haber sido ya obtenidas, en cuyo caso habrán sido descartadas.

4. Conclusiones

Se ha definido una nueva jerarquía de calculabilidad para las funciones difusas, teniendo en cuenta -a diferencia de lo habitual en la literatura- su posible carácter parcial. Cada clase de la jerarquía ha quedado caracterizada en términos de cierto tipo de cálculos realizados por W -máquinas, y puesta en relación con la jerarquía de Gerla [1]. Algunos resultados análogos pueden obtenerse en relación con las funciones finito-calculables [2].

Referencias

- [1] GERLA, G.: Turing L-Machines and Recursive Computability for L-Maps, *Studia Logica* XLVIII 2 (1989), 179-192.
- [2] MORALES BUENO, R. *et al.*: Estudio de una clase de W -funciones calculables. En *Actas del III Congreso Español sobre Tecnologías y Lógica fuzzy*, Santiago de Compostela 1993, 153-160
- [3] MORALES BUENO, R. *et al.*: An extension of the L programming language and its relationship to fuzzy computability. Remitido a *Fuzzy Sets and Systems*.
- [4] SANTOS, E.S.: Fuzzy algorithms. *Information and Control* 17 (1970), 326-339.

Propiedades de transformación en conjuntos fuzzy

M.L. Eraso

Universidad Pública de Navarra
Campus de Arrosadía s/n
31006 Pamplona, España
e-mail: leraso@si.upna.es

M.I. Portilla

Universidad Pública de Navarra
Campus de Arrosadía s/n
31006 Pamplona, España
e-mail: mportilla@si.upna.es

Resumen

Se estudian en los conjuntos fuzzy de $[0,1]^E$ sobre un referencial finito $E \neq \emptyset$ con $|E|=n$, la transformación dada por $\tau_\epsilon(\mu_A(x)) = (1-\epsilon)\mu_A(x) + \chi_A \epsilon$ $0 \leq \epsilon \leq 1$ introducidas por Fadini, Ventre y Cella [4]. Se analizan sus propiedades homomórficas respecto a t-normas y t-conormas y su actuación respecto a distintos complementos fuzzy en $[0,1]$. Por otra parte se estudian los efectos de tales operadores en las formas de entropía en $[0,1]^E$ con las definiciones más usuales en la bibliografía.

Palabras clave: t-normas, t-conormas, complementarios fuzzy, entropía.

1 Introducción.

Fadini, Ventre y Cella definen en [4] una transformación τ_ϵ de funciones de pertenencia de conjuntos fuzzy, de forma que si A es un conjunto fuzzy en $[0,1]^E$ con referencial finito $E \neq \emptyset$, y con función de pertenencia $\mu_A(x)$ para cada $x \in E$, la expresión de la aplicación τ_ϵ viene dada por:

$$\tau_\epsilon: [0,1] \rightarrow [0,1]$$

$$\tau_\epsilon(\mu_A(x)) = (1-\epsilon)\mu_A(x) + \chi_A \epsilon \quad 0 \leq \epsilon \leq 1$$

donde χ_A es la función característica de $\underline{A}$, el conjunto ordinario más próximo al conjunto fuzzy A , es decir: $\chi_A = 0$ si $\mu_A(x) \leq 0,5$, $\chi_A = 1$ si $\mu_A(x) > 0,5$.

De esta forma la aplicación τ_ϵ define la transformación de un conjunto fuzzy en otro conjunto fuzzy según la definición siguiente:

Definición 1 Dado un conjunto fuzzy A de $[0,1]^E$ se define su transformado por τ_ϵ como

$$\tau_\epsilon(A) = \{(x, \tau_\epsilon(\mu_A(x))), x \in E\} \in [0,1]^E$$

Se pueden observar dos casos particulares en esta definición que conviene resaltar. El primero, $\epsilon=0$, para el cual $\tau_0(A) = A$, corresponde a la transformación

identidad y se elimina del estudio en lo sucesivo. El segundo es $\epsilon=1$, en el que $\tau_1(A) = \underline{A}$, es decir, τ_1 transforma cualquier difuso en un conjunto ordinario (el más próximo a él), mientras que $\tau_\epsilon(A)$ $0 < \epsilon < 1$ es siempre un difuso si A lo es. Por tanto el caso $\epsilon=1$ se estudiará como caso singular en algunos resultados.

2 Propiedades de la transformación τ_ϵ

Estos operadores τ_ϵ cumplen, entre otras, las propiedades siguientes [4]:

Proposición 1 La transformación τ_ϵ para $0 < \epsilon \leq 1$

verifica $\forall A \in [0,1]^E$

$$\tau_\epsilon(\mu_A(x)) \leq \mu_A(x) \leq 0,5 \quad \tau_\epsilon(\mu_A(x)) \geq \mu_A(x) > 0,5$$

con igualdad únicamente para $\mu_A(x) = 0$ ó 1 .

Proposición 2 La transformación τ_ϵ para $0 < \epsilon < 1$

verifica $\forall A \in [0,1]^E$

$$\{\tau_\epsilon(\mu_A(x)) \mid x \in E\} \cap (0,5-\epsilon/2, 0,5+\epsilon/2] = \emptyset$$

Otra propiedad importante de la función τ_ϵ como transformación de conjuntos fuzzy de $[0,1]^E$, es la de conservación de la unión y la intersección habituales de conjuntos fuzzy [4]:

Proposición 3 La transformación τ_ϵ con $0 < \epsilon \leq 1$

verifica $\forall A, B \in [0,1]^E$

$$\tau_\epsilon(A \cup B) = \tau_\epsilon(A) \cup \tau_\epsilon(B) \quad \tau_\epsilon(A \cap B) = \tau_\epsilon(A) \cap \tau_\epsilon(B)$$

con $\mu_{A \cup B}(x) = \max\{\mu_A(x), \mu_B(x)\}$

$$\mu_{A \cap B}(x) = \min\{\mu_A(x), \mu_B(x)\}$$

Luego τ_ϵ $0 < \epsilon \leq 1$ es homomorfismo en $\{[0,1]^E, \cup, \cap\}$.

3 Conservación de t-normas y t-conormas de conjuntos fuzzy por τ_ϵ

De la misma forma se puede estudiar si la generalización de uniones e intersecciones por t-conormas y t-normas se conservan al realizar la transformación por τ_ϵ .

Es conocido (ver [5], [8], [11]) que una aplicación $T: [0,1] \times [0,1] \rightarrow [0,1]$ verificando

$$T1.- T(a,1)=a \quad \forall a \in [0,1]$$

$$T2.- T(a,b)=T(b,a) \quad \forall a,b \in [0,1]$$

$$T3.- T(a,b) \leq T(c,d) \quad \text{si } a \leq c \quad b \leq d \quad \forall a,b,c,d \in [0,1]$$

$$T4.- T(a,T(b,c))=T(T(a,b),c) \quad \forall a,b,c \in [0,1]$$

es una t-norma y una aplicación

$$S: [0,1] \times [0,1] \rightarrow [0,1] \quad \text{verificando}$$

$$S1.- S(a,0)=a \quad \forall a \in [0,1]$$

$$S2.- S(a,b)=S(b,a) \quad \forall a,b \in [0,1]$$

$$S3.- S(a,b) \leq S(c,d) \quad \text{si } a \leq c \quad b \leq d \quad \forall a,b,c,d \in [0,1]$$

$$S4.- S(a,S(b,c))=S(S(a,b),c) \quad \forall a,b,c \in [0,1]$$

es una t-conorma.

A partir de ellas, para $A, B \in [0,1]^E$, se pueden construir los nuevos difusos

$$T(A,B) = \{(x, T(\mu_A(x), \mu_B(x))), x \in E\} \in [0,1]^E$$

$$S(A,B) = \{(x, S(\mu_A(x), \mu_B(x))), x \in E\} \in [0,1]^E$$

De esta forma para estudiar si τ_ϵ conserva las t-normas y t-conormas anteriores para $0 < \epsilon \leq 1$, se debe analizar si se verifica:

$$\tau_\epsilon(T(A,B)) = T(\tau_\epsilon(A), \tau_\epsilon(B)) \quad \forall A, B \in [0,1]^E \quad (1)$$

$$\tau_\epsilon(S(A,B)) = S(\tau_\epsilon(A), \tau_\epsilon(B)) \quad \forall A, B \in [0,1]^E \quad (2)$$

para cualquier par (T,S) de t-norma T y t-conorma S o bien si solo lo satisface el par $(\min, \max)$.

Debemos notar que la mayoría de las t-normas y t-conormas que recogen Gupta y Qi [5], Klir y Folger [8], Zimmermann [11] no verifican (1) y (2). Veamos algunos ejemplos en los que denotaremos los elementos de E por $\{x_1, x_2, \dots, x_n\}$:

• **Producto algebraico** $T(a,b)=ab$

Tomando los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,6$ y $\mu_A(x_i)=\mu_B(x_i)=0$ para $i=2, \dots, n$ vemos que no se cumple (1) ya que $\forall \epsilon \quad 0 < \epsilon \leq 1$

$$\tau_\epsilon(T(0,6,0,6))=\tau_\epsilon(0,36)=(1-\epsilon)0,36 \neq$$

$$T(\tau_\epsilon(0,6), \tau_\epsilon(0,6))=((1-\epsilon)0,6+\epsilon)^2$$

• **Suma algebraica** $S(a,b)=a+b-ab$

Tomando los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,5$ y $\mu_A(x_i)=\mu_B(x_i)=0$ para $i=2, \dots, n$ vemos que no se cumple (2) ya que $\forall \epsilon \quad 0 < \epsilon \leq 1$

$$\tau_\epsilon(S(0,5,0,5))=\tau_\epsilon(0,75)=(1-\epsilon)0,75+\epsilon \neq$$

$$S(\tau_\epsilon(0,5), \tau_\epsilon(0,5))=S((1-\epsilon)0,5, (1-\epsilon)0,5)=(1-\epsilon) - \frac{(1-\epsilon)^2}{4}$$

• **Diferencia acotada** $T(a,b)=\max(a+b-1, 0)$

Tomando los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,6$ y $\mu_A(x_i)=\mu_B(x_i)=0$ para $i=2, \dots, n$ vemos que no se cumple (1) ya que $\forall \epsilon \quad 0 < \epsilon \leq 1$

$$\tau_\epsilon(T(0,6,0,6))=\tau_\epsilon(0,2)=(1-\epsilon)0,2 \neq$$

$$T(\tau_\epsilon(0,6), \tau_\epsilon(0,6))=T((1-\epsilon)0,6+\epsilon, (1-\epsilon)0,6+\epsilon)=$$

$$\max\{(1-\epsilon)1, 2+2\epsilon-1, 0\}=0,2+0,8\epsilon$$

• **Suma acotada** $S(a,b)=\min(a+b, 1)$

Tomando los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,3$ y $\mu_A(x_i)=\mu_B(x_i)=0$ para $i=2, \dots, n$ vemos que no se cumple (2) ya que $\forall \epsilon \quad 0 < \epsilon \leq 1$

$$\tau_\epsilon(S(0,3,0,3))=\tau_\epsilon(0,6)=(1-\epsilon)0,6+\epsilon \neq$$

$$S(\tau_\epsilon(0,3), \tau_\epsilon(0,3))=S((1-\epsilon)0,3, (1-\epsilon)0,3)=(1-\epsilon)0,6$$

• **Producto Hamacher** $T(a,b)=\frac{ab}{a+b-ab}$

Tomando los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,6$ y $\mu_A(x_i)=\mu_B(x_i)=0$ para $i=2, \dots, n$ vemos que no se cumple (1) ya que $\forall \epsilon \quad 0 < \epsilon \leq 1$

$$\tau_\epsilon(T(0,6,0,6))=\tau_\epsilon(3/7)=\frac{(1-\epsilon)3}{7} \neq T(\tau_\epsilon(0,6), \tau_\epsilon(0,6))=$$

$$=T((1-\epsilon)0,6+\epsilon, (1-\epsilon)0,6+\epsilon)=\frac{(1-\epsilon)0,6+\epsilon}{1,4-0,4\epsilon}$$

• **Suma Hamacher** $S(a,b)=\frac{a+b-2ab}{1-ab}$

Tomando los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,5$ y $\mu_A(x_i)=\mu_B(x_i)=0$ para $i=2, \dots, n$ vemos que no se cumple (1) ya que $\forall \epsilon \quad 0 < \epsilon \leq 1$

$$\tau_\epsilon(S(0,5,0,5))=\tau_\epsilon(2/3)=\frac{2(1-\epsilon)}{3} + \epsilon \neq$$

$$S(\tau_\epsilon(0,5), \tau_\epsilon(0,5))=\frac{(1-\epsilon)-(1-\epsilon)^2 0,5}{1-(1-\epsilon)^2 0,25}$$

• **t-norma Hamacher** $T(a,b)=\frac{\lambda ab}{1-(1-\lambda)(a+b-ab)} \quad \lambda > 0$

Tomando los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,6$ y $\mu_A(x_i)=\mu_B(x_i)=0$ para $i=2, \dots, n$ vemos que no se cumple (1) ya que $\forall \epsilon \quad 0 < \epsilon \leq 1$ y $\forall \lambda > 0$

$$\tau_\epsilon(T(0,6,0,6))=\tau_\epsilon(9\lambda/(4+21\lambda))=(1-\epsilon)\frac{9\lambda}{4+21\lambda} \neq$$

$$T(\tau_\epsilon(0,6), \tau_\epsilon(0,6))=\frac{\lambda(0,6+0,4\epsilon)^2}{1-(1-\lambda)(0,6+0,4\epsilon)(1,4-0,4\epsilon)}$$

Como caso particular para $\lambda=0,5$ el Producto Einstein

$$T(a,b)=\frac{ab}{2-(a+b-ab)} \quad \text{tampoco verifica (1).}$$

• **t-conorma Hamacher** $S(a,b)=\frac{\lambda(a+b)+ab(1-2\lambda)}{\lambda+ab(1-\lambda)} \quad \lambda > 0$

Tomando los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,5$ y $\mu_A(x_i)=\mu_B(x_i)=0$ para $i=2, \dots, n$ vemos que no se cumple (2) ya que $\forall \epsilon \quad 0 < \epsilon \leq 1$ y $\forall \lambda > 0$

$$\tau_\epsilon(S(0,5,0,5))=\tau_\epsilon((1+2\lambda)/(1+3\lambda))=(1-\epsilon)\frac{1+2\lambda}{1+3\lambda} + \epsilon \neq$$

$$S(\tau_\epsilon(0,5), \tau_\epsilon(0,5))=\frac{\lambda(1-\epsilon)+(1-\epsilon)^2 0,25(1-2\lambda)}{\lambda+(1-\lambda)(1-\epsilon)^2 0,25}$$

Como caso particular para $\lambda=0,5$ la Suma Einstein

$$S(a,b)=\frac{a+b}{1+ab} \quad \text{tampoco verifica (2).}$$

• **t-norma Dubois** $T(a,b)=\frac{ab}{\max(a,b,\lambda)} \quad \lambda \in [0,1]$

Si $\lambda=0$, $T(a,b)=\min(a,b)$ cumple (1) por proposición 3.

Para cada $\lambda \in (0,1]$ si tomamos los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,4\lambda$ y $\mu_A(x_i)=\mu_B(x_i)=0$ para $i=2,\dots,n$ no se cumple (1) ya que $\forall \varepsilon \ 0<\varepsilon<1$
 $\tau_\varepsilon(T(0,4\lambda,0,4\lambda))=\tau_\varepsilon(0,16\lambda)=(1-\varepsilon)0,16\lambda \neq$

$$T(\tau_\varepsilon(0,4\lambda),\tau_\varepsilon(0,4\lambda))=\frac{((1-\varepsilon)0,4\lambda)^2}{\max((1-\varepsilon)0,4\lambda,\lambda)}=(1-\varepsilon)^2 0,16\lambda$$

• **t-conorma Dubois** $S(a,b)=1-\frac{(1-a)(1-b)}{\max(1-a,1-b,\lambda)}$
 $\lambda \in [0,1]$

Si $\lambda=0$, $S(a,b)=\max(a,b)$ cumple (2) por proposición 3.
 Para cada $\lambda \in (0,1]$ si tomamos los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=1-0,4\lambda$ y $\mu_A(x_i)=\mu_B(x_i)=0$ para $i=2,\dots,n$ no se cumple (2) ya que $\forall \varepsilon \ 0<\varepsilon<1$
 $\tau_\varepsilon(S(1-0,4\lambda,1-0,4\lambda))=\tau_\varepsilon(1-0,16\lambda)=1-(1-\varepsilon)0,16\lambda \neq$

$$S(\tau_\varepsilon(1-0,4\lambda),\tau_\varepsilon(1-0,4\lambda))=1-\frac{((1-\varepsilon)0,4\lambda)^2}{\max((1-\varepsilon)(1-0,4\lambda)+\varepsilon,\lambda)}$$

$$=1-(1-\varepsilon)^2 0,16\lambda$$

Proposición 4 Si $0<\varepsilon<1$, entonces

$$\tau_\varepsilon(T_\omega(A,B))=T_\omega(\tau_\varepsilon(A),\tau_\varepsilon(B)) \quad \forall A,B \in [0,1]^E$$

$$\tau_\varepsilon(S_\omega(A,B))=S_\omega(\tau_\varepsilon(A),\tau_\varepsilon(B)) \quad \forall A,B \in [0,1]^E$$

con T_ω producto drástico y S_ω suma drástica.

Demostración. La t-norma anterior viene dada por:

$$\text{Producto drástico} \quad T_\omega(a,b)=\begin{cases} a & \text{si } b=1 \\ b & \text{si } a=1 \\ 0 & \text{en otro caso} \end{cases}$$

y verifica la propiedad (1) para $0<\varepsilon<1$ ya que:

$$\begin{aligned} \text{i) Para } x \in E \text{ tal que } \mu_B(x)=1 \\ \tau_\varepsilon(T_\omega(\mu_A(x),\mu_B(x)))&=\tau_\varepsilon(\mu_A(x))=T_\omega(\tau_\varepsilon(\mu_A(x)),1)= \\ &=T_\omega(\tau_\varepsilon(\mu_A(x)),\tau_\varepsilon(\mu_B(x))) \end{aligned}$$

ii) Análogamente para $x \in E$ tal que $\mu_A(x)=1$.

iii) Para $x \in E$ tal que $\mu_A(x) \neq 1$ y $\mu_B(x) \neq 1$

$$\begin{aligned} \tau_\varepsilon(T_\omega(\mu_A(x),\mu_B(x)))&=\tau_\varepsilon(0)=0= \\ &=T_\omega(\tau_\varepsilon(\mu_A(x)),\tau_\varepsilon(\mu_B(x)))=0 \end{aligned}$$

puesto que $\tau_\varepsilon(a)=1$ únicamente para $a=1$.

La t-conorma anterior viene dada por:

$$\text{Suma drástica} \quad S_\omega(a,b)=\begin{cases} a & \text{si } b=0 \\ b & \text{si } a=0 \\ 1 & \text{en otro caso} \end{cases}$$

y verifica la propiedad (2) para $0<\varepsilon<1$ ya que:

$$\begin{aligned} \text{i) Para } x \in E \text{ tal que } \mu_B(x)=0 \\ \tau_\varepsilon(S_\omega(\mu_A(x),\mu_B(x)))&=\tau_\varepsilon(\mu_A(x))=S_\omega(\tau_\varepsilon(\mu_A(x)),0)= \\ &=S_\omega(\tau_\varepsilon(\mu_A(x)),\tau_\varepsilon(\mu_B(x))) \end{aligned}$$

ii) Análogamente para $x \in E$ tal que $\mu_A(x)=0$.

iii) Para $x \in E$ tal que $\mu_A(x) \neq 0$ y $\mu_B(x) \neq 0$

$$\begin{aligned} \tau_\varepsilon(S_\omega(\mu_A(x),\mu_B(x)))&=\tau_\varepsilon(1)=1= \\ &=S_\omega(\tau_\varepsilon(\mu_A(x)),\tau_\varepsilon(\mu_B(x)))=1 \end{aligned}$$

puesto que $\tau_\varepsilon(a)=0$ únicamente para $a=0$. ♦

Notemos que si $\varepsilon=1$ la proposición anterior no es cierta ya que si denotamos por $\{x_1, x_2, \dots, x_n\}$ los elementos de E y tomamos los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,6$ $\mu_A(x_i)=\mu_B(x_i)=0$ $i=2,\dots,n$ no se verifica la expresión (1) para $\varepsilon=1$ ya que

$$\tau_1(T_\omega(0,6,0,6))=\tau_1(0)=0 \neq$$

$$T_\omega(\tau_1(0,6),\tau_1(0,6))=T_\omega(1,1)=1$$

y tomando los difusos A y B tales que $\mu_A(x_1)=\mu_B(x_1)=0,5$ $\mu_A(x_i)=\mu_B(x_i)=0$ $i=2,\dots,n$ no se verifica la expresión (2) para $\varepsilon=1$ ya que

$$\tau_1(S_\omega(0,5,0,5))=\tau_1(1)=1 \neq$$

$$S_\omega(\tau_1(0,5),\tau_1(0,5))=S_\omega(0,0)=0$$

En esta línea disponemos de la siguiente información:

Proposición 5 Si $\varepsilon=1$ una t-norma T verifica (1) si y solo si

$$T(\mu_A(x),\mu_B(x))>0,5 \Leftrightarrow \mu_A(x)>0,5 \text{ y } \mu_B(x)>0,5$$

$$\forall x \in E \quad \forall A,B \in [0,1]^E$$

y una t-conorma S verifica (2) si y solo si

$$S(\mu_A(x),\mu_B(x)) \leq 0,5 \Leftrightarrow \mu_A(x) \leq 0,5 \text{ y } \mu_B(x) \leq 0,5$$

$$\forall x \in E \quad \forall A,B \in [0,1]^E$$

Demostración. Si T verifica (1) para $\varepsilon=1$

$$T(\tau_1(A),\tau_1(B))=\tau_1(T(A,B)) \quad \forall A,B \in [0,1]^E$$

$$\Leftrightarrow T(\chi_A, \chi_B)=\chi_{T(A,B)} \quad \forall A,B \in [0,1]^E$$

$$\Leftrightarrow \begin{cases} T(1,1)=1 & \text{si } \mu_A(x)>0,5 \text{ y } \mu_B(x)>0,5 \\ T(0,0)=T(1,0)=T(0,1)=0 & \text{en otro caso} \end{cases}$$

$$= \begin{cases} 1 & \text{si } T(\mu_A(x),\mu_B(x))>0,5 \\ 0 & \text{en otro caso} \end{cases}$$

$$\forall A,B \in [0,1]^E \quad \forall x \in E$$

Si S verifica (2) para $\varepsilon=1$

$$S(\tau_1(A),\tau_1(B))=\tau_1(S(A,B)) \quad \forall A,B \in [0,1]^E$$

$$\Leftrightarrow S(\chi_A, \chi_B)=\chi_{S(A,B)} \quad \forall A,B \in [0,1]^E$$

$$\Leftrightarrow \begin{cases} S(0,0)=0 & \text{si } \mu_A(x) \leq 0,5 \text{ y } \mu_B(x) \leq 0,5 \\ S(0,1)=S(1,0)=S(1,1)=1 & \text{en otro caso} \end{cases}$$

$$= \begin{cases} 0 & \text{si } S(\mu_A(x),\mu_B(x)) \leq 0,5 \\ 1 & \text{en otro caso} \end{cases}$$

$$\forall A,B \in [0,1]^E \quad \forall x \in E \quad \blacklozenge$$

4 Conservación de complementos de conjuntos fuzzy por τ_ε

Para el caso del complementario de un conjunto fuzzy también se va a estudiar la relación entre la transformación por τ_ε del complementario h de un difuso A de $[0,1]^E$ y el complementario de su transformación por τ_ε y comprobar si verifica:

$$\tau_\varepsilon(h(A))=h(\tau_\varepsilon(A)) \quad \forall A \in [0,1]^E \quad (3)$$

Para generalizar emplearemos aplicaciones negación (ver Klir y Folger [8]) $h : [0,1] \rightarrow [0,1]$ no crecientes tales que $h(0)=1$ y $h(1)=0$, de forma que para un difuso $A \in [0,1]^E$ escribimos su negación o complemento como $h(A) = \{(x, h(\mu_A(x))), x \in E\} \in [0,1]^E$

El caso más habitual es la pseudo-complementación en $[0,1]$, $h(a)=1-a$, en la que se prueba lo siguiente:

Proposición 6 Si $h(a)=1-a$, entonces para $0 < \varepsilon \leq 1$
 $h(\tau_\varepsilon(\mu_A(x))) = \tau_\varepsilon(h(\mu_A(x))) \Leftrightarrow \mu_A(x) \neq 0,5 \quad \forall x \in E$

Demostración. Para $0 < \varepsilon \leq 1$

$$h(\tau_\varepsilon(\mu_A(x))) = \tau_\varepsilon(h(\mu_A(x))) \Leftrightarrow$$

$$1 - \tau_\varepsilon(\mu_A(x)) = \tau_\varepsilon(1 - \mu_A(x)) \Leftrightarrow$$

$$1 - [(1-\varepsilon)\mu_A(x) + \chi_{\Delta} \varepsilon] = (1-\varepsilon)(1 - \mu_A(x)) + \chi_{h(\Delta)} \varepsilon \Leftrightarrow$$

$$\chi_{\Delta} + \chi_{h(\Delta)} = 1 \Leftrightarrow \mu_A(x) \neq 0,5 \quad \forall x \in E \quad \blacklozenge$$

En el caso más general de complementación fuzzy llegamos a la conclusión siguiente:

Proposición 7 Sea h una negación tal que $h(a) \neq 0$ para algún $a \in (0,1)$. Entonces si h verifica (3) $\forall \varepsilon$ $0 < \varepsilon \leq 1$ se tiene $h(0,5) > 0,5$.

Demostración. Para $\varepsilon=1$.

$$h(\tau_1(A)) = \tau_1(h(A)) \quad \forall A \in [0,1]^E \Leftrightarrow$$

$$\begin{aligned} \forall A \in [0,1]^E \quad \forall x \in E \quad & \begin{cases} h(0)=1 & \text{si } \mu_A(x) \leq 0,5 \\ h(1)=0 & \text{si } \mu_A(x) > 0,5 \end{cases} \\ & = \begin{cases} 0 & \text{si } h(\mu_A(x)) \leq 0,5 \\ 1 & \text{si } h(\mu_A(x)) > 0,5 \end{cases} \end{aligned}$$

$$\Rightarrow h(0,5) > 0,5$$

Para $0 < \varepsilon < 1$. Tenemos

$$h(\tau_\varepsilon(A)) = \tau_\varepsilon(h(A)) \quad \forall A \in [0,1]^E \quad \forall \varepsilon \quad 0 < \varepsilon < 1$$

Denotemos los elementos de $E = \{x_1, x_2, \dots, x_n\}$ y sea A tal que $\mu_A(x_1) = 0,5$ $\mu_A(x_i) = 0$ $i=2, \dots, n$

entonces se verificará (3) $\forall \varepsilon \quad 0 < \varepsilon < 1$ si y solo si

$$h(\tau_\varepsilon(0,5)) = \tau_\varepsilon(h(0,5)) \quad \forall \varepsilon \quad 0 < \varepsilon < 1$$

Si suponemos $h(0,5) = 0,5$ tenemos

$$h((1-\varepsilon)/2) = (1-\varepsilon)/2 \quad \forall \varepsilon \quad 0 < \varepsilon < 1$$

$$\Leftrightarrow h(a) = a \quad \forall a \in (0, 0,5)$$

que es contradictorio con h decreciente, con lo cual $h(0,5) \neq 0,5$.

Si suponemos $h(0,5) < 0,5$ tenemos

$$h((1-\varepsilon)/2) < (1-\varepsilon)/2 \quad \forall \varepsilon \quad 0 < \varepsilon < 1$$

$$\Leftrightarrow h(a) < a \quad \forall a \in (0, 0,5)$$

Por ser h decreciente podemos escribir

$$0 \leq h(c) \leq h(b) \leq h(a) < a \leq b \leq 0,5 \leq c \leq 1$$

que lleva de nuevo a contradicción puesto que la función negación es tal que $h(a) \neq 0$ para algún $a \in (0,1)$. $\blacklozenge$

Si imponemos que la función h sea además continua (condición que cumplen las complementaciones más usuales) tenemos:

Proposición 8 Para cualquier función negación continua h se cumple que $\forall \varepsilon$ con $0 < \varepsilon < 1$

$$\exists A \in [0,1]^E \text{ tal que } \tau_\varepsilon(h(A)) \neq h(\tau_\varepsilon(A))$$

Demostración. Por reducción al absurdo supongamos que se verifica (3) $\forall \varepsilon \quad 0 < \varepsilon < 1$

Entonces por la proposición 7 $h(0,5) > 0,5$

Por ser h continua y además decreciente

$$\exists a \in (0,5, h^{-1}(0,5)) \text{ tal que } h(a) > 0,5$$

Sea A tal que $\mu_A(x_1) = a$ $\mu_A(x_i) = 0$ $i=2, \dots, n$

entonces se verifica (3) $\forall \varepsilon \quad 0 < \varepsilon < 1$ si y solo si

$$h(\tau_\varepsilon(a)) = \tau_\varepsilon(h(a)) \quad \forall \varepsilon \quad 0 < \varepsilon < 1$$

$$\Leftrightarrow h((1-\varepsilon)a + \varepsilon) = (1-\varepsilon)h(a) + \varepsilon \quad \forall \varepsilon \quad 0 < \varepsilon < 1$$

Por ser h decreciente y $a < (1-\varepsilon)a + \varepsilon \quad \forall a \in [0,1]$ y

entonces tenemos por lo anterior $\forall \varepsilon \quad 0 < \varepsilon < 1$

$$h(a) \geq h((1-\varepsilon)a + \varepsilon) = (1-\varepsilon)h(a) + \varepsilon > h(a)$$

que es imposible. $\blacklozenge$

Proposición 9 Sea h una negación discontinua tal que $h(a) \neq 0$ para algún $a \in (0,1)$. Entonces si h verifica (3) $\forall \varepsilon$ con $0 < \varepsilon \leq 1$ se tiene

$$I) \quad \forall a \in (0,5, 1] \quad h(a) \leq 0,5$$

II) h tiene al menos una discontinuidad en $0,5$.

Demostración.

I) Por reducción al absurdo si suponemos que

$$\exists a \in (0,5, 1] \text{ tal que } h(a) > 0,5$$

y procedemos como en la demostración de la proposición anterior llegamos a

$$h((1-\varepsilon)a + \varepsilon) = (1-\varepsilon)h(a) + \varepsilon \quad \forall \varepsilon \quad 0 < \varepsilon < 1$$

que es imposible.

Para $\varepsilon=1$ si se verifica (3) tenemos $\forall A \in [0,1]^E$

$$h(\tau_1(A)) = \begin{cases} h(0)=1 & \text{si } \mu_A(x) \leq 0,5 \\ h(1)=0 & \text{si } \mu_A(x) > 0,5 \end{cases}$$

$$= \tau_1(h(A)) = \begin{cases} 0 & \text{si } h(\mu_A(x)) \leq 0,5 \\ 1 & \text{si } h(\mu_A(x)) > 0,5 \end{cases}$$

y entonces $h(a) \leq 0,5 \Leftrightarrow a \in (0,5, 1]$

II) Por proposición 7 $h(0,5) > 0,5$ y por I) $h(a) \leq 0,5$

$\forall a \in (0,5, 1]$, luego $h(0,5^+) \leq 0,5$ y $h(0,5) > 0,5$ como queríamos demostrar. $\blacklozenge$

Nota: Para $\varepsilon=1$ es cierto el recíproco del apartado I) de la proposición 9 (ver demostración):

h verifica (3) $\varepsilon=1$ si y solo si $h(a) \leq 0,5 \Leftrightarrow a \in (0,5, 1]$

Según lo expuesto las únicas negaciones que pueden verificar (3) $\forall \varepsilon \quad 0 < \varepsilon < 1$ se encuentran entre las discontinuas al menos en $0,5$ que verifican $h(0,5) > 0,5$ y $h(a) \leq 0,5 \quad \forall a \in (0,5, 1]$. Por ejemplo las funciones

$$h(x) = \begin{cases} 1-\alpha a & \text{si } a \leq 0,5 & \alpha \in [0,1[\\ \beta(1-a) & \text{si } a > 0,5 & \beta \in [0,1] \end{cases}$$

cumplen dichas condiciones y verifican (3) $\forall \epsilon \ 0 < \epsilon \leq 1$.

Sin embargo dichas condiciones no son suficientes para $\epsilon \ 0 < \epsilon < 1$ ya que basta considerar las negaciones

$$h(x) = \begin{cases} 1-\alpha a^2 & \text{si } a \leq 0,5 & \alpha \in [0,2[\\ \beta(1-a^2) & \text{si } a > 0,5 & \beta \in [0,2/3] \end{cases}$$

para comprobar fácilmente que cumple las condiciones anteriores y no verifica (3) en general ya que si $\mu_A(x) = a \leq 0,5$ tenemos

$$\tau_\epsilon(h(a)) = \tau_\epsilon(1-\alpha a^2) = (1-\epsilon)(1-\alpha a^2) + \epsilon \neq$$

$$h(\tau_\epsilon(a)) = 1-\alpha \tau_\epsilon(a)^2 = 1-\alpha(1-\epsilon)^2 a^2 \quad \forall \epsilon \ 0 < \epsilon < 1$$

y si $\mu_A(x) = b > 0,5$ tenemos

$$\tau_\epsilon(h(b)) = \tau_\epsilon(\beta(1-b^2)) = (1-\epsilon)\beta(1-b^2) \neq$$

$$h(\tau_\epsilon(b)) = \beta(1-\tau_\epsilon(b)^2) = \beta(1-((1-\epsilon)b + \epsilon)^2) \quad \forall \epsilon \ 0 < \epsilon < 1$$

5 Modificación de la Entropía por τ_ϵ

Para estudiar el efecto de τ_ϵ en relación con distintas formas de entropías π en $[0,1]^E$, con $|E|=n$, se utilizan las entropías más generales y usuales en la bibliografía para universos finitos y se realiza la comparación de entropías $\pi(A)$ y $\pi(\tau_\epsilon(A))$, con $0 < \epsilon < 1$ ya que $\tau_1(A)$ es un conjunto ordinario y por tanto $\pi(\tau_1(A)) = 0$.

En primer lugar se analiza el efecto de τ_ϵ en la entropía según las definiciones axiomáticas, ya que la mayoría de las definiciones generales y particulares de entropía verifican la definición axiomática de De Luca y Termini [2] o la de medida de borrosidad de Tipo 2 de Klir y Folger [8].

Recordemos la definición axiomática de entropía de De Luca y Termini en la que una función

$\pi: [0,1]^E \rightarrow \mathbb{R}^+$ es una entropía si verifica:

$\pi 1.- \pi(A) = 0 \quad \forall A$ subconjunto ordinario de E

$\pi 2.- \pi(A)$ máximo si $\mu_A(x) = 0,5 \quad \forall x \in E$

$\pi 3.- \forall A, B \in [0,1]^E$

$$\left. \begin{array}{l} \mu_B(x) \geq \mu_A(x) \geq 0,5 \\ \mu_B(x) \leq \mu_A(x) \leq 0,5 \end{array} \right\} \Rightarrow \pi(A) \geq \pi(B)$$

Nota: Algunos autores como Xuecheng [15] incluyen en la definición axiomática de entropía un cuarto axioma:

$\pi 4.- \pi(A) = \pi(h(A)) \quad \forall A \in [0,1]^E$

con h pseudo-complementario $h(a) = 1-a$

Proposición 10 Si una entropía π verifica la definición axiomática de De Luca y Termini, entonces

$$\pi(A) \geq \pi(\tau_\epsilon(A)) \quad \forall A \in [0,1]^E \quad \forall \epsilon \ 0 < \epsilon < 1$$

Demostración. Por proposición 1

$$\tau_\epsilon(\mu_A(x)) \leq \mu_A(x) \leq 0,5 \quad \tau_\epsilon(\mu_A(x)) \geq \mu_A(x) > 0,5$$

$$\forall A \in [0,1]^E \quad \forall \epsilon \ 0 < \epsilon < 1 \quad \forall x \in E$$

Con lo cual aplicando el axioma $\pi 3$ de la definición axiomática de De Luca y Termini tenemos

$$\pi(A) \geq \pi(\tau_\epsilon(A)) \quad \forall A \in [0,1]^E \quad \forall \epsilon \ 0 < \epsilon < 1 \quad \blacklozenge$$

La definición axiomática anterior la verifican definiciones particulares como la Entropía no probabilística de De Luca y Termini [5] y los Índices de Borrosidad de Kauffmann [10] [11] en las que se pueden obtener resultados más precisos:

Proposición 11 La entropía no probabilística de De Luca y Termini verifica

$$\pi(A) > \pi(\tau_\epsilon(A)) \quad \forall A \in [0,1]^E \quad \forall \epsilon \ 0 < \epsilon < 1$$

Demostración. La expresión de la entropía no

$$\text{probabilística es } \pi(A) = K \sum_{i=1}^n S(\mu_A(x_i))$$

$$\text{con } K \in \mathbb{R}^+ \text{ y } S(a) = -a \ln a - (1-a) \ln(1-a) \quad \forall a \in [0,1]$$

Por proposición 1

$$\tau_\epsilon(\mu_A(x_i)) \leq \mu_A(x_i) \leq 0,5 \quad \tau_\epsilon(\mu_A(x_i)) \geq \mu_A(x_i) > 0,5$$

$\forall A \in [0,1]^E \quad \forall \epsilon \ 0 < \epsilon < 1 \quad i=1, \dots, n$ con igualdad únicamente para $\mu_A(x_i) = 0$ ó 1

Y por ser la función de Shannon estrictamente creciente en $[0, 0,5]$ y estrictamente decreciente en $(0,5, 1]$, con único máximo $s(0,5)$ y $S(0)=S(1)=0$ tenemos:

$$S(\tau_\epsilon(\mu_A(x_i))) < S(\mu_A(x_i)) \quad i=1, \dots, n$$

$$\forall A \in [0,1]^E \quad \forall \epsilon \ 0 < \epsilon < 1$$

Con lo cual aplicando el sumatorio y multiplicando por una constante K positiva:

$$\pi(A) > \pi(\tau_\epsilon(A)) \quad \forall A \in [0,1]^E \quad \forall \epsilon \ 0 < \epsilon < 1 \quad \blacklozenge$$

Proposición 12 Los índices de borrosidad de Kauffmann verifican

$$\pi(\tau_\epsilon(A)) = (1-\epsilon)\pi(A) \quad \forall A \in [0,1]^E \quad \forall \epsilon \ 0 < \epsilon < 1$$

Demostración. La expresión de los índices de borrosidad de Kauffmann es, para $\omega \geq 1$

$$\pi_\omega(A) = K \left(\sum_{i=1}^n |\mu_A(x_i) - \mu_{\underline{A}}(x_i)|^\omega \right)^{1/\omega}$$

con $K \in \mathbb{R}^+$ y $\underline{A}$ el conjunto ordinario más próximo al conjunto fuzzy A (o sea $\mu_{\underline{A}}(x) = \chi_{\underline{A}}$).

Por la proposición 2 $\tau_\epsilon(A) = \underline{A}$ y entonces

$$\begin{aligned} \pi_\omega(\tau_\epsilon(A)) &= K \left(\sum_{i=1}^n |\tau_\epsilon(\mu_A(x_i)) - \chi_{\underline{A}}|^\omega \right)^{1/\omega} \\ &= K \left(\sum_{i=1}^n |(1-\epsilon)\mu_A(x_i) + \chi_{\underline{A}}\epsilon - \chi_{\underline{A}}|^\omega \right)^{1/\omega} \end{aligned}$$

$$= K \left(\sum_{i=1}^n | (1-\varepsilon) (\mu_A(x_i) - \chi_A) |^\omega \right)^{1/\omega}$$

$$= (1-\varepsilon) \pi(A) \quad \forall A \in [0,1]^E \quad \forall \varepsilon \quad 0 < \varepsilon < 1 \quad \blacklozenge$$

Para introducir la definición axiomática de medida de borrosidad de tipo 2 de Klir y Folger [12] debemos recordar que toda función complemento h posee al menos un equilibrio $e \in [0,1]$ tal que $h(e)=e$, que además es único si h es continua. De esta forma una función $\pi: [0,1]^E \rightarrow \mathbb{R}^+$ es una medida de borrosidad de tipo 2 si verifica:

- $\pi 1.- \pi(A) = 0 \quad \forall A$ subconjunto ordinario de E
 $\pi 2.- \pi(A)$ máximo si $\mu_A(x) = e \quad \forall x \in E$
 con e equilibrio de la función complemento h .
 $\pi 3.- |\mu_A(x_i) - h(\mu_A(x_i))| \geq |\mu_B(x_i) - h(\mu_B(x_i))|$
 implica $\pi(A) \leq \pi(B) \quad \forall A, B \in [0,1]^E$
 con h función complementario fuzzy

Proposición 13 Si π es una medida de borrosidad de tipo 2 que verifica la definición axiomática de Klir y Folger para h pseudo-complementario, entonces

$$\pi(A) \geq \pi(\tau_\varepsilon(A)) \quad \forall A \in [0,1]^E \quad \forall \varepsilon \quad 0 < \varepsilon < 1$$

Demostración. Por proposición 1

$$\tau_\varepsilon(\mu_A(x)) \leq \mu_A(x) \leq 0,5 \quad \tau_\varepsilon(\mu_A(x)) \geq \mu_A(x) > 0,5$$

$$\forall A \in [0,1]^E \quad \forall \varepsilon \quad 0 < \varepsilon < 1$$

$$\text{Luego } |\mu_A(x_i) - h(\mu_A(x_i))| = |2\mu_A(x_i) - 1| \geq$$

$$|\tau_\varepsilon(\mu_A(x_i)) - h(\tau_\varepsilon(\mu_A(x_i)))| = |2\tau_\varepsilon(\mu_A(x_i)) - 1|$$

Con lo cual aplicando el axioma $\pi 3$ de la definición 6 tenemos

$$\pi(A) \geq \pi(\tau_\varepsilon(A)) \quad \forall A \in [0,1]^E \quad \forall \varepsilon \quad 0 < \varepsilon < 1 \quad \blacklozenge$$

Conclusiones

El estudio sobre el efecto de la transformación τ_ε definida por Fadini, Ventre y Cella en [4] nos lleva a la conclusión de que no sólo la intersección (definida con el mínimo) y unión (con el máximo) se conservan con la transformación, sino que también la t -norma Producto drástico y la t -conorma Suma drástica se conservan. Sin embargo ninguna negación continua se conserva al transformarla con τ_ε y tampoco todas las discontinuas.

Por otra parte la transformación τ_ε consigue reducir la indeterminación que conlleva un difuso, ya que las entropías que verifican los axiomas de De Luca y Termini o los de medida de borrosidad de tipo 2 de Klir y Folger disminuyen al aplicar la transformación.

Trabajos futuros

El trabajo que seguirá al aquí presentado irá orientado a la caracterización de las t -normas y t -conormas que verifiquen las propiedades (1) y (2), y de las negaciones que verifiquen (3), así como al análisis de formas concretas de medidas de borrosidad

de tipo 2 y su efecto en relación con las transformaciones τ_ε . Por otra parte se analizarán las propiedades de convergencia de sucesiones de difusos $\{\tau_\varepsilon^n(\mu_A(x))\}_{n \in \mathbb{N}}$ y entropías $\{\pi(\tau_\varepsilon^n(\mu_A(x)))\}_{n \in \mathbb{N}}$

Agradecimientos

Queremos agradecer al Dr. Pedro Burillo su orientación y supervisión de los resultados, así como su interés en colaborar en trabajos futuros.

References

- [1] BURILLO, P., AND FUENTES, R. Perturbations on L-Fuzzy Sets. *Aportaciones Matemáticas en Memoria del Profesor Onieva*, Universidad de Cantabria, Santander, (1991), 43-54.
- [2] DE LUCA, A., AND TERMINI, S. Entropy and energy measures of a fuzzy set. *Advances in Fuzzy Sets Theory and Applications* (North Holland, 1979), 321-338.
- [3] DUBOIS, D., AND PRADE, H. *Fuzzy sets and Systems. Theory and Applications*. Academic Press, 1980.
- [4] FADINI, A., VENTRE, A., AND CELLA, C. Su certi operatori per insieme fuzzy. *Estratto degli "Atti" dell'Accademia Pontiniana Nuova Serie*, Vol. XXXVIII (1978), 1-12.
- [5] GUPTA, M. M., AND QI, J. Theory of T-norms and fuzzy inference methods. *Fuzzy Sets and Systems*, 40 (1991), 431-450.
- [6] KAUFMANN, A. *Introduction a la theorie des sous-ensembles flous a l'usage des ingénieurs. Tomo I. Elements theoriques de base*. Masson, 1977.
- [7] KAUFMANN, A. *Introduction a la theorie des sous-ensembles flous a l'usage des ingénieurs. Tomo IV. Compléments et nouvelles applications*. Masson, 1977.
- [8] KLIR, G. J., AND FOLGER, T.A. *Fuzzy sets, uncertainty and information*. Prentice-Hall International, 1988.
- [9] RIERA, T., AND TRILLAS, E. Entropies in Finite Fuzzy Sets. *Information Sciences*, 15 (1978), 159-168.
- [10] XUECHENG, L. Entropy, distance measure and similarity measure of fuzzy sets and their relations. *Fuzzy Sets and Systems*, 52 (1992), 305-318.
- [11] ZIMMERMANN, H.J. *Fuzzy set theory and its applications*. Kluwer Academic Publishers, 1991.

Aplicación de Algoritmos Evolutivos para el Problema de la Triangulación en Redes Causales.

Luis Daniel Hernández Molinero
Dpto. Informática y Sistemas
Facultad de Informática
Universidad de Murcia
e-mail: ldaniel@dif.um.es

Manuel Jorge Bolaños Carmona
Dpto. CC. Computación e I.A.
E.T.S. de Informática
Universidad de Granada
e-mail: mjbcb@robinson.ugr.es

Resumen

Inference problem in causal networks can be seen as the building of a cluster tree in order to allow the passing of local messages between adjacent clusters. The usual way of obtain this tree is by triangulation of the corresponding moral graph and the following use of its cliques to build the tree. The main question is the efficiency of the computations involved in the propagation, which is better as smaller is the size of the cliques. This size is determined by the form of the triangulation. For measures more complex then probabilities, as fuzzy measures, the problem of the efficiency is worse.

In this paper we propose a evolutive algorithms based methodology in order to obtain the best triangulation and the results are compared with those coming from other knowed methods.

Keywords: Causal networks, triangulation of graphs, fuzzy measures, evolutive algorithms.

1 Introducción

Shachter, Andersen y Szolovits [19] han mostrado que los diferentes métodos exactos probabilísticos para la propagación en redes causales que aparecen en la literatura [11, 16, 14, 20, 3, 9, 21] son todos casos particulares de un algoritmo sencillo y general que han llamado Algoritmo Clustering. En base a a este algoritmo se demuestra que la eficiencia en la inferencia probabilística depende de la factorización de la distribución de probabilidad conjunta bajo hipótesis de independencia condicional. Tanto la factorización como la independencia se muestran de forma implícita en una estructura gráfica llamada árbol cluster. El punto en común entre los distintos métodos exactos es que todos construyen un grafo de cluster, y la diferencia estriba en que cada uno busca el que resulta mas adecuado para sus operaciones de inferencia. Es decir, las diferencias entre los métodos pueden entenderse como diferentes aproximaciones para desarrollar las mismas tareas en el algoritmo general. Así mismo, se han obtenido resultados para medidas difusas, en particular para medidas de evidencia [22, 1, 2], que también necesitan de la construcción de un árbol cluster para la propagación, si bien el problema de

eficiencia aparece con mayor gravedad.

La construcción del árbol cluster se basa en la triangulación del grafo moral obtenido a partir de la estructura dirigida original; los cluster que integran dicho árbol se obtienen directamente del proceso de triangulación. Dado que la posterior inferencia será tanto mas óptima cuanto menor sea el tamaño de los clusters que se forman, el punto clave para que un método exacto sea eficiente está en la realización de una buena triangulación. Si bien la construcción de la triangulación óptima es un problema NP-duro [24], existen algoritmos que obtienen buenas triangulaciones usando una cantidad de tiempo razonable [18, 23, 5, 12].

En este trabajo proponemos nuevas técnicas de triangulación que se comparan con los algoritmos mas destacados en estos ultimos años usando un generador de grafos dirigidos acíclicos. El trabajo se ha organizado como sigue. En la siguiente sección presentamos brevemente los conceptos fundamentales relacionados con la triangulación. A continuación se presenta la técnica de los programas evolutivos y se relaciona con el problema de la triangulación. Por último se indica como se han evaluado y comparado la técnica presentada con los métodos heurísticos determinísticos conocidos y finalmente presentamos las conclusiones.

2 La Triangulación de Grafos. Notación y definiciones.

Supongamos un conjunto de variables $X = \{X_1, X_2, \dots, X_n\}$, donde cada variable X_i toma valores discretos en un conjunto Ω_i . Un modo de expresar las independencias probabilísticas entre las variables es por medio de un grafo cuya misión es representar a través de su topología las independencias de las variables [16]. Normalmente la representación consiste en establecer una relación uno a uno entre el conjunto $N = \{1, 2, \dots, m\}$ de vértices del grafo y el conjunto X de variables (i.e. $i \leftrightarrow X_i$ y $m = n$). Estos grafos $G = (N, R)$ son acíclicos dirigidos y reciben, entre otros, los nombres de red causal [10], red de creencia [17] o diagrama de influencia probabilística [20].

Sin embargo para la propagación de incertidumbre probabilística se utiliza un árbol cluster. Éste es un

árbol join¹ construido a partir de una red causal y en el que cada familia² de la red se encuentra completa en al menos un cluster. Si el árbol de clusters está formado por los cliques³ de un grafo diremos que es un **árbol de cliques**. En la práctica, la construcción del árbol cluster se realiza mediante los siguientes pasos:

1. Construcción del **grafo moral** de la red bayesiana. El grafo moral de un grafo dirigido acíclico se construye añadiendo arcos no dirigidos entre cada par de nodos padres de un nodo dado e ignorando la dirección de los arcos del grafo original.
2. Construcción de un **grafo cordal** a partir del grafo moral. Un grafo no dirigido se dice que es **cordal** o **triangulado** si y solo si cada ciclo de longitud cuatro o mas tiene un arco entre cada par de nodos no adyacentes del ciclo.
3. Obtención de los cliques del grafo triangulado y construcción del árbol

La **eliminación** de un vértice es el proceso por el cual se añaden las aristas necesarias para que el vértice y sus nodos adyacentes formen un subgrafo completo y el posterior borrado del vértice y los arcos incidentes en él. Notaremos por C_i al subgrafo completo que se produce en la eliminación del nodo i . La **técnica de la triangulación** puede verse como un proceso que consiste en añadir arcos extras al grafo original producidos por la eliminación de los vértices del grafo uno a uno. Desde este punto de vista la técnica de la triangulación de grafos consiste en establecer ordenaciones de los vértices que especifican la secuencia en la cual deberán de eliminarse. Una ordenación se considera óptima con respecto a las demás cuando el coste medio de los cluster que produce dicha ordenación es menor que el coste producido por el resto de las secuencias. En el contexto de las redes causales el coste de un cluster C_i formado por la eliminación del nodo i suele venir determinado por alguna de estas medidas [12, 13]:

$$\text{Tamaño} \equiv T(C_i) = \sum_{j \in C_i} 1 = |C_i| \quad (1)$$

$$\text{Tamaño natural} \equiv S(C_i) = \prod_{j \in C_i} c(j) \quad (2)$$

$$\text{Peso} \equiv W(C_i) = \sum_{j \in C_i} w(j) \quad (3)$$

¹Un árbol join es un árbol no dirigido cuyos nodos son clusters (i.e. conjuntos de variables) y verifica que si una variable aparece en dos clusters, aparece también en los clusters que integran el camino que los une.

²En un grafo dirigido se llama familia de un nodo i al conjunto $\{i\} \cup \Pi(i)$, donde $\Pi(i)$ denota a los padres de i .

³Un clique de un grafo G es un subgrafo completo maximal de G .

donde $c(j)$ denota el número de casos de X_j y $w(j) = \log_2 c(j)$. El grafo triangulado obtenido por la eliminación de los vértices del grafo $G = (N, R)$ para una ordenación σ está formado por el par $G_\sigma^t = (N, R \cup T)$, donde $T = T(G, \sigma)$ denota al conjunto de nuevas aristas que se han añadido al grafo original; los arcos de $T(G, \sigma)$ reciben el nombre de "arcos de relleno" (i.e. fill edge). Podemos encontrar tres metodologías fundamentales para obtener este conjunto T o lo que es lo mismo, la secuencia de eliminación σ :

1. Las basadas en resultados exactos teóricos. Buscan el modo de producir la mejor eliminación identificando el conjunto T para conseguir triangulaciones minimales⁴ [18, 23, 5].
2. Las basadas en funciones heurísticas. Establecen un criterio de ordenación basado en la regla de búsqueda "el siguiente nodo a eliminar es el que minimiza $f()$ " donde $f()$ es la función de una o varias medidas sobre un conjunto de nodos del grafo $G = (N, R)$. Las medidas mas utilizadas en la literatura suelen venir dadas en función de las funciones-coste de las expresiones (1), (2) y (3) [12]. A estos métodos los llamaremos métodos heurísticos determinísticos.
3. Las basadas en resultados estocásticos. Se basan en partir de un posible conjunto de soluciones factibles, para a partir de ellas ir mejorando por métodos aleatorios y en pasos sucesivos las soluciones óptimas obtenidas en cada iteración [13]. A estos métodos los llamaremos métodos heurísticos estocásticos.

3 Programas Evolutivos y Triangulación.

En estos últimos años han surgido nuevos tipos de técnicas para la resolución de problemas combinatorios: Los programas evolutivos (P.E.) [15]. Estos algoritmos son modificaciones mejoradas y generalizadas de los algoritmos genéticos [4, 6, 8] utilizados para problemas de optimización y difieren de éstos en que se crean estructuras de datos, y funciones, adaptadas al problema combinatorio. Un P.E. se desarrolla del siguiente modo

1. Se parte de un conjunto de posibles soluciones $S(t)$ (inicialmente $t = 0$) que llamaremos **población inicial**.
2. Se evalúa cada uno de los elementos σ de $S(t)$ según una determinada función $eval(\sigma)$.
3. Si se cumple la condición de parada, entonces parar; en caso contrario hacer $t \leftarrow t + 1$ e ir al siguiente paso.

⁴Se dice que T es minimal si para cualquier subconjunto T' de T , el grafo $(N, R \cup T')$ no está triangulado (i.e. no hay arcos de relleno redundantes)

4. Seleccionar un nuevo conjunto de individuos $S(t)$ a partir de $S(t-1)$. En general la selección se realiza eligiendo aleatoriamente a los individuos de la nueva población de dos en dos, y para cada pareja se puede hacer o no el intercambio de la información contenida en ambos (si esto se hace, hablaríamos del **operador cruce**). La selección de esos pares de individuos se realiza por separado atendiendo, para cada uno de ellos, a una determinada **probabilidad de selección**, y el proceso se lleva a cabo, con o sin cruce, hasta que las parejas seleccionadas completen el número de elementos de la nueva población.
5. Hacer posibles modificaciones en los individuos de $S(t)$. Se trata de realizar alteraciones individuales (si procede) en cada uno de los individuos, por lo que también se denomina a este paso **operador mutación**. El que un individuo quede o no alterado viene dado por una determinada **probabilidad de mutación** que, en general, se pone en juego para todos y cada uno de los elementos de la población del paso t .
6. Ir al paso 2.

Para el diseño de un algoritmo como éste, es necesario al menos tomar decisiones respecto a los parámetros y condiciones siguientes:

1. Representación del problema.
2. Número y determinación de individuos necesarios para construir una población inicial.
3. Función de valoración a utilizar.
4. Probabilidad de selección
5. Operador de cruce
6. Probabilidad de mutación
7. Operador de mutación
8. Condición de parada.

A continuación mostramos como puede enfocarse cada uno de estos pasos para el problema particular de la triangulación.

• **REPRESENTACIÓN DEL PROBLEMA.** Si partimos de un grafo $G = (N, R)$ con $N = \{1, 2, \dots, n\}$, grafo moral de alguna red causal, la representación mas natural de un orden de eliminación sería mostrar un lista de los nodos del grafo de tal modo que el nodo i se encuentra en la posición k sii el nodo i es el k -ésimo en ser eliminado. Es decir, un individuo del espacio de búsqueda viene representado por un elemento del conjunto de permutaciones de n -elementos, o en otras palabras, el espacio de búsqueda es el conjunto de permutaciones de n -elementos

$$S_n = \{(\sigma(1), \sigma(2), \dots, \sigma(n)) \mid \sigma(i) \in N\}$$

Ahora bien, como la eliminación previa de los nodos que verifican alguna de estas condiciones

1. El nodo no tiene hijos.
2. El nodo no tiene ningún padre y un solo hijo.
3. El nodo tiene un solo padre y un solo hijo, y además, ese padre verifica esta condición o la anterior.

no producen arcos de relleno en la triangulación, quiere decir que el grafo obtenido es un grafo "representante" del grafo moral original [7]. Según lo expuesto, tan solo deberemos de establecer el problema de la triangulación para aquellos nodos "sobrevivientes" a esta eliminación previa y en consecuencia nuestro espacio de búsqueda se encuentra en un conjunto S_m con $m \leq n$. Notese que la secuencia completa de una triangulación $\sigma \in S_n$ viene dada por $\sigma = (\sigma', \sigma'')$ con $\sigma' \in S_{n-m}$, $\sigma'' \in S_m$ y $\sigma(i) \neq \sigma(j) \forall i, j$.

• **NÚMERO Y DETERMINACIÓN DE INDIVIDUOS DE LA POBLACIÓN INICIAL.** El número inicial de individuos puede establecerse atendiendo a una función $f()$, como por ejemplo, alguna que venga determinada por el grado de complejidad del grafo. En la práctica nosotros hemos tomado $f(n) = (n \times m)/100$, con n el número de nodos del grafo moral, es decir, un cierto porcentaje m del número n .

La determinación de los individuos iniciales puede realizarse mediante la generación aleatoria de individuos $\sigma \in S_n$. Sin embargo, podría ser conveniente que no existan dos individuos que tengan el primer elemento igual; es decir, si σ y σ' son dos individuos de la población inicial $S(0)$, entonces $\sigma(1) \neq \sigma'(1)$, y ello para dar gran diversidad a la población inicial, especialmente al comienzo de la secuencia.

• **FUNCIÓN DE EVALUACIÓN DE LOS INDIVIDUOS.** Las funciones de valoración de los individuos para nuestro caso viene dada por la naturaleza del problema; es una función de la forma:

$$eval(\sigma) = \frac{M(G_\sigma^t)}{m_\sigma} = \frac{1}{m_\sigma} \sum_{C_i} M(C_i) = \overline{M}(G_\sigma^t)$$

donde $\sigma \in S_n$, G_σ^t denota al grafo triangulado obtenido para la ordenación σ ; m_σ denota al número de clusters obtenidos para esa ordenación y $M(C_i)$ denota a alguna de las medidas de las expresiones (1), (2) y (3). Notese que nuestro objetivo es encontrar aquella ordenación σ que produzca el menor tamaño o peso medio.

• **PROBABILIDAD DE SELECCIÓN.** La selección de los individuos se realiza en general atendiendo al valor $eval()$ obtenido en el paso anterior de tal modo que los individuos mejores tengan mayor probabilidad de ser seleccionados. Como es fácilmente comprensible, existen muy diversos criterios para establecer esa probabilidad. Para nuestro problema, nos parece razonable realizar una selección basada en "un criterio elitista". Dada una población $S(t)$, en la que suponemos que los individuos se encuentran ordenados de menor a mayor tamaño (o peso) medio, generaremos la población $S(t+1)$ del siguiente modo:

1. Se selecciona el mejor $X\%$ de los individuos de $S(t)$ (i.e. los primeros $X\%$ individuos) y pasan a la siguiente población de forma directa (i.e. ocupan directamente las $X\%$ primeras posiciones de $S(t+1)$). A éstos individuos los denotamos por $SM(t) \equiv$ Población Mejor.
2. Se selecciona el peor $Y\%$ de los individuos, y éstos son totalmente descartados. Los denotamos por $SP(t) \equiv$ Población Peor.
3. El $(100-X-Y)\%$ restante lo denotaremos por $SI(t) \equiv$ Población Intermedia. Esta población intermedia, junto con la población mejor, lo que representa en total un $(100-Y)\%$, se utilizan para generar el resto $(100-X)\%$ de la siguiente población $S(t+1)$. Ello se realiza mediante la adopción de un determinado criterio probabilístico. Éste podría ser, en nuestro caso, seleccionar a cada individuo con probabilidad

$$P_s(\sigma^i = \sigma) = \begin{cases} \frac{1/eval(\sigma)}{\sum_{i=1}^k 1/eval(\sigma_i)} & \text{si } \sigma \in SM(t) \cup SI(t) \\ 0 & \text{si } \sigma \in SP(t) \end{cases}$$

donde $P_s(\sigma^i = \sigma)$ denota la probabilidad de que el individuo que ocupe el lugar i (σ^i) en la nueva población $S(t+1)$ sea el individuo $\sigma \in S(t)$ y donde $t \in \{1, 2, \dots, k\}$ denota el recorrido del $(100-Y)\%$ de la población (i.e. $SM(t) \cup SI(t) \subset S(t)$). Hay que notar que los que producen una mayor evaluación $eval$ son los peores (p.e. producen mayores cluster) por ello daremos mayor probabilidad a los que produzcan menor valor en la evaluación.

• **OPERADOR CRUCE.** La selección de los individuos para integrar la población $S(t+1)$ que acabamos de exponer puede ser realizada de dos en dos, y para cada pareja, y según una cierta probabilidad que puede ser fija, se realiza o no un "cruce" de información de la siguiente forma. Consideremos dos individuos σ y σ' :

$$\begin{aligned} \sigma &= (\sigma(1), \sigma(2), \dots, \sigma(m)) \\ \sigma' &= (\sigma'(1), \sigma'(2), \dots, \sigma'(m)) \end{aligned}$$

que proceden de la población $S(t)$ y nuestro objetivo es conseguir dos nuevos individuos d y d' para la población $S(t+1)$. Estos individuos se obtienen realizando un intercambio de información entre σ y σ' , buscando algún modo por el que los nuevos elementos d y d' contengan "lo mejor" de los que los generan.

Por tanto hay que establecer cual es la mejor información que un individuo puede traspasar a otro. En nuestro caso, se trata de establecer cual es la mejor subsecuencia de nodos de un individuo σ que producen un menor tamaño; pero como la eficiencia de una subordenación depende de la secuencia previa, consideramos que la mejor información que debe traspasar un individuo a otro está formada por sus primeros elementos.

De este modo, si la importancia de una ordenación se encuentra en sus nodos iniciales, consideraremos

que un cruce consiste en tomar los k -primeros elementos de una ordenación σ y reordenarlos según el orden que define el otro individuo σ' , con k elegido aleatoriamente entre 1 y m .

• **PROBABILIDAD DE MUTACIÓN.** Para cada individuo de la nueva población $S(t+1)$ (resultado o no del cruce de dos elementos de $S(t)$), éste puede ser o no modificado atendiendo a una cierta probabilidad de mutación. Como en nuestro problema crece la complejidad en función de $n!$ y el programa evolutivo tiene que competir con los métodos conocidos, y teniendo en cuenta que el costo real se encuentra en la valoración de los individuos, deberíamos de optar por considerar poblaciones pequeñas y en consecuencia con una alta probabilidad de mutación constante o creciente. Esto es para evitar "probablemente" el que nos centremos en extremos locales del objetivo al usar una cantidad pequeña de individuos representativos del espacio de búsqueda total.

• **OPERADOR DE MUTACIÓN.** El operador de mutación clásico nos dice que debemos de alterar un alelo (valor de una posición del individuo) por otro valor del alfabeto considerado (i.e. $M = \{1, 2, \dots, m\}$). En nuestro caso, como no pueden existir dos alelos iguales (supondría eliminar un nodo dos veces) y utilizamos todos los valores del alfabeto, si alteramos un alelo (i.e. $\sigma(1)$) y lo cambiamos de valor deberemos de considerar el alelo que contenía este valor (i.e. $\sigma(2)$) y modificarlo por otro (i.e. $\sigma(3)$), lo cual supone ir a este tercer alelo (i.e. $\sigma(3)$) y cambiarlo de valor (i.e. $\sigma(4)$) lo que supone ir a un quinto, y así sucesivamente. Notar que esto genera un procedimiento que produciría la alteración completa del cromosoma. Obviamente hay que establecer un límite para que esto no se produzca y lo llamaremos nivel de mutación. En este sentido, y como caso particular, existe un operador de mutación básico que se conoce por *inversión* que lo que hace es cambiar dos valores de posición.

• **CRITERIO DE PARADA.** Si en el desarrollo del programa evolutivo observamos que los costos de los mejores individuos no se modifican o difieren en una cantidad pequeña podemos concluir que el algoritmo ha terminado. Sin embargo, como el algoritmo ha de ser competitivo con los métodos heurísticos determinísticos (en velocidad) hemos optado por fijar el número de poblaciones a generar.

4 Resultados.

Los resultados se han desarrollado con un programa en lenguaje C++ y se ha trabajado tanto sobre PCs como en SUN Workstation. El modo de evaluar los algoritmos evolutivos ha consistido en generar aleatoriamente un conjunto de grafos cuyos nodos se enumeran desde el 1 hasta un n determinado. El modo en el que se eligen los vecinos de un nodo y los casos de las variables involucradas, se realiza de forma interactiva con el usuario.

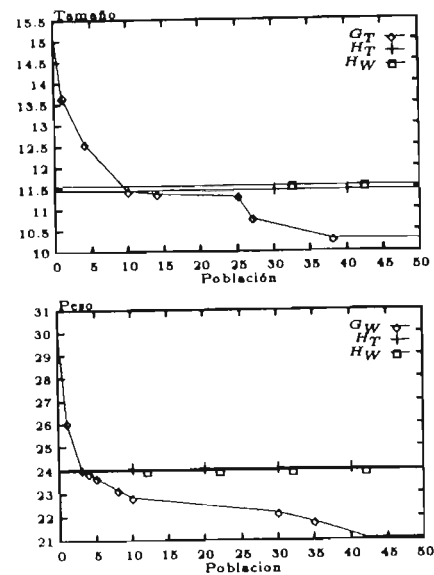
En este trabajo presentamos un ejemplo de los resultados obtenidos sobre una familia de grafos que

- 50 generaciones o iteraciones.
- 10 individuos (o posibles soluciones) en cada población.
- En cada iteración pasan a la siguiente población el 50 % de la población (5 individuos) y se descartan directamente el 10% (1 individuo).
- Probabilidad de cruce igual a 0.8
- Probabilidad de mutación a 0.3
- Nivel de mutación a 12 intercambios.

$l=0$	T	W	C	V
H_T	5.891	13.269	29.9	175.9
H_W	5.901	13.277	29.86	175.92
G_T	5.895	13.278	29.8	175.36
G_W	5.899	13.282	29.86	175.88
$l=1$	T	W	C	V
H_T	3.249	7.189	44.2	141.18
H_W	3.271	6.976	44.16	142.04
G_T	3.207	7.104	43.4	137
G_W	3.247	6.780	43.34	138
$l=2$	T	W	C	V
H_T	15.464	34.665	20.72	316
H_W	15.480	34.548	20.74	316.7
G_T	12.191	27.348	11.72	139.04
G_W	12.231	27.370	11.8	140
$l=20$	T	W	C	V
H_T	12.253	27.704	25.8	313.96
H_W	12.279	27.537	25.76	314
G_T	10.931	24.697	16.5	179.38
G_W	10.990	24.690	16.36	178.04
$l=100$	T	W	C	V
H_T	9.948	22.479	29.2	286.78
H_W	9.966	22.309	29.2	287.32
G_T	9.295	20.932	22.1	202.48
G_W	9.401	21.057	22.14	204.88

Tabla 1: Un test sobre 500 grafos.

reflejan en gran medida los resultados obtenidos en otras familias. La que aquí se presenta consta de 500 grafos con 50 nodos cada uno. Los parámetros utilizados para su generación son los siguientes: A cada nodo se le asigna un número de casos siguiendo una distribución de Poisson de media 5.0 y un mínimo de 2 casos. El número de hijos generados para cada nodo se realiza atendiendo a una distribución uniforme de media 5 y éstos se eligen entre los siguientes a uno dado en el orden de etiquetado. Estos 500 grafos se han dividido en 5 grupos de 100 grafos cada uno atendiendo a un parámetro l que afecta a la probabilidad con que se elegirán los siguientes nodos como hijos del nodo que se esté considerando. El primer grupo responde a aquellos grafos donde los hijos son los nodos inmediatamente siguientes al considerado (i.e. $l=0$); en el segundo grupo los nodos siguientes a uno dado son igual de probables en ser elegidos (i.e. $l=1$) y en los restantes grupos los nodos mas cercanos a uno dado son mas probables de ser elegidos cuanto mayor sea el valor de l (i.e. $l > 1$).

Figura 1: Evolución para un grafo de la familia $l = 20$

En la tabla 1 puede verse los parámetros evolutivos utilizados y los resultados obtenidos para cada uno de los grupos de esta familia de grafos. El significado de las abreviaturas que aparecen en las filas y columnas es como sigue:

1. Parámetro que afecta a la probabilidad con que se seleccionan los nodos siguientes como vecinos a uno dado.

T . Tamaño medio de los clusters.

W . Peso medio de los clusters.

C . Número medio de clusters generados.

V . Número medio de variables en el grafo cluster.

H_T . Técnica heurística basada en menor tamaño.

H_W . Técnica heurística basada en menor peso.

G_T . Técnica genética basada en menor tamaño.

G_W . Técnica genética basada en menor peso.

5 Conclusiones

Los algoritmos evolutivos permiten llegar a triangulaciones que algunos método heurísticos determinísticos no pueden alcanzar, llegando incluso a reducir en un 56% el número de cluster necesarios y en casi un 44% el número de variables a manejar (p.e. ver tabla 1 para $l = 2$). Este resultado es análogo al obtenido por Kjærulf con la simulación por enfriamiento [13]. De hecho, haciendo un análisis exhaustivo de los grafos estudiados hemos observado que los algoritmos evolutivos partían en general de configuraciones con un alto coste reduciéndose éste en muy pocas evoluciones; es más, han bastado un número pequeño de iteraciones y un número muy pequeño de individuos en la

población para llegar a mejores triangulaciones que las que alcanzan los métodos heurísticos determinísticos (p.e. ver figura 1).

Si el parámetro tiempo de ejecución no se considera esencial recomendamos el uso de los P.E.. Es más, podemos simplificar drásticamente el número de ejecuciones si realizamos una hibridación del algoritmo consiste en usar como población inicial la generada por los individuos-solución de uno o varios métodos heurísticos determinísticos; esta población podría entonces entenderse como la resultante de haber aplicado el P.E. un número previo de iteraciones.

Referencias

- [1] Cano J.E., Delgado, M., Moral, S. An axiomatic framework for the propagation of uncertainty in directed acyclic graphs. *International Journal of Approximate reasoning*, (8):253-280, 1993.
- [2] Cano Ocaña, J.E. *Propagación de probabilidades inferiores y superiores en grafos*. PhD thesis, Dpto. de C.C. e I.A. Facultad de Ciencias. Universidad de Granada, 1993.
- [3] Cooper, G.F. Bayesian belief-network inference using recursive decomposition. Technical Report KSL-88-27, Knowledge systems laboratory, Stanford University, California, 1990.
- [4] Davis, L. *Genetic algorithms and simulated annealing*. London: Pitman, 1987.
- [5] Fujisawa, T. and Orino, H. An efficient algorithm of finding a minimal triangulation of a graph. *IEEE International Symposium on Circuits and Systems*, pages 172-175, 1974.
- [6] Goldberg, D.E. *Genetic algorithms in search, optimization, and machine learning*. Addison-Wesley, 1989.
- [7] Hernández Molinero, L.D., Bolaños Carmona, M.J. New algorithms to triangulate causal networks. Technical report, Universidad de Granada, 1994. (To appear).
- [8] Holland, J.H. *Adaptation in natural and artificial systems*. The University of Michigan Press, Ann Arbor, 1975.
- [9] Jensen, F.V., Lauritzen, S.L., and Olesen, K.G. Bayesian updating in causal probabilistic networks by local computations. *Computational Statistics Quarterly*, pages 269-282, 1990.
- [10] Jensen, F.V., Olesen, K.G. and Andersen, S.K. An algebra of bayesian belief universes for knowledge based systems. *Networks*, (20):637-659, 1990.
- [11] Kim, J.H. and Pearl, J. A computational model for causal and diagnostic reasoning in inference engines. *8th. International Joint Conference on Artificial Intelligence, Karlsruhe, West Germany*, 1983.
- [12] Uffe Kjærulff. Triangulation of graphs-algorithms giving total state space. Technical Report R 90-09, Department of Mathematics and Computer Science, Institute for electronic Systems, Aalborg University, 1990.
- [13] Uffe Kjærulff. Optimal descompositon of probabilistic networks by simulated annealing. *Statistics and Computing*, (2):1-21, 1992.
- [14] Lauritzen, S.L. and Spiegelhalter, D.J. Local computations with probabilistic on graphical structures and their application to expert systems. *J.R. statistics Society B.*, 2(50):157-224, 1988.
- [15] Michalewicz, Z. *Genetic algorithms + data structures = evolution programs*. Springer Verlag, 1992.
- [16] Pearl, J. A constraint-propagation approach to probabilistic reasoning. In L.N. Kanal and J.F. Lemmer, editor, *Uncertainty in Artificial Intelligence* (pp 357-370). Amsterdam: North Holland, 1986.
- [17] Pearl, J. Fusion, propagation and structuring in belief networks. *Artificial Intelligence*, 3(29):241-288, 1986.
- [18] Rose, D.J., Tarjan, R.E. and Lueker, G.S. Algorithmic aspects of vertex elimination on graphs. *SIAM Journal on Computing*, (5):266-283, 1976.
- [19] Schachter R.D., Andersen S.K., Szlovits, P. The equivalence of exact methods for probabilistic inference on belief network. *Submitted to Artificial Intelligence*, 1991.
- [20] Shachter, R.D. Probabilistic inference and influence diagrams. *Operations Research*, 36(4):589-604, July-August 1988.
- [21] Shafer, G. Shenoy, P.P. Probability propagation. *Annals of Mathematics and Artificial Intelligence*, (2):327-351, 1990.
- [22] Shenoy, P.P., Shafer, G. Axioms for probability and belief functions propagation. In R.D. Shachter, T.S. Levitt, L.N. Kanal, J.F. Lemmer, editor, *Uncertainty in artificial intelligence*, number 4, pages 169-198. Elsevier science publisher B.V. (North-Holland), 1990.
- [23] Tarjan, R.E. and Yannakakis, M. Simple linear-time algorithms to test chordality of graphs, acyclicity of hypergraphs, and selectively reduce acyclic hypergraphs. *SIAM Journal on Computing*, 3(13):566-579, 1984.
- [24] Wen, W.X. Optimal decomposition of belief networks. *Proceedings of the Sixth Workshop on Uncertainty in Artificial Intelligence*, (Cambridge, MA):245-256, 1990.

Una estructura elemental de conocimiento fuzzy

Angel Núñez¹ y Jesús Delegido²

(1) Instituto Tecnológico "Blasco Ibañez". C/
Antig Regne de Valencia 46. Valencia 46005.
Tno 96-3732712

(2) Departamento de Termología. Facultat de
Farmàcia. Universitat de València. C/ V. Andrés
Estellés. Burjassot. Valencia 46100. Tno 96-
3864902. FAX 96-3864813

Resumen

Este trabajo está dedicado a la búsqueda y al análisis de la estructura reticular mínima que soporta la vaguedad.

Palabras clave: experto, borroso, retículo.

1 Introducción

La posibilidad, límite y validez, origen y naturaleza del conocimiento, constituyen la base fundamental de los estudios gnoseológicos, como intento permanente de comprender el hecho cognoscitivo y, en consecuencia, el avance de la ciencia. Ya entre los antiguos, desde Platón y Aristóteles hasta nuestros días, se han venido debatiendo esos grandes temas del conocimiento, y casi sin ningún resultado concluyente (que no sea otro que aquél que advierte que el problema continua abierto).

Así como el desarrollo de las ciencias de lo social y de la naturaleza ha venido profundamente afectado por la controversia epistemológica (fundamentalmente: ¿cómo puedo conocer?), no ha sido así en el caso de la matemática. El pensamiento matemático es un esquema abstracto construido bajo el principio del *tercio excluso*. Como es bien conocido, este principio soporta una lógica formal absoluta, capaz de deslindar el hecho cierto (o verdadero) de aquél que no lo es, y dentro de la cual, la evaluación de sus predicados desde el retículo booleano, se realiza mediante una relación de

equivalencia expresada como sigue: x pertenece al conjunto X si existe un y en X , tal que x es igual a y . El operador evaluador, que se denomina *función característica*, es invariante y único, por tanto está en condición de producir total certeza y exactitud.

El modelo de la matemática, que denominamos ordinario, surgido de lo indicado anteriormente, tiene tanta transcendencia que si quitásemos el 1 y el 0 de nuestro acervo cultural, produciríamos un drama social de consecuencias imprevisibles. Pero, al mismo tiempo, parece observarse cada día un mayor consenso en que, posiblemente, el modelo de la matemática ordinaria no sea el más adecuado para abordar cierta clase de problemas, como por ejemplo las cuestiones relativas al análisis, emulación y aplicación artificial del conocimiento humano, ya que éste viene caracterizado en mayor medida por lo aproximado que por lo exacto. En su intento de aprehender la realidad, va casi siempre acompañado por la incertidumbre, lo que, por otro lado, no impide que su avance sea efectivo.

Partiendo de estas consideraciones, la reflexión que ofrecemos en este trabajo tiene como objetivo proponer una estructura reticular mas débil que supere las valoraciones del retículo ordinario. Se trata pues, de un intento de escrutar la posibilidad de construir, a estas alturas del problema, algún tipo de sistematización básica, en una aproximación a

lo que se viene describiendo como lógicas del conocimiento humano. En esta línea, hacemos uso de un modelo de experto (o ente cognoscente), elegido para que valore un conjunto de alternativas existentes en un problema P , alternativas que se presentan como posibles soluciones al problema. A dicho "experto" se le asocia una función experta mediante la cual evalúa sus propias valoraciones. En tales valoraciones el experto va a usar siempre la unidad dialéctica conocimiento-ignorancia. Esta elección tiene como soporte las condiciones que se especifican con detalle a continuación. Como es obvio, siempre que hagamos referencia al "experto", estaremos hablando, en realidad, de su modelo lógico.

2 Definición

Sea X un conjunto de alternativas $\{x\}$ ofertadas como posibles soluciones a un determinado problema P . Sea R el cuerpo real y $[0,1] \subset R$

Definición 2.1. Diremos que la función:

$$e: X \rightarrow]0,1[\quad (1)$$

es una función experta o función de conocimiento asociada a un experto e , si se verifican las dos condiciones siguientes:

1ª) $e(x)$ está bien definida en X .

2ª) $e(x) \leq \alpha < 1$ donde α es el grado de conocimiento que el experto e tiene de la alternativa $x \in X$

De la definición se deduce que el *sop* $e(x)$ es todo X . Por tanto el experto mantiene un cierto grado de conocimiento de todos los elementos de X .

$$A = \{(x_i, e(x_i)) \mid x_i \in X, e(x_i) \leq \alpha_i < 1\} \quad (2)$$

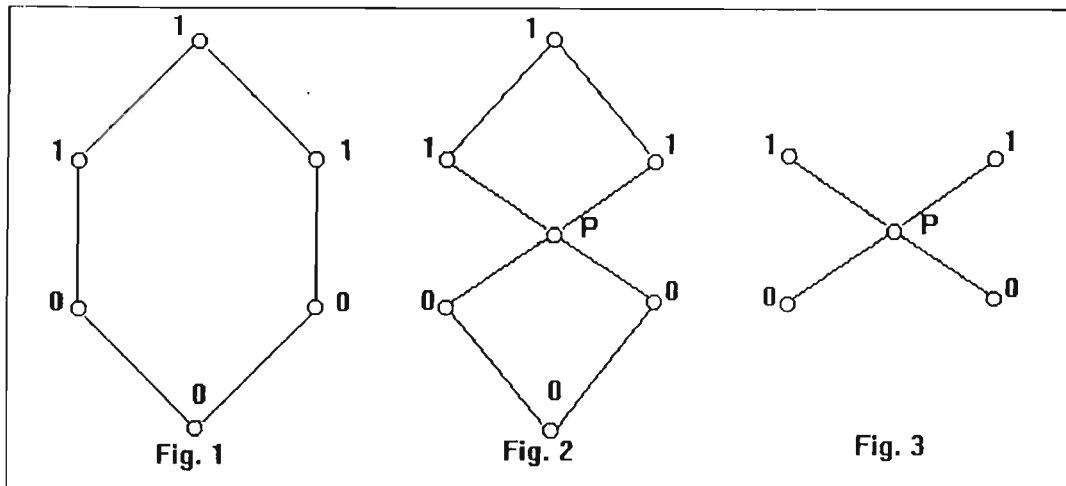
Se considera aquí que el conocimiento humano es en parte estructural y otra, la mayor, adquirida. Aquí sólo nos interesa la primera parte de la consideración, entre otras razones, porque sin ésta no puede hablarse de conocimiento. De otra parte es obvio que el predicado "conocer" o tener conocimiento, es un

predicado borroso. Si unimos los dos enunciados de las partes "estructural" y "conocimiento", con la cualidad de borrosidad que hemos asociado a éste último enunciado, parece aconsejable el uso de estructuras reticulares como una de las herramientas mas apropiada capaz de acercarnos al estudio del conocimiento humano, y definir, con cierto grado de aproximación, algunas cualidades que ofrecen dichas estructuras a la hora de ofertar tipos de valoración sobre un conjunto X , cuyos elementos, como definimos en (1), son las diferentes alternativas existentes sobre un problema P . La gradación en las estructuras reticulares booleanas nos es posible.

Observemos que la Fig. 1, con una estructura de seis elementos, representa lo que podríamos llamar conocimiento ordinario, ya que en la estructura reticular que queda al eliminar 1 como máximo y 0 como mínimo nos aparece el álgebra del complementario: es decir, el conocimiento ordinario.

El retículo completo siguiente al ordinario que expresa la Fig. 2 y su estructura reticular nacida por (1) nos muestra una estructura que no puede expresarse mediante un producto cartesiano, como sucedió en la Fig. 1. La pregunta es inmediata ¿qué modelo de conocimiento representa esa estructura reticular?. Podemos afirmar que, aún pudiendo alcanzar P , ya sea por *sup inf* o por *inf sup*, el sistema reticular es irreducible ya que P no puede alcanzar ni el 1 ni el 0. Por tanto nos encontramos con la mínima estructura reticular que supera a la ordinaria. Es decir, entre la "verdad" y la no "verdad" existe un elemento que comparte las partes.

Situemos esta estructura reticular en tareas de representar conocimiento. Para ello vamos a usar, como herramienta de análisis, la misma unidad dialéctica, contingencia-necesidad, propia del conocimiento natural, traducida aquí como unidad conocimiento-ignorancia. La razón de usar este tipo de unidad se sustenta, además de lo señalado anteriormente, en su adaptabilidad al razonamiento fuzzy.



De estas consideraciones y volviendo a la Fig. 1, podemos establecer en $(1,0) \times (0,1)$ que $(1,0)$ representa la categoría del conocimiento y la $(0,1)$ la categoría de la ignorancia.

Visto lo anterior, dichas categorías actúan en paralelo, dando paso al concepto de complementario. Aparece así el esquema reticular que soporta lo que llamamos el conocimiento ordinario o conocimiento categórico.

De otra parte y fijándonos en la Fig. 3, P representa un **punto de indecisión**, ya que dicho punto P tiene tanto "conocimiento como ignorancia", separando en tres estadios la estructura reticular. Si consideramos que nuestro experto "e" lleva como estructura de conocimiento la Fig. 3, hablaremos de que en el experto domina el conocimiento cuando sus valoraciones se sitúan por encima del "punto P". El entrecomillado anterior viene a indicar que ese "punto" representa la zona de cambio de tendencia en el dominio la unidad. Mas adelante volvemos sobre el análisis de ese punto.

Escrito lo anterior, estamos en condiciones de definir la estructura de conocimiento del experto "e" ligado al sistema reticular de la Fig. 3.

Definición 2.2. Diremos que el experto "e" valora con una estructura de tipo-5 si el sistema reticular básico admite un punto de indecisión. Es decir:

$$e(x) \begin{cases} \text{supremo} \\ \text{mas o menos supremo} \\ \text{ínfimo} \end{cases} \quad (3)$$

Esta estructura permite expresar valoraciones como sigue:

$$e(x) = \begin{cases} \lambda < 1 \\ \text{mas o menos } 0 \leq x_1 \leq \lambda < 1 \\ 0 \leq \lambda \leq x_1 \end{cases} \quad (4)$$

Por tanto, podemos afirmar que el experto asociado al sistema reticular Fig. 3 genera vaguedad, y así podemos presentar la idoneidad de la definición (1). O dicho de otra manera: la aparición de estructuras reticulares del tipo Fig.3 aceptan definiciones como la dada en (1). Más aún, puesto que el soporte del experto asociado a la estructura reticular Fig. 3 es el mínimo de los soportes que ofrece lo ordinario, podemos afirmar que $e(x)$ representa la extensión inmediatamente superior a lo ordinario. Esto es: la más sencilla extensión de lo ordinario a la vaguedad.

De interés es también la consideración referida al hecho de poder adaptar nuestras exigencias a través del experto "e". La exigencia puede venir desde el supremo o desde el ínfimo.

Si nuestras exigencias son por el supremo, tendremos que:

$$e(x) = \begin{cases} \begin{cases} \text{supremo-2} \\ \text{mas o menos supremo-2} \\ \text{supremo} \\ \text{mas o menos supremo} \\ \text{ínfimo} \end{cases} \end{cases} \quad (5)$$

Si por el contrario nuestras exigencias vienen por el ínfimo tendremos que:

$$e(x) = \begin{cases} \begin{cases} \text{supremo} \\ \text{mas o menos supremo} \\ \text{ínfimo} \\ \text{mas o menos ínfimo-2} \\ \text{ínfimo-2} \end{cases} \end{cases} \quad (6)$$

Al primero le denominaremos experto de nivel-2 supremo. De igual manera el segundo será de nivel-2 ínfimo.

3 Ejemplo de actuación de esta estructura reticular

En el estudio de los predicados vagos aparece un hecho concreto que la estructura reticular ordinaria no puede analizar. Se trata del término "Muy". Como es conocido "bien" y el "muy bien", en las estructuras reticulares ordinarias aparecen solapadas ya que si el predicado "es bien" $e(x) = 1$, el "no es bien" $e(x) = 0$. Si ahora añadimos el término "Muy", estaremos con "muy bien" cuyo $e(x) = 1$. Obviamente "bien" no es lo mismo que "muy bien". Sin embargo si hacemos intervenir la estructura 5, después de conocer la estructura 4, observamos que

$$e_2(x) = \begin{cases} \begin{cases} \text{supremo-2} \\ \text{mas o menos supremo-2} \\ \text{supremo} \\ \text{mas o menos supremo} \\ \text{ínfimo} \end{cases} \end{cases} = \begin{cases} \begin{cases} r' \\ r \leq k_2 \leq r' \\ r \\ r_e \leq k_1 \leq r \\ r_0 \end{cases} \end{cases}$$

$r', r_1, k_1, k_2, \in]0,1[$ (7)

Por tanto podemos definir el $e_{MB}(x) = e_1(x) \cdot e_2(x) = r \cdot r'$

Si r es muy aproximado a r' consideraremos que: $e_{MB} = r^2$

$$e_{MB}(x) = r^2 \quad (8)$$

Este mismo resultado se ha definido en la lógica borrosa a partir de la iniciativa de lo que se entiende como "sentido común". Por tanto no sería desproporcionado decir que el experto "e" actúa con "sentido común", o, de otra forma, que la estructura Fig. 3 es una estructura reticular apropiada para valoraciones del "sentido común".

4 Análisis del punto P

Hemos afirmado que la estructura reticular representado en la Fig. 3 rompe el concepto de complementario ordinario, y da pie a unos resultados menos categóricos, que no puede ser valorados, sino mas bien intuido. Ese resultado ventajoso lo aporta el punto P.

Dada la importancia de dicho punto en las tareas que nos ocupan, es necesario profundizar algo mas en la reflexión que se hizo anteriormente.

Si nos situamos en la estructura reticular Fig. 4 y designamos por f el elemento de la unidad de la contingencia y por q el elemento de lo necesario, y designamos a f como f -de conocimiento y g , como g -de ignorancia, podemos definir lo siguiente.

Definición 4.1 Diremos que los elementos (f, q) constituyen una complementaridad borrosa si se cumple que:

$$0 < f + g < 1$$

De la definición se deduce que cuando $f \sim g$

$$f + g \cong 2f \cong 2g \cong 0.5 \quad (11)$$

Este punto aparece en la lógica borrosa como punto, cuyo $\mu_P(x) = 1/2$. También es objeto de discusión en la teoría de expertos donde aparecen autores que a la valoración de P dan cero, y otras que mantienen el $1/2$.

Así Fadini y Basilo al definir las relaciones de preferencias como una relación binaria R en

$X \times X$ que verifica:

$$C(x, y) + C(y, x) = 1 \quad \forall (x, y) \in X \times X, X \text{ conjunto de alternativa.} \quad (12)$$

Evidentemente cuando $y = x$, de (11) tenemos que

$$C(x, y) + C(y, x) = 1 \rightarrow 2C(x_1, x) = 1 \rightarrow C(x_1, x) = C(y_1, y) = 0.5 \quad (13)$$

Otros autores, sin embargo a $C(x, x)$ le asocian el valor "0" (Orlovsky, Spillman, Bezdez y otros). Ellos aducen para dar este valor, que como en $e(x, x)$ no se puede distinguir qué variable es la que domina, o qué variable domina, el valor a asignar sería el carácter tanto,

$$C(x_1, x) = 0 \quad \text{con } (x_1, x) \in X \times X \quad (14)$$

Ambas consideraciones se fundamentan en la necesidad del complementario. Veamos cómo actúa la estructura reticular del experto, Fig. 4, sobre cuestiones de dominio de un elemento de unidad sobre el otro.

En primer lugar nuestra escala de valoración $]0, 1[$ se sitúa como indica la Fig. 4.

De esta figura se obtiene que cuando $a \geq 0.5$ la f domina a la g , es decir, el conocimiento domina la ignorancia. Por tanto,

$$f \equiv 0.5 + \alpha; g \equiv 0.5 - \alpha \quad (15)$$

De igual manera se trataría cuando la ignorancia dominara al conocimiento. Estamos pues en un sistema de representación de la unidad dialéctica con el punto 0.5 como punto de cambio. Si ahora situamos el punto P en cero estaríamos en la Fig. 5, en la cual se observa que la estructura reticular no cambia y que sólo

se reduce a un cambio de escala, ya que en estas tareas del conocimiento, los sectores de dominio, que es lo fundamental, no varían. Por tanto, en las formulaciones anteriores no puede haber ningún tipo de controversia ya que la importancia del punto P no está en su carácter cuantitativo, porque, como se indica antes, su grado de conocimiento-ignorancia puede ser alcanzado por un $\inf(\sup)$ o un $\sup(\inf)$, sino por su fundamento cualitativo al impedir P que la estructura reticular pueda reducirse a otra más simple. Por tanto el valor numérico que se da a P , bien por definición, como se hace en (13), bien por conveniencia, mantienen la misma relevancia, por tanto, desde este punto de vista no hay ningún contencioso.

5 Algunas consideraciones sobre P

Es claro que la estructura reticular que venimos analizando constituye una estructura vaga en tanto, como hemos escrito anteriormente, generaliza la estructura reticular de los ordinarios. Hasta este punto no nos habíamos detenido, excepto en la fórmula (9), al estudio de P considerado desde su propia estructura vaga. En lo que sigue, tratamos el comportamiento de P desde los diferentes cortes producidos en la estructura reticular que soporta al experto.

Para ello observemos que hace nuestro experto "e" a la hora de valorar una determinada alternativa.

Sea x una de las alternativas de X , e analiza

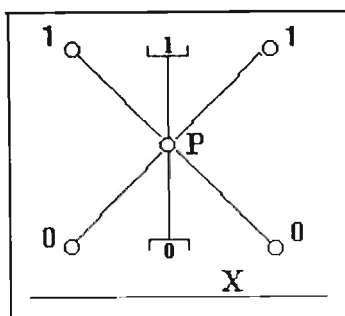


Fig. 4

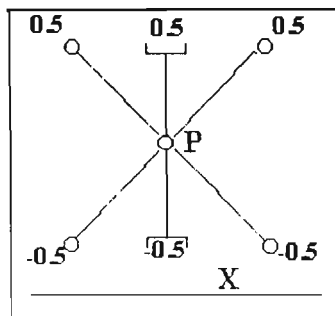


Fig. 5

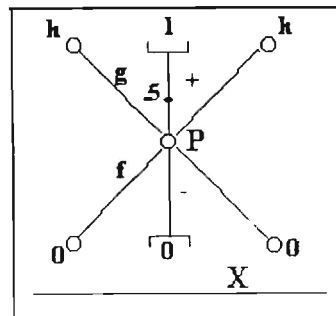


Fig. 6

dicha alternativa

$$x \rightarrow]0,1[\times]0,1[$$

desde su estructura reticular asociada.

La respuesta de tal análisis se da en función de:

$$]0,1[\times]0,1[\rightarrow]0,1[$$

$$(\lambda, g) \rightarrow \alpha < 1$$

donde (f,g) están obligadas por P a que se verifique:

$$f + g \equiv 1$$

Escrito lo anterior podemos afirmar que el mayor interés desde el punto de vista borroso, se sitúa en aquellos valores de (f, g) que no estén muy aproximados al 1 ni al 0 ya que de lo contrario parece prudente situar la valoración del problema en lo ordinario como expresa la Fig. 6.

En dicha figura se observa que hemos producido un h-corte en cada unidad que designamos por h. Esto nos da una estructura reticular cuya diferencia está, en la situación del punto P, que en este caso hace que 0.5 se sitúe en la zona + (designamos con + la zona que predomina la f y con - la que predomina la g), por tanto ya no es un punto de indecisión. De la misma forma la (9) en este sistema reticular, ha de cumplir que,

$$f' + g' \equiv h$$

y por tanto,

$$f' \equiv g = h / 2$$

donde h es el corte dado en ambas categorías en +. En este caso el punto de indecisión P se sitúa por debajo de 0.5. Si por el contrario los cortes se dieran en la zona -, el punto P correspondiente a cada corte estaría ahora en la zona +; Como ambos cortes no pueden alcanzar los pares (0,0) y (1,1) es por lo que hablamos en la expresión (9) de cuasi complementariedad de (f,g), expresión empleada para designar que en los casos más extremos esa estructura puede ser sustituida por una estructura reticular de carácter ordinario. Así mismo, indiquemos como el punto de indecisión P está regulado por h-corte, de tal manera que podemos ajustar la estructura a diversas necesidades, sin que esto suponga cambio de estructura. Sólo con el corte

en la dualidad podemos controlar a P. Por tanto podemos concluir que, el experto definido en (1) lleva asociado un α -corte, cuyo punto de indecisión es $\alpha/2$. Y, por tanto la relación asociada a e(x) es:

$$f + g \equiv \alpha$$

Consideración final

Para finalizar este primer trabajo orientado en lo que venimos describiendo estructuras reticulares del conocimiento, parece que la "herramienta fuzzy" es, en estos momentos, la que ofrece mas posibilidades a la hora de tratar problemas análogos al propuesto en este trabajo. Decir también que en la formulación dada en (1) relativa a la estructura de valoración que usa el experto, no sólo está justificada desde el punto de vista reticular (ver bibliografía adjunta), sino que es una realidad congruente con la vaguedad, impuesta por el sentido común.

Bibliografía

- [1] BEZDEK, J. C., Ed.: *Analysis of Fuzzy Information*, vols. 1, 2 y 3: *Applications in Engineering and Science*. Boca Ratón, FL: CRC, 1987.
- [2] BRACHMAN, R. J.: The basics of knowledge representation and reasoning, *AT&T Tech. J.*, vol. 67, (1988), 25-40.
- [3] DE KLEER, J., y BROWN, J.: A qualitative physics based on confluences. *Artificial Intell.*, vol. 24, (1984), 7-84.
- [4] NÚÑEZ GÓMEZ, A.: *Retículos Completos Completamente Distributivos*. Univ. Alcalá de Henares, 1990. (Tesis Doctoral).
- [5] ZADEH, L. A.: *Fuzzy Sets and Applications* (Selected Papers, edited by R. R. Yager, S. Dvchinnikov, R. M. Tong, H. T. Nguyen), John Wiley, Nueva York, 1987.

Semioraciones y semicadenas (Estudio de algunos aspectos)

Alejandro Sobrino

José Angel Olivas

Area de Lógica y Filosofía de la Ciencia.
Campus Universitario, s.n. E-15706
Santiago de Compostela.
Fax 981. 52.18.18
e-mail: usc.lo.asc@cesga.es

PMM Centro de Estudios
Avda. de Viñuelas, 17. 28760
Tres Cantos. Madrid
Telf. 91. 8035900

Resumen

Este trabajo tiene como objetivo profundizar en un tema de vaguedad gramatical: las semioraciones. Una semioración puede caracterizarse como una oración que no es, desde un punto de vista gramatical, absolutamente correcta, pero tampoco incorrecta (en el sentido de que la comprende perfectamente un hablante oyente de la lengua en la que es proferida). Este estudio indaga:

(a) en una posible caracterización de las semioraciones en el ámbito de la gramática de una lengua, lo que ayuda a profundizar en el concepto de gramaticalidad;

(b) en el correlato de las semioraciones en los lenguajes formales, las semicadenas, que resultan de ejecutar producciones que se desvían de aquello que estrictamente permiten las reglas de una gramática.

Palabras clave: vaguedad, gramáticas fuzzy, semioración, semicadena.

1. Vaguedad lingüística y vaguedad gramatical

El tipo de vaguedad lingüística más usual, y también al que más se alude, es aquel que tiene su origen en la referencia; es decir, en el proceso mediante el cual una palabra significa a un objeto. A este tipo de vaguedad lingüística se le denomina, de manera más específica, vaguedad semántica. Desde este punto de vista, se dice que una palabra es vaga cuando no denota de una manera precisa o inequívoca a su referente; esto es, cuando existen objetos para los cuales es dudoso

que queden significativamente bien representados por esa palabra. Cuando estas palabras lo son de colecciones o clases de objetos, dan nombre a los conjuntos borrosos y propician una semántica formal borrosa [Zadeh, 1971].

Aunque menos corrientes que la vaguedad semántica, existen otros tipos de vaguedad lingüística, como la vaguedad fonológica, presente, p. ej., en el contexto del reconocimiento fónico de los usos dialectales de una lengua; la vaguedad morfológica, acompañante de la deficiente caracterización de clases de palabras en virtud de sus terminaciones; la vaguedad sintáctica, que hace alusión, p. ej., a ausencia de regularidades en el formato de las cláusulas de complemento indirecto; o, por fin, la vaguedad gramatical, presente en clasificaciones gramaticales, donde los datos se agrupan de manera más bien 'squishy' que precisa [Ross, 1972].

Este trabajo estudiará aspectos relacionados con la vaguedad gramatical; en concreto, atenderá al estudio de las semioraciones, oraciones que, desde un punto de vista gramatical, no son totalmente correctas, pero tampoco incorrectas; esto es, oraciones que son más bien gramaticales que no gramaticales o agramaticales, pero, aún así, no totalmente gramaticales. El estudio de las semioraciones tiene una doble relevancia:

(I) En el ámbito del lenguaje natural, para profundizar en el concepto de oración gramaticalmente casi correcta, así como en la comprensión que un hablante oyente de una lengua tiene de estas oraciones a pesar de su relativa malformación. Abundar, en este contexto, en la distinción de gramaticalidad y aceptabilidad.

(II) En el ámbito de los lenguajes formales, para fundamentar el desarrollo de gramáticas

formales que permitan la generación de lenguajes en los que las cadenas se desvían de aquello que estrictamente permiten las reglas de producción de esa gramática.

2. Notas para una teoría de las semioraciones

Un estudio sobre la vaguedad gramatical es, en buena parte, una indagación acerca de si determinados objetos lingüísticos se pueden agrupar en categorías gramaticales con bordes nítidos o si, por el contrario, hay objetos susceptibles de pertenecer a más de una categoría, de ser miembros de una categoría con más claridad que de otra. Si esta última hipótesis fuese correcta, entonces debe admitirse la existencia de categorías gramaticales con bordes difusos, a las que los objetos lingüísticos pertenecerán con un grado; grado que representa a algunas de las propiedades gramaticales de ese hecho.

Un análisis de la gramática inglesa muestra esta particularidad [Ross, 1972]. Representando en cursiva algunas de las categorías centrales de esta gramática, y en escritura normal a las categorías de transición entre las categorías centrales,

Verbo > *participio presente* > *participio pasado* > *Adjetivo* > *preposición* > *nombre adjetival* > *Nombre*. (> indica un orden parcial).

se constata una gradación entre ellas. En las categorías centrales, las oposiciones funcionan de una manera binaria: un hecho lingüístico pertenece o no pertenece a la categoría considerada. Las categorías de transición, en cambio, no tienen este carácter dicotómico, como muestra de modo relevante el caso del participio: si es de presente, supera algunas pruebas gramaticales propias de los verbos, pero si es de pasado, comparte más características con la categoría 'adjetivo' que con la categoría 'verbo'.

El ejemplo visto anima a distinguir en las clasificaciones gramaticales unos elementos centrales, con pertenencia feta, y otros periféricos, con pertenencia gradual a la categoría. Un modelo gramatical que permita esta agrupación elástica, no rígida, de hechos lingüísticos, hará suya nociones como: 'construcciones sintácticas inestables', 'unidades morfológicas poco integradas', 'esquemas morfológicos semiproductivos' o 'se-

mioraciones', entre otras. Discutamos, a continuación, algunas características de las semioraciones.

La conceptualización de una semioración presupone, respecto a la gramaticalidad, una alternativa a la oposición gramatical/no gramatical. En este contexto, el sentido doble de la partícula 'no' permite una triple distinción: (i) oraciones gramaticalmente correctas, (ii) oraciones gramaticalmente no correctas e (iii) oraciones gramaticalmente incorrectas o no gramaticales. A las oraciones que se identifican con la tipología (ii), le denominamos semi-oraciones (S-O). Un ejemplo de semioración es: *Una idea surgió a mí* o *Un bocadillo comí*.

En estos ejemplos se puede observar como la malformación es ocasionada por una deficiente ordenación sintáctica. Pero no todas las semi-oraciones son ocasionadas por problemas sintácticos. Irregularidades en el estrato semántico explican, como se verá más abajo, la ausencia de gramaticalidad completa, p. ej., en la oración *El juicio corre deprisa*.

Una teoría de las semioraciones debe explicar, entonces, como el conocimiento que el hablante tiene de la estructura de una lengua le capacita para entender frases que no son totalmente gramaticales, pero tampoco gramaticales. Respecto al conocimiento que un hablante posee de su gramática y a la capacidad explicativa que se le presupone a ésta para agrupar a todos los hechos lingüísticos, una S-O genera una doble perplejidad:

- Si la S-O no es derivable, entonces la gramática es incompleta,
- Si la S-O es derivable, la gramática falla,

perplejidad que debe resolver, de una manera no dicotómica, una teoría que se ocupe de estas entidades lingüísticas.

Una forma de catalogar oraciones producidas por irregularidades semánticas es distinguir distintos niveles de descripción estructural de una oración [Chomsky, 1964]. Llamémosle a cada uno de estos niveles, formativos (*formatives*) [Bolinger, 1961]. Supongamos, entonces, que disponemos de m niveles de formativos $f^1_1, f^1_2, \dots, f^1_n, \dots, f^2_1, f^2_2, \dots, f^2_n, \dots, f^m_1, f^m_2, \dots, f^m_n$, de tal forma que los subíndices 1, 2, ..., n indican mayor grado de especificidad en la descripción estructural. Sea la oración ya citada,

El juicio corre deprisa (O1)

oración con una buena ordenación sintáctica, pero de la que no se puede decir que sea gramaticalmente tan correcta como la oración *El atleta corre deprisa*. Su anomalía se encuentra, por tanto, en el componente semántico. Veámoslo.

Especifiquemos, a continuación, diferentes niveles de formativos con referencia a los tipos de hechos lingüísticos que describen. Admítase que, en el nivel 1, f^1_1 , representa a palabras. En el nivel 2, f^1_2 a los nombres, f^2_2 a los verbos, f^3_2 a los adverbios y f^4_2 a los artículos. Por último, en el nivel 3 se representan clases de formativos que son subclases de los presentados en 2: f^3_1 representa a los nombres de objetos inanimados, f^3_2 a los verbos de movimiento f^3_3 , a los adverbios de tiempo y f^4_3 a los artículos determinados.

(O1) tiene, por tanto, tres posibles niveles de representación estructural (acordes con los formativos distinguidos). En el nivel 1, f^1_1 , repetido tantas veces como palabras hay. En el nivel 2, $f^4_2 + f^2_1 + f^2_2 + f^3_3$, que equivale a Art + N + V + Adv. Por último, en el nivel 3, $f^3_{4a} + f^3_{1m} + f^3_{2t} + f^3_{3d}$, donde a, m, t y d representan subcategorizaciones (animado, de movimiento...) de las categorías diferenciadas a nivel 2.

Dada una gramática que disponga de estos niveles de descripción, es posible decir algo acerca de la gramaticalidad de (O1):

-El nivel 1 es muy genérico y daría como gramatical a cualquier sucesión de palabras, incluida, por tanto, también (O1).

-Los niveles 2 y 3 constituyen sucesivos refinamientos de 1. El nivel 2 revela una primera estructura de (O1): Art + N + V + Adv. Dado que ésta representa a una buena ordenación sintáctica, es entonces posible, según la gramática del español, generar (O1).

-El nivel 3, en cambio, a pesar de que comparte con el nivel 2 la descripción estructural de (O1), impide la generación de esta oración, porque hay una regla de estructura de constituyentes que indica que un verbo de movimiento sólo acompaña a nombres con la marca 'animado', como ocurre, p. ej., en la oración *El atleta corre deprisa*. El nivel 2 permite generar la semioración (O1), porque su nivel de análisis gramatical no es suficientemente granular. (O1) no contraviene en el nivel 2 la regla que en el nivel 3 transgrede.

Acorde con esta forma de descripción gramatical, se puede decir que una semioración es

una oración que es gramatical (nivel 1 y 2) y no gramatical (nivel 3) a la vez.

No obstante, aún en un mismo nivel de descripción estructural, existen oraciones que no son totalmente gramaticales respecto a otras que sí lo son. Son desviaciones que se caracterizan a nivel puramente sintáctico. Tal es el caso de *Un bocadillo comí* (con referencia a la oración gramatical, *Comí un bocadillo*).

Una forma usual de explicar las S-O desde un punto de vista sintáctico, contempla tres formas de desviación de lo que podría considerarse como oraciones nucleares, obtenidas de derivaciones sintácticas normalizadas. Los tipos más sencillos de variantes son descritos como 'inversión', 'adición' y 'supresión', que se pueden resumir en el siguiente formato, (propio de las reglas de producción):

Formato	Nombre
.....A.....B →B.....A	Inversión
.....A.....B →A.....	Supresión
.....A..... →B.....A	Adición

Ejemplos (en el mismo orden):

Comí un bocadillo → Un bocadillo comí

Comí un bocadillo → Comí bocadillo

Comí en un bar → Te comí en un bar

Obsérvese como en el lado izquierdo de las producciones hay frases totalmente gramaticales, mientras que en el lado derecho hay semioraciones.

De modo paradigmático, la propiedad conocida como supresión (*deletion*), muestra que, para determinado tipo de palabras, no puede darse una regla universal acerca de la posibilidad de suprimirlas o no en las oraciones sin que se altere la gramaticalidad (vaguedad) o el significado (ambigüedad) de la oración. [Ross, 1973]. El siguiente ejemplo muestra casos de ambigüedad gramatical. (Se usará la siguiente notación: (*) por 'puede elidirse' y *() por 'no puede elidirse').

Se estudiará la supresión de tres clases de palabras: preposiciones, adjetivos y nombres.

• *Supresión para la preposición.*

- Tu perfume (*a) me sorprendió.
- Fue sorprendente *(para) mi tu perfume.

• *Supresión para los adjetivos.*

- Sara es como (*a) una muñeca.
- Sara está [cerca] *(a) una muñeca.

• *Supresión para los nombres.*

- Siento [el] (*hecho) [de] que vinieses.
- Siento [el] *(hecho) de ayer.

La propiedad de supresión muestra, por tanto, como algunos elementos gramaticales ('a', 'hecho') pueden ser a veces suprimidos o no suprimidos en las frases sin pérdida de gramaticalidad y como en otras ocasiones esta característica no se mantiene, como muestra la oración del ejemplo, *Siento el hecho de ayer*.

Hasta aquí se han expuesto algunas formas de catalogar o generar semioraciones. No obstante, una teoría de las semioraciones debe tener como objetivo dar una medida del grado de buena formación de una frase, una valoración de su grado de desviación gramatical. Una propuesta plausible debería hacerlo depender del grado de adición/supresión y del grado de inversión del elemento o elementos que originan la desviación gramatical.

El grado de inversión o accesibilidad es el grado de movilidad del elemento o hecho lingüístico dentro de una oración (el elemento considerado se pone en cursiva).

1. i. Es posible que Jorge se matricule de francés *el próximo año*.

ii. *El próximo año* es posible que Jorge se matricule de francés.

2. i. El hombre con bigote *que vi ayer* es el primo del inspector.

ii. El hombre con bigote es el primo del inspector *que vi ayer*.

En 1 se tiene el máximo grado de accesibilidad para la expresión '*el próximo año*'. Por ello, el sentido y referencia de las oraciones del ejemplo (1i) y (1ii) no cambia a pesar de que la ordenación de sus elementos léxicos sea diferente. En 2, en cambio, el grado de accesibilidad de la oración de relativo '*que vi ayer*' es mínimo; al cambiarla de posición se cambia el sentido de la frase, como muestra una comparación entre (2i) y (2ii).

El grado de adición es es grado de posibilidad de incrustación de un elemento gramatical en la frase. P. ej.,

Oración base: Es posible que llueva.

3. (i). Es posible que *mañana* llueva.

(ii). *Mañana* es posible que llueva.

Oración base: Siento que Juan sea así.

4. (i). Siento *siempre* que Juan sea así.

(ii). Siento que Juan *siempre* sea así.

En 4 la posibilidad de que la frase cambie significativamente de sentido alterando la ocurrencia de la palabra en cursiva es máxima, circunstancia que no se mantiene en 3.

Nótese como en estos ejemplos, las palabras en cursiva no originan casos de vaguedad gramatical, sino de ambigüedad gramatical, que es posible reconocer inmediatamente si se genera el árbol de cada una de las oraciones. Alteran simplemente el significado de las oraciones, no su gramaticalidad. Si la incrustación y la movilidad de un elemento originan, no una alteración en el significado (la oración seguiría en ese caso estando acorde con las reglas gramaticalmente establecidas), sino la transgresión de alguna de estas reglas, entonces estaríamos ante un caso de vaguedad gramatical. Cuanto menor sea el grado de incrustación y de movilidad, mayor será la transgresión, de manera que una semioración sería tanto menos gramatical cuanto más se aproxime el producto de estos dos factores a cero.

Por último, una cuestión metodológica: es interesante analizar si la comprensión de una semioración es explicable porque supone una desviación pequeña respecto a una regularidad gramatical o si se debe a nuestra habilidad cognitiva. Es decir, si las semioraciones no son gramaticales, pero son aceptables como oraciones porque permiten comunicar a una buena parte de los hablantes de un lenguaje o si permiten comunicar porque presentan parecido con las oraciones totalmente gramaticales. Mientras 'gramaticalidad' es un término propio del ámbito de la lingüística, 'aceptabilidad' es una palabra de la sociolingüística, de manera que podría pensarse, p. ej., en evaluar una oración como semioración y hasta un grado dependiendo de la aceptabilidad que esa oración tuviese en una comunidad lingüística. Esta evaluación no sería posible para una oración totalmente gramatical, que es tal independientemente de su aceptación. Esto entronca

con una cuestión teórica venerable, pero que no va a ser discutida en este trabajo: Si la gramática debe ser normativa o, por el contrario, debe responder a los hechos lingüísticos tal y como se presentan y, por tanto, también a sus desviaciones.

Lo relatado hasta aquí sólo pretende ser un esbozo inicial de algunos elementos para una teoría de las semioraciones. La complejidad del lenguaje natural hace que la tarea de su reconstrucción sea una labor mucho más amplia que los simples apuntes que se indican en este trabajo.

3. Semigramaticalidad y semicadenas

El objetivo de este apartado es definir casos de vaguedad gramatical en el contexto de un lenguaje generado por una gramática formal.

Dada una gramática formal, $G = \langle N, T, S, P \rangle$, donde:

N es el vocabulario no terminal

T es el vocabulario terminal

S es el símbolo inicial y

P es el conjunto de reglas de producción,

se considera que se pueden producir cadenas imprecisas (semicadenas), de modo análogo a como se generaban semioraciones, de tres formas:

Por,

- **Inversión.** Se genera una sarta análoga a una sarta gramatical, con la diferencia de que tiene sus elementos cambiados de orden (la naturaleza de estos elementos, es decir, si son elementos del vocabulario no terminal o del vocabulario terminal, lo indicará el tipo de gramática que se este definiendo -sensible al contexto, libre de contexto...-)

- **Supresión.** Se genera una cadena análoga a una cadena gramatical, excepto en que tiene algún símbolo terminal menos que la cadena gramatical con la que se compara.

- **Adición.** Se genera una cadena análoga a una cadena gramatical, excepto en que tiene algún símbolo terminal más que la cadena gramatical con la que se compara.

- **Ejemplo.**

Sea $G = \langle \{A, B\}, \{a, b, c\}, A, P \rangle$, donde las P son,

1. $A \rightarrow aB$

2. $B \rightarrow cAa$

3. $A \rightarrow c$

Veamos, a continuación, que producciones son posibles,

Sin desviación

- Reglas de producción
(las mismas)

- Cadenas
– cca (2, 3)
– acca (1, 2, 3)

Con desviación por:

* *Inversión*

- Reglas de producción

1i. $A \rightarrow Ba$

2i. $B \rightarrow aAc$

3i. $A \rightarrow c$

- Cadenas

– acc (2i, 3i)

– acca (1i, 2i, 3i)

* *Supresión*

- Reglas de producción

1s. $(A \rightarrow B_ \mid A \rightarrow _a)$

2s. $(B \rightarrow _Ac \mid B \rightarrow a_c \mid B \rightarrow aA_)$

3s. $A \rightarrow _.$

- Cadenas

– c (2s, 3s)

– a (1s)

– ac (1s, 2s)

– aa (1s, 2s)

* *Adición*

- Reglas de producción

1a. $A \rightarrow \{a \mid b \mid c\}aB \mid aB\{a \mid b \mid c\}$

2a. $B \rightarrow \{a \mid b \mid c\}cAa \mid \{cAa\{a \mid b \mid c\}\}$

3a. $A \rightarrow \{a \mid b \mid c\} \mid c \mid \{c\{a \mid b \mid c\}\}$

- Cadenas (Sólo se ponen algunas)

– acaca (2a, 3a)

– acbca (2a, 3a)

– accca (2a, 3a)

...

Una cuestión que surge de modo natural es si se puede dar algún tipo de cuantificación de la desviación de una cadena y cómo. Una posible (y parcial) respuesta es hacerla depender de algún tipo de razón entre el número de símbolos de la

cadena y los símbolos anómalos que ésta contenga, como p. ej.:

grado de vaguedad de una cadena = $(|ST| - |SA| \div |ST|)_{\text{inversión}} \times (|ST| - |SA| \div |ST|)_{\text{supresión}} \times (|ST| - (|SA| \div |ST|)_{\text{adición}})$, donde:

ST = Número total de símbolos

SA = Número anómalo de símbolos

• Ejemplo:

Si se toma como cadena gramatical (generada sin desviación alguna) a: – cca, el cálculo de la gramaticalidad para las semicadenas (generadas con desviación por 'inversión' y 'adición y supresión', respectivamente): – acc y – accca, se haría de la siguiente forma:

Para – acc sería: $3-2 \div 3 \times 3-0 \div 3 \times 3-0 \div 3 = 1 \div 3 = 0.33$

Para –acaca sería: (considerando que se añada aca y se suprime c): $6-3 \div 6 \times 6-1 \div 6 \times 6-0 \div 6 = 3 \div 6 \times 5 \div 6 = 0.41$

Este esbozo de cálculo cuantitativo de la desviación de una cadena constituye una aproximación muy provisional al tema. Un estudio más elaborado superaría, sin embargo, el alcance de este trabajo.

4. Conclusiones

En esta contribución se indican algunos elementos a tener en cuenta en el posible diseño de una teoría de las semioraciones. En cualquier caso, la complejidad del lenguaje natural hace que la aproximación emprendida aquí tenga un carácter sólo provisional. Se ofrece también una breve nota acerca de cómo la noción de semigramaticalidad puede ser tratada en una gramática formal a través de la noción de semicadena.

Agradecimientos

Los autores agradecen el soporte económico de la D.G.I.C.Y.T (proyecto PB93-0394).

Al libro *Fundamentos de Informática*, de Gregorio Fernández y Fernando Sáez Vacas, le agradecemos que nos haya introducido de una manera tan agradable, y no obstante rigurosa, al ámbito de los lenguajes formales, así como a

otros apasionantes temas teóricos de la computación.

5. Bibliografía

- BOLINGER, D [1961], *Generality, Gradience and All-or-None*. La Haya. Mouton.
- CHOMSKY, N. [1964], 'Degrees of Grammaticalness', en J. A. Fodor y J. J. Katz (eds.), *The Structure on Language. Readings in the Philosophy of Language*. Prentice-Hall, Inc. Englewood Cliffs.
- ROSS, J. R. [1972], 'The Category Squish: Endstation Hauptwort', *Chicago Linguistic Society*, 8, pp. 303-328.
- ROSS, J. R. [1974], 'Nouniness', en Fujimura (ed.), *Three Dimensions of Linguistic Theory*. Tokyo, TEC, pp. 138-257.
- SCHNEIDER, M.; LIM, H. y SHOAF, W. [1993], 'The recognition of imperfect strings generated by fuzzy context sensitive grammars', *Fuzzy Sets and Systems*, 62, pp. 21-29.
- ZADEH, L. A. [1971], 'Quantitative fuzzy semantics', *Information Sciences*, 3, pp. 159-176.

Independencia en la Teoría de la Posibilidad*

Luis M. de Campos, Juan F. Huete

Departamento de Ciencias de la Computación e I.A.

Universidad de Granada

18071-Granada, España

e-mail: lci(jhg)@robinson.ugr.es

Resumen

Este trabajo se centra en el estudio del concepto de independencia para medidas de posibilidad. Se consideran diferentes alternativas, y se analizan sus propiedades. Además se realiza un estudio formal de las condiciones que, de verificarse, conducirían a definiciones apropiadas de independencia para posibilidades.

Palabras clave: independencia condicional, medidas de posibilidad, condicionamiento.

1 Introducción

El concepto de irrelevancia o independencia entre sucesos o variables se ha identificado como uno de los más importantes a tener en cuenta para realizar tareas de razonamiento en dominios de conocimiento amplios y/o complejos de una forma eficiente. La dependencia es una relación que establece un cambio potencial en nuestra creencia actual sobre un suceso o una variable de un dominio de conocimiento, como consecuencia de un cambio específico de nuestro estado de conocimiento sobre otro suceso o variable de dicho dominio. Por tanto, si una variable X es considerada como independiente de otra variable Y , dado un estado de conocimiento Z , entonces nuestra creencia sobre X no variará como consecuencia de adquirir información adicional sobre Y . En otras palabras, la independencia permite modularizar el conocimiento de forma que solo es necesario consultar la información relevante para la cuestión particular en la que estamos interesados, en lugar de tener que examinar una base de conocimiento completa. De este modo, los sistemas de razonamiento que tienen en cuenta consideraciones sobre la independencia pueden ganar en eficiencia. Este es el caso de las redes causales o redes de creencia ([10, 11, 17]), que codifican las relaciones de independencia mediante grafos.

*Este trabajo ha sido financiado por la DGICYT, Proyecto PB92-0939

En el contexto de sistemas que manejan conocimiento con incertidumbre, el concepto de independencia e independencia condicional solo ha sido estudiado profundamente para medidas de probabilidad (ver, por ejemplo, [3, 9, 15]), aunque ya existen algunas aportaciones para otros formalismos de tratamiento de la información con incertidumbre ([1, 2, 14]), así como trabajos que consideran la independencia desde un punto de vista abstracto ([11, 12, 16]).

El objetivo de este trabajo es investigar diversas formas de definir relaciones de independencia en el contexto de la teoría de la posibilidad [18, 6].

El trabajo se estructura en 7 secciones: en la sección 2 se introducen las propiedades que habitualmente se considera que una relación de independencia debería verificar, y que conducen a un esquema de representación mediante grafos. Los conceptos de marginalización y condicionamiento para medidas de posibilidad que se emplean en el trabajo se describen en la sección 3. La sección 4 estudia diferentes definiciones alternativas de independencia para posibilidades, basadas en un concepto de condicionamiento propio de la teoría de Dempster-Shafer. Manteniendo todavía el mismo concepto de condicionamiento, en la sección 5 se hace un breve estudio de las propiedades que una relación de similitud entre posibilidades debería verificar para que su empleo en la definición de independencia resultase adecuado. En la sección 6 se consideran otras definiciones de independencia, pero ahora tomando como base otro concepto de condicionamiento, más usual en un contexto posibilístico. Finalmente, la sección 7 contiene algunas conclusiones.

2 Axiomática para las relaciones de independencia

El estudio en profundidad de la idea de independencia condicional en la teoría de la probabilidad (también en teoría de bases de datos o en teoría de grafos) ha conducido a la identificación de una serie de propiedades que parece razonable exigir a toda relación que intente capturar la noción intuitiva de independencia, y que por tanto se han postulado como

una axiomática para las relaciones de independencia [11]. Si denotamos por $I(X, Y|Z)$ la aserción 'X es independiente de Y dado Z', donde X, Y y Z representan conjuntos (disjuntos) de variables de un determinado dominio de conocimiento, tales propiedades son las siguientes:

- A1 Independencia trivial:
 $I(X, \emptyset|Z)$
- A2 Simetría.
 $I(X, Y|Z) \Rightarrow I(Y, X|Z)$
- A3 Descomposición:
 $I(X, Y \cup W|Z) \Rightarrow I(X, Y|Z)$
- A4 Unión débil.
 $I(X, Y \cup W|Z) \Rightarrow I(X, W|Z \cup Y)$
- A5 Contracción:
 $I(X, Y|Z)$ y $I(X, W|Z \cup Y) \Rightarrow I(X, Y \cup W|Z)$
- A6 Intersección.
 $I(X, Y|Z \cup W)$ y $I(X, W|Z \cup Y) \Rightarrow I(X, Y \cup W|Z)$

donde X, Y, Z, W son conjuntos disjuntos arbitrarios de variables.

De todas estas propiedades, tal vez la menos universal es la de Intersección (por ejemplo, solo las distribuciones de probabilidad estrictamente positivas la verifican, mientras que cualquier distribución de probabilidad verifica las restantes propiedades).

Dependiendo de qué propiedades se verifiquen en una situación concreta, para un modelo de independencia concreto, es posible construir diferentes tipos de grafos, por diferentes medios, que exhiban (a través de los criterios de separación o d-separación en grafos) todas o algunas de las independencias representadas en el modelo de partida ([11, 17]).

3 Medidas de posibilidad. Marginalización y condicionamiento

Si consideramos como referencial un conjunto finito $X = \{x_1, x_2, \dots, x_n\}$ (o bien una variable X cuyos valores posibles son $\{x_1, x_2, \dots, x_n\}$), una medida de posibilidad definida sobre X es una función de conjunto

$$\Pi : \mathcal{P}(X) \rightarrow [0, 1],$$

tal que

- 1 $\Pi(X) = 1$,
- 2 $\Pi(A \cup B) = \max(\Pi(A), \Pi(B))$, $\forall A, B \subseteq X$

Asociada con la medida de posibilidad Π , se define la distribución de posibilidad

$$\pi : X \rightarrow [0, 1]$$

mediante

$$\pi(x) = \Pi(\{x\}), \forall x \in X.$$

La información contenida en la distribución de posibilidad es suficiente para reconstruir la medida de posibilidad, puesto que

$$\forall A \subseteq X, \Pi(A) = \max_{x \in A} \pi(x).$$

Dada una distribución de posibilidad bidimensional π definida sobre el producto cartesiano $X \times Y$, la distribución de posibilidad marginal sobre X, π_X (por simplicidad, y siempre que no de lugar a confusión, eliminaremos el subíndice), se define siempre mediante

$$\pi_X(x) = \max_{y \in Y} \pi(x, y), \forall x \in X.$$

En cambio, el concepto de distribución de posibilidad condicional no es tan universal, existen diferentes alternativas para su definición. Nosotros consideraremos principalmente las siguientes:

1.- La distribución de posibilidad sobre X condicionada al suceso $[Y = y]$, $\pi_d(\cdot|y)$, se define mediante

$$\pi_d(x|y) = \frac{\pi(x, y)}{\pi(y)} = \frac{\pi(x, y)}{\max_{x' \in X} \pi(x', y)}$$

que no es más que una de las definiciones habituales de condicionamiento para medidas de plausibilidad de Dempster-Shafer [4, 13], concretamente la regla de Dempster de condicionamiento, para el caso particular de medidas de posibilidad.

2.- La distribución de posibilidad sobre X condicionada al suceso $[Y = y]$, $\pi_h(\cdot|y)$, se define mediante

$$\pi_h(x|y) = \begin{cases} \pi(x, y) & \text{si } \pi(x, y) < \pi(y) \\ 1 & \text{si } \pi(x, y) = \pi(y) \end{cases}$$

Una versión ligeramente diferente de esta definición fue propuesta en [8], como la solución de la ecuación

$$\pi(x, y) = \pi(x|y) \wedge \pi(y), \forall x \in X.$$

π_h resulta ser la solución menos específica de esta ecuación [7].

Obviamente estas definiciones de distribuciones de posibilidad marginales y condicionales pueden generalizarse de forma inmediata al caso n-dimensional.

4 Definiciones de independencia posibilística

Consideremos un conjunto finito de variables $\mathcal{U}$, sobre el que disponemos de una distribución de posibilidad n-dimensional π . Dados tres subconjuntos (disjuntos) de variables de $\mathcal{U}$, X, Y y Z, representaremos mediante $I(X, Y|Z)$ la sentencia X es condicionalmente independiente de Y dado Z, para el modelo de

independencia asociado a la distribución π . Denotaremos por x, y, z los valores genéricos que las variables respectivas pueden tomar. Los valores de, por ejemplo $Y \cup Z$, se denotarán simplemente mediante yz .

La forma más obvia de definir la independencia condicional es proceder de forma similar al caso probabilista, es decir, mediante la factorización de la distribución conjunta de X, Y, Z . Esta idea es la considerada en [14], en el contexto más general de los sistemas basados en valuaciones. En nuestro caso, esta idea, empleando el condicionamiento de Dempster, da lugar a:

$$I(X, Y|Z) \Leftrightarrow \pi_d(x|yz) = \pi_d(x|z),$$

es decir, el conocimiento sobre el valor de la variable Y no modifica nuestra información sobre los valores de X , una vez que se conoce el valor de la variable Z ; x, y, z son valores arbitrarios de las variables X, Y, Z , con la única restricción de que las medidas condicionales implicadas estén definidas, es decir, que $\pi(yz) \neq 0$.

Esta definición cumple las propiedades A1-A5, y si la distribución de posibilidad π es estrictamente positiva, también cumple A6.

El problema con esta definición es, en nuestra opinión, que es excesivamente estricta, habida cuenta que se exige la igualdad de las distribuciones, y que éstas representan un conocimiento impreciso. El problema se agrava en el caso (habitual) de que las distribuciones tengan que ser estimadas, a partir de datos o de juicios humanos.

La definición anterior supone que nuestra información sobre X permanece inalterada al condicionar a Y . Una idea alternativa podría ser considerar que se da la independencia cuando al condicionar a Y no se gana información adicional sobre los valores de X (pero se podría llegar a perder). El concepto de que una distribución de posibilidad sea más o menos informativa que otra es capturado adecuadamente por la siguiente definición de inclusión [5]: Dadas dos distribuciones de posibilidad π y π' , se dice que π' está incluida en (proporciona menos información que) π si y solo si $\pi(x) \leq \pi'(x) \forall x$.

Empleando la relación de inclusión entre posibilidades, la idea anterior se puede expresar mediante

$$I(X, Y|Z) \Leftrightarrow \pi_d(x|yz) \geq \pi_d(x|z)$$

Desgraciadamente esta definición no cumple la propiedad A4 (unión débil), aunque sí cumple A1-A3 y A5. El problema creemos que se encuentra en el hecho de que no se ha llevado hasta sus últimas consecuencias la idea de independencia como no ganancia de información: si al condicionar se pierde información, parece más conveniente 'quedarnos como estábamos'. Esto puede ser debatible, pero representa una especie de regla por defecto: si para un contexto muy específico se carece de información, emplear la información disponible en un contexto menos específico. En términos prácticos, esta idea implica un cambio en la definición de condicionamiento:

$$\pi_d(x|y) = \begin{cases} \pi(x) & \text{si } \pi_d(x|y) \geq \pi(x) \forall x \\ \pi_d(x|y) & \text{si } \exists x' \text{ tal que } \pi_d(x'|y) < \pi(x') \end{cases}$$

Empleando este condicionamiento, la nueva definición de independencia es

$$I(X, Y|Z) \Leftrightarrow \pi_d(x|yz) = \pi_d(x|z).$$

Esta definición satisface las propiedades A1 y A3-A6 (incluso para distribuciones no estrictamente positivas se cumple la propiedad de intersección). No verifica la simetría (pero ésta es una propiedad fácil de recuperar: basta definir una nueva relación $I'(\cdot, \cdot)$ mediante $I'(X, Y|Z) \Leftrightarrow I(X, Y|Z)$ y $I(Y, X|Z)$, para obtener la simetría).

Otro enfoque muy diferente para definir el concepto de independencia consiste en considerar que las distribuciones $\pi_d(x|yz)$ y $\pi_d(x|z)$ sean similares en algún sentido. Así pues, si $\simeq$ es una relación en el conjunto de las distribuciones de posibilidad, se definiría la independencia mediante

$$I(X, Y|Z) \Leftrightarrow \pi_d(x|yz) \simeq \pi_d(x|z).$$

Existen diversas alternativas para definir la relación $\simeq$; veamos algunas de ellas:

Si pensamos que una distribución de posibilidad esencialmente establece una ordenación entre los valores que la variable sobre la que la distribución proporciona información puede tomar, considerando que la cuantificación de los grados de posibilidad es secundaria, entonces podríamos decir que dos distribuciones de posibilidad son similares si establecen la misma ordenación. Más formalmente, podemos definir la relación $\simeq$ mediante

$$\pi \simeq \pi' \Leftrightarrow \forall x, x' [\pi(x) < \pi(x') \Leftrightarrow \pi'(x) < \pi'(x')].$$

La independencia definida mediante esta relación (que podríamos llamar de 'isoordenación') verifica las propiedades A1 y A3-A5, y también A6 para distribuciones estrictamente positivas; sin embargo no verifica la simetría.

Otra alternativa es hablar de similaridad entre distribuciones cuando los grados de posibilidad de las distribuciones para cada valor sean parecidos. Más concretamente, sea m un entero positivo cualquiera, y sean $\{\alpha_k\}_{k=0, \dots, m}$ tales que $\alpha_0 < \alpha_1 < \dots < \alpha_{m-1} < \alpha_m$, con $\alpha_0 = 0$ y $\alpha_m = 1$. Si denotamos $I_k = [\alpha_{k-1}, \alpha_k]$, $k = 1, \dots, m-1$, y $I_m = [\alpha_{m-1}, \alpha_m]$, entonces definimos la relación $\simeq$ mediante

$$\pi \simeq \pi' \Leftrightarrow \forall x \exists k \in \{1, \dots, m\} \text{ tal que } \pi(x), \pi'(x) \in I_k$$

Esta definición es equivalente a la siguiente, establecida en términos de los α -cortes de la distribución

$$\pi \simeq \pi' \Leftrightarrow C(\pi, \alpha_k) = C(\pi', \alpha_k) \forall k = 1, \dots, m-1$$

donde $C(\pi, \alpha) = \{x | \pi(x) \geq \alpha\}$. La idea es también equivalente a discretizar el intervalo $[0, 1]$ (el rango de variación de las distribuciones), y decir que dos distribuciones son similares si sus respectivas discretizaciones coinciden. La relación de independencia definida en estos términos verifica las propiedades A1 y A3-A6 (aunque A6 solo para distribuciones estrictamente positivas), y de nuevo falla la simetría.

Por último, una tercera alternativa para definir $\simeq$ está basada en la siguiente idea: considerar un umbral α_0 , a partir del cual se considera interesante discriminar entre los grados de posibilidad de dos distribuciones, de forma que los valores cuyos grados de posibilidad sean inferiores al umbral no se consideran relevantes. En términos de los α -cortes de las distribuciones, esta relación $\simeq$ se expresaría de la siguiente forma:

$$\pi \simeq \pi' \Leftrightarrow C(\pi, \alpha) = C(\pi', \alpha) \quad \forall \alpha \geq \alpha_0,$$

que resulta equivalente a

$$\begin{aligned} \pi \simeq \pi' \Leftrightarrow & \quad C(\pi, \alpha_0) = C(\pi', \alpha_0) \\ & \text{y} \quad \pi(x) = \pi'(x) \quad \forall x \in C(\pi, \alpha_0) \end{aligned}$$

De nuevo, la independencia definida en términos de esta relación $\simeq$ verifica las propiedades A1, A3- A5 (y A6 para distribuciones estrictamente positivas), y no cumple A2.

5 Condiciones suficientes para los axiomas de independencia

En esta sección continuamos con la idea de definir la independencia $I(X, Y|Z)$ en términos de una relación de similaridad $\simeq$ entre las distribuciones de posibilidad condicionadas $\pi_d(x|yz)$ y $\pi_d(x|z)$, pero pretendemos estudiar qué tipo de propiedades de $\simeq$ son suficientes para garantizar que la relación de independencia así definida cumpla algunas de las propiedades o axiomas considerados anteriormente.

En primer lugar, es obvio que la Independencia trivial (A1) se cumplirá si $\simeq$ es una relación reflexiva. También es evidente que la transitividad de $\simeq$ garantiza la propiedad A5 (Contracción). Si además, $\simeq$ es simétrica, entonces puede deducirse que se verifica A3 (Descomposición) si y solo si se verifica A4 (Unión débil). Por tanto, parece que las relaciones de equivalencia $\simeq$ son buenas candidatas para definir la independencia.

Una condición suficiente para que se verifique A3 es que $\simeq$ cumpla la siguiente propiedad:

Sea $\{\pi_s\}$ una familia de distribuciones de posibilidad, y $\{\lambda_s\}$ una familia de números reales positivos. Entonces

$$\frac{\pi_s}{\lambda_s} \simeq \pi \Rightarrow \frac{\max_s \pi_s}{\max_s \lambda_s} \simeq \pi \quad (1)$$

Además, en caso de que las distribuciones sean estrictamente positivas, el cumplimiento de la propiedad anterior también garantiza que se verifica A6.

Por tanto, toda relación de independencia posibilística definida en términos de una relación $\simeq$ que sea de equivalencia y verifique (1) cumple las propiedades de Independencia trivial, Descomposición, Unión débil, Contracción e Intersección (aunque esta última solo para distribuciones estrictamente positivas). La única propiedad que queda fuera de este contexto es la Simetría, lo cual resulta curioso, puesto que es una

de las propiedades de la independencia aparentemente más intuitiva. Evidentemente, las tres relaciones $\simeq$ definidas en la sección anterior son de equivalencia y cumplen (1).

6 Cambiando el operador de condicionamiento

En esta sección emplearemos π_h como operador de condicionamiento en lugar de π_d , que ha sido el utilizado en las dos secciones anteriores. Continuando con la idea de independencia como no ganancia de información tras condicionar, podemos entonces definir

$$I(X, Y|Z) \Leftrightarrow \pi_h(x|yz) \geq \pi_h(x|z).$$

En este caso, a diferencia de lo que ocurría con π_d , la definición anterior verifica las propiedades A1- A5, pero en general no verifica A6 (Intersección). Además, esta definición resulta ser equivalente a la siguiente:

$$I(X, Y|Z) \Leftrightarrow \pi(xyz) = \pi(xz) \wedge \pi(yz).$$

Si consideramos el caso particular de independencia marginal (es decir, cuando $Z = \emptyset$), entonces obtenemos:

$$I(X, Y|\emptyset) \Leftrightarrow \pi(xy) = \pi(x) \wedge \pi(y),$$

es decir, el conocido concepto de no interactividad para medidas de posibilidad o conjuntos difusos. Por tanto, nuestro concepto de independencia es lo que podríamos llamar 'no interactividad condicional'.

Obsérvese que para este tipo de condicionamiento la única operación necesaria es la de comparación. Por tanto podríamos fácilmente considerar distribuciones de posibilidad valuadas en conjuntos diferentes del intervalo $[0, 1]$: bastaría usar un conjunto $(\mathcal{L}, \preceq)$, donde

$$\mathcal{L} = \{L_0, L_1, \dots, L_n\}$$

y

$$L_0 \preceq L_1 \preceq \dots \preceq L_n$$

es decir, un conjunto totalmente ordenado (por ejemplo, etiquetas lingüísticas), y definir medidas de posibilidad mediante

$$\Pi : \mathcal{P}(X) \rightarrow \mathcal{L}$$

verificando:

1. $\Pi(X) \preceq L_n$,
2. $\Pi(A \cup B) = \vee_{\preceq}(\Pi(A), \Pi(B))$, $\forall A, B \subseteq X$,

donde $\vee_{\preceq}$ es el operador máximo (supremo) asociado al orden $\preceq$.

En esas condiciones podemos definir el condicionamiento y la independencia de exactamente la misma forma que antes, obteniendo las mismas propiedades.

Otra vía diferente de extensión (análoga a la que hemos considerado en la sección 5 para el otro condicionamiento) consistiría en considerar una relación $\approx$

sobre el conjunto de las medidas de posibilidad, y definir la independencia mediante

$$I(X, Y|Z) \Leftrightarrow \pi(xyz) \approx \pi(xz) \wedge \pi(yz).$$

En este caso se puede probar que para que esta definición de independencia cumpla los axiomas A1–A5, es condición suficiente que $\approx$ sea una relación de equivalencia compatible con la marginalización y la combinación de distribuciones de posibilidad (empleando el operador mínimo como operador de combinación).

Por otra parte si, continuando el paralelismo con la sección 4, definimos π_{h_c} mediante

$$\pi_{h_c}(x|y) = \begin{cases} \pi(x) & \text{si } \pi_h(x|y) \geq \pi(x) \forall x \\ \pi_h(x|y) & \text{si } \exists x' \text{ tal que } \pi_h(x'|y) < \pi(x') \end{cases}$$

y la relación de independencia

$$I(X, Y|Z) \Leftrightarrow \pi_{h_c}(x|yz) = \pi_{h_c}(x|z),$$

entonces se cumplen todas las propiedades A1–A6 excepto la Simetría. Esta relación resulta ser más restrictiva que la no interactividad condicional, pues si se verifica la independencia con esta definición, también se verifica la no interactividad condicional, pero el recíproco no es cierto.

7 Conclusiones

El concepto de independencia, al igual que para otros formalismos de tratamiento de información con incertidumbre, es también importante para la teoría de la posibilidad. Sin embargo, y salvo contadas excepciones, creemos que ha sido un tema no muy estudiado hasta la fecha. En este trabajo hemos propuesto diversas alternativas para la definición de independencia posibilística, basadas casi todas ellas en la idea de condicionamiento, aunque empleando diferentes formulaciones. Además hemos estudiado sus características, contrastándolas con un conjunto de axiomas bien conocidos, que tratan de capturar la idea intuitiva de independencia condicional. Es evidente que queda mucho por hacer en este tema, desde profundizar mucho más en las cuestiones aquí estudiadas (como por ejemplo, las consecuencias de una definición no simétrica de independencia), hasta considerar otros puntos de vista, no necesariamente basados en la idea de condicionamiento. También es importante abordar el problema de la construcción de grafos de independencias para posibilidades, y los mecanismos de propagación en tales estructuras.

Agradecimientos

Los autores están agradecidos a uno de los revisores anónimos por sus interesantes sugerencias, que han servido para mejorar este trabajo.

Bibliografía

- [1] CAMPOS, L.M. DE, HUETE, J.F. Independence concepts in upper and lower probabilities, Uncertainty in Intelligent Systems, eds. B. Bouchon-Meunier, L. Valverde, R.R. Yager, North-Holland (1993), 49–59.
- [2] CAMPOS, L.M. DE, HUETE, J.F. Learning non probabilistic belief networks, Symbolic and Quantitative Approaches to Reasoning and Uncertainty, Lecture Notes in Computer Science 747, eds. M. Clarke, R. Kruse, S. Moral, Springer Verlag (1993), 57–64.
- [3] DAWID, A.P. Conditional independence in statistical theory, J.R. Statist. Soc. Ser. B 41 (1979), 1–31.
- [4] DEMPSTER, A.P. Upper and lower probabilities induced by a multivalued mapping, Ann. Math. Stat. 38 (1967), 325–339.
- [5] DUBOIS, D., PRADE, H. A set-theoretic view of belief functions, International Journal of General Systems 12 (1986), 193–226.
- [6] DUBOIS, D., PRADE, H. Possibility Theory: An approach to computerized processing of uncertainty, Plenum Press, 1988.
- [7] DUBOIS, D., PRADE, H. Belief revision and updates in numerical formalisms—An overview, with new results for the possibilistic framework, Proceedings of the 13th IJCAI, Morgan and Kaufmann (1993), 620–625.
- [8] HISDAL, E. Conditional possibilities, independence and noninteraction, Fuzzy Sets and Systems 1 (1978), 283–297.
- [9] LAURITZEN, S.L., DAWID, A.P., LARSEN, B.N., LEIMER, H.-G. Independence properties of directed Markov fields, Networks 20 (1990), 491–505.
- [10] PEARL, J. Fusion, propagation and structuring in belief networks, Artificial Intelligence 29 (1986), 241–288.
- [11] PEARL, J. Probabilistic reasoning in intelligent systems: Networks of plausible inference, Morgan and Kaufmann, 1988.
- [12] PEARL, J., GEIGER, D., VERMA, T. Conditional independence and its representation, Kybernetika 25 (1989), 33–44.
- [13] SHAFER, G. A mathematical theory of evidence, Princeton University Press, 1976.
- [14] SHENOY, P.P. Conditional independence in uncertainty theories, Proceedings of the Eighth Conference on Uncertainty in Artificial Intelligence, eds. D. Dubois, M.P. Wellman, B. D'Ambrosio, P. Smets, Morgan and Kaufmann (1992), 284–291.
- [15] SPOHN, W. Stochastic independence, causal independence and shieldability, Journal of Philosophical Logic 9 (1980), 73–99.

- [16] STUDENÝ, M. Attempts at axiomatic description of conditional independence, *Kybernetika* 25, Supplement to n. 3 (1989), 72-79.
- [17] VERMA, T., PEARL, J. Causal networks: Semantics and expressiveness, *Uncertainty in Artificial Intelligence 4*, eds. R.D. Shachter, T.S. Levitt, L.N. Kanal, J.F. Lemmer, North-Holland (1990), 69-76.
- [18] ZADEH, L.A. Fuzzy sets as a basis for a theory of possibility, *Fuzzy Sets and Systems* 1 (1978), 3-28.

An Approach to First-order Multiple-valued Logic Programming *

Gonzalo Escalada-Imaz¹

Felip Manyà^{1,2}

⁽¹⁾Institut d'Investigació en Intel·ligència Artificial
Centre d'Estudis Avançats de Blanes (IIIA-CSIC)
C. Santa Bàrbara s/n, 17300 Blanes, Spain

⁽²⁾Dpt. Informàtica i E. Industrial
Universitat de Lleida
Ap. Correus 471, 25080 Lleida, Spain

Abstract

We present a first order interpreter for multiple-valued logic programming. This interpreter has been designed in such a way that the computation of unifiers and exploration of the search space are carried out efficiently. Another feature of this interpreter is that it can deal with a parametric family of multiple-valued logics since different interpretation functions for conjunction and implication can be used.

1 Introduction

Logic programming based on uncertainty logics (possibilistic, probabilistic, evidential), as well as on fuzzy or multiple-valued logics is a growing area of interest in Artificial Intelligence [1, 3, 6, 7, 9, 10, 14].

In this paper we present a first order interpreter for multiple-valued logic programming. This interpreter has been designed in such a way that the computation of unifiers and exploration of the search space are carried out efficiently. Another feature of this interpreter is that it can deal with a parametric family of multiple-valued logics since different interpretation functions for conjunction and implication can be used. In our proposal terms are only constants and variables and we consider as future work to extend the interpreter to handle functions too.

Section 2 defines the family of multiple-valued logics that the interpreter can deal with giving their syntax, semantics and logical inference. Section 3 presents the interpreter beginning from a simple version and, after discussing its inefficiencies concerning the computation of unifiers and the exploration of the search space, an improved version of the interpreter is described. Finally, some conclusions are given in the last section.

2 First order multiple-valued logic

The logics considered in this paper have the following features: the set of truth values is the real interval $[0, 1]$; the interpretation function for the negation connective is $N(a) = 1 - a$, for the conjunction connective can be any T-norm and for the implication connective is defined by residuation wrt the T-norm. For each

T-norm, a different logic is obtained and so the interpreter deals with a parametric family of multiple-valued logics [6].

We present the syntax and semantics for multiple-valued definite clauses whose terms are constants or variables. We also give a sound and complete calculus.

Definition 1 (term) A term is a constant or a variable.

Definition 2 (mv-atomic formula) If p is a predicate symbol and $t_1, \dots, t_n$ are terms then $p(t_1, \dots, t_n)$ is an atomic formula. If $p(t_1, \dots, t_n)$ is an atomic formula and $\alpha \in [0, 1]$ then $(p(t_1, \dots, t_n); \alpha)$ is a mv-atomic formula.

A program P in this language is formed by a finite set of multiple-valued clauses (C, ρ) divided into facts and rules which are defined as follows:

- **fact:** it is a ground mv-atomic formula.
- **rule:** it is a couple $(q \leftarrow p_1, \dots, p_k; \alpha)$ where $q, p_1, \dots, p_k$ are atomic formulae and $\alpha \in [0, 1]$.

Definition 3 (interpretation) An interpretation $\mathcal{I}$ of a first-order multiple-valued language is a non empty domain D and a mapping I that associates:

- each constant symbol a with an element $I(a) \in D$;
- each predicate symbol p with a mapping $I(p) : D^n \rightarrow [0, 1]$, where n is the arity of p .

Definition 4 (valuation) A valuation $\varphi_{\mathcal{I}}$ with respect to an interpretation $\mathcal{I}$ is a mapping from variables to the domain of $\mathcal{I}$.

Definition 5 (semantics of terms) Let $\mathcal{I}$ be an interpretation, $\varphi_{\mathcal{I}}$ a valuation and t a term. The meaning $\varphi_{\mathcal{I}}^*(t)$ of t is an element of the domain of $\mathcal{I}$ defined as follows:

- if t is a constant symbol a then $\varphi_{\mathcal{I}}^*(t) = I(a)$;
- if t is a variable symbol x then $\varphi_{\mathcal{I}}^*(t) = \varphi_{\mathcal{I}}(x)$;

In the following, $\mathcal{I} \models (S; \alpha)$ denotes that the interpretation $\mathcal{I}$ is a model of the sentence $(S; \alpha)$

*Research partially supported by the project PB91-0334-C03-03 funded by the DGICYT

Definition 6 (semantics of wff's) Let $\mathcal{I}$ be an interpretation and $\varphi_{\mathcal{I}}$ a valuation. The meaning of a wff is defined as follows:

- $\mathcal{I} \models (p(t_1, \dots, t_n); \alpha)$ iff $I(p)(\varphi_{\mathcal{I}}^*(t_1), \dots, \varphi_{\mathcal{I}}^*(t_n)) \geq \alpha$ for all the possible valuations;
- $\mathcal{I} \models (p(t_1, \dots, t_n) \leftarrow q(t'_1, \dots, t'_m); \alpha)$ iff $I_T(I(p)(\varphi_{\mathcal{I}}^*(t_1), \dots, \varphi_{\mathcal{I}}^*(t_n)), I(q)(\varphi_{\mathcal{I}}^*(t'_1), \dots, \varphi_{\mathcal{I}}^*(t'_m))) \geq \alpha$ for all the possible valuations;
- $\mathcal{I} \models (q(t_1, \dots, t_k) \leftarrow p_1(t_1^1, \dots, t_{m_1}^1), \dots, p_n(t_1^n, \dots, t_{m_n}^n); \alpha)$ iff $I_T(I^*(p_1), I^*(p_2), \dots, I^*(p_n), \dots), I^*(q)) \geq \alpha$ where $I^*(p_i) = (I(p_i)(\varphi_{\mathcal{I}}^*(t_1^i), \dots, \varphi_{\mathcal{I}}^*(t_{m_i}^i)))$

Definition 7 (logical consequence) Let P be a program and (p, α) a fact. (p, α) is a logical consequence of P , noted $P \models (p, \alpha)$, iff all the models of P are also models of (p, α) .

Logical inference: the interpreter is based on two inference rules, the **multiple valued modus ponens**

$$\frac{(p \leftarrow p_1, \dots, p_n; \alpha) \quad (p'_1; \beta_1), \dots, (p'_n; \beta_n) \quad p_1\sigma = p'_1\sigma, \dots, p_n\sigma = p'_n\sigma}{(p; T(\alpha, \beta_1, \dots, \beta_n))}$$

where σ is an mgu of p_i and p'_i , $1 \leq i \leq n$, and the **weakening**,

$$\frac{(p; \alpha), \alpha > \alpha'}{(p; \alpha')}$$

and the axiom $\vdash (p; 0)$.

Theorem 1 (soundness and completeness) Let P be a program and q a goal. $P \models (q; \alpha)$ iff $P \vdash (q; \alpha)$, where the $\vdash$ relation is formed by the previous inference rules and axiom.

3 A first order interpreter

In this section we design an interpreter based on the previous calculus. Given a program P and a goal q , the aim of the interpreter is to find a substitution θ and a truth value α such that $P \models (q\theta; \alpha)$ and $\alpha \geq \beta$, β being a given truth threshold.

First, we describe an interpreter that can be seen as an adaptation of the Prolog procedure to the multiple-valued case. Second, we describe a more efficient interpreter and discuss its improvements.

Data structures: for each program clause C we use the following data structure:

- $\text{consequent}(C) = p$, where p is the consequent of a mv-rule $(p \leftarrow p_1, \dots, p_n; \alpha)$ or the proposition of a mv-fact $(p; \alpha)$.
- $\text{antecedents}(C) = \{p_1, \dots, p_n\}$ is the set of antecedents of a mv-rule $(p \leftarrow p_1, \dots, p_n; \alpha)$ or the empty set for mv-facts.
- $\text{val}(C) = \alpha$ represents the truth value attached to C .

The formal description of the algorithm to interpret logic programs is the following:

```

INTERPRETER(goal,  $\theta$ ,  $\alpha$ )
  if goal = {} then halt( $\{\theta, \alpha\}$ )
  (assume that goal has the form goal =  $\{q_1, q_2, \dots, q_n\}$ )
  for each  $C \in P$  such that
     $\sigma = \text{mgu}(q_1, \text{consequent}(C))$  and  $\sigma \neq \text{fail}$  do:
       $\alpha' \leftarrow T(\alpha, \text{val}(C)); \theta' \leftarrow \theta \cdot \sigma;$ 
      goal'  $\leftarrow (\text{antecedents}(C) \cup \{q_2, \dots, q_n\})\sigma$ 
      if  $\alpha' \geq \beta$  then INTERPRETER(goal',  $\theta'$ ,  $\alpha'$ )
  return(fail)

```

Given a program P , a threshold β and a goal q , the function $\text{INTERPRETER}(\text{goal}, \theta, \alpha)$ looks for a substitution θ and a truth value α such that $P \models (q\theta; \alpha)$ and $\alpha \geq \beta$. Initially, the interpreter is called with the following parameters $\text{INTERPRETER}(\{q\}, \{\}, 1)$.

The function INTERPRETER runs in a depth-first manner applying the modus ponens inference rule in a reverse order, i.e. a literal p is deduced if all the antecedent literals of a rule whose consequent unifies with p are deduced. Thus, in every state $\text{INTERPRETER}(\{q_1, \dots, q_n\}, \theta, \alpha)$, the function selects a program clause $(C, \text{val}(C))$ such that $\text{consequent}(C)$ unifies, via an mgu σ , with the leftmost atom q_1 of the goal, and new state $\text{INTERPRETER}(\text{antecedents}(C) \cup \{q_2, \dots, q_n\}\sigma, \theta \cdot \sigma, T(\alpha, \text{val}(C)))$ is obtained. If either no program clause unifies with the leftmost atom of the goal or the truth value attached to the goal is smaller than β then the function backtracks (from the current state, only states with truth value smaller than α can be reached). This process continues until a solution state with an empty set of literals is reached.

Example

Let P be the following logic program:

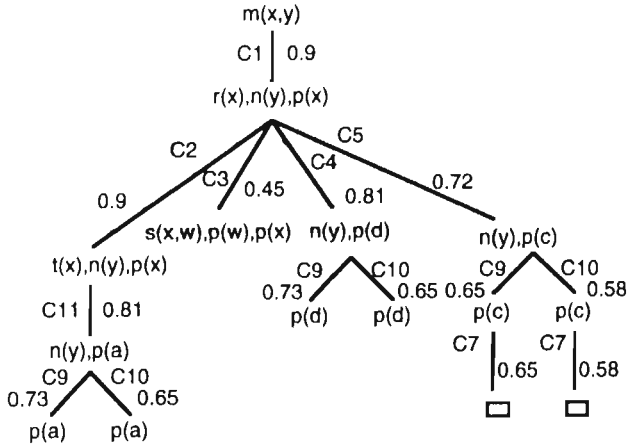
```

C1 : (m(x, y)  $\leftarrow$  r(x), n(y), p(x); 0.9)
C2 : (r(z)  $\leftarrow$  t(z); 1)
C3 : (r(u)  $\leftarrow$  s(u, w), p(w); 0.5)
C4 : (r(d); 0.9)
C5 : (r(c); 0.8)
C6 : (s(b, d); 0.1)
C7 : (p(c); 1)
C8 : (p(b); 1)
C9 : (n(b); 0.9)
C10 : (n(c); 0.8)
C11 : (t(a); 0.9)

```

The T-norm used in the following examples is $T(a, b) = a * b$. Assume that we would need to determine a solution of the problem $P \models (m(x, y), 0.5)$. Then making the call $\text{INTERPRETER}(\{m(x, y)\}, \{\}, 1)$, the interpreter returns the substitution $\{x/c, y/b\}$ and the truth value 0.66. Nevertheless, this is not the only solution. This is the first solution found when the search space is traversed in a depth first strategy. The other solution is the substitution $\{x/c, y/c\}$ and its truth value 0.58. Figure 1 shows the search space corresponding to this problem $P \models (m(x, y), 0.5)$.

The definition of INTERPRETER can be modified to obtain all the solutions by changing the first line by "if goal = {} then print($\{\theta, \alpha\}$); return" and the last line by "return"

Figure 1: search space for $P \models (m(x, y); 0.5)$

The interpreter described above is based on unification steps, depth first search and chronological backtracking. To see better how it works, the different states successively reached by the interpreter are displayed below.

- $e_1 = (\{m(x, y)\}, \{\}, 1,)$
- $e_2 = (\{r(x), n(y), p(x)\}, \{\}, 0.9)$
- $e_3 = (\{t(x), n(y), p(x)\}, \{z/x\}, 0.9)$
- $e_4 = (\{n(y), p(a)\}, \{z/a, x/a\}, 0.81)$
- $e_5 = (\{p(a)\}, \{z/a, x/a, y/b\}, 0.73) \text{ fail}$
- $e_6 = (\{p(a)\}, \{z/a, x/a, y/c\}, 0.65) \text{ fail}$
- $e_7 = (\{s(x, w), p(w), p(x)\}, \{u/x\}, 0.45) \text{ fail}$
- $e_8 = (\{n(y), p(d)\}, \{x/d\}, 0.81)$
- $e_9 = (\{p(d)\}, \{x/d, y/b\}, 0.73) \text{ fail}$
- $e_{10} = (\{p(d)\}, \{x/d, y/c\}, 0.65) \text{ fail}$
- $e_{11} = (\{n(y), p(c)\}, \{x/c\}, 0.72)$
- $e_{12} = (\{p(c)\}, \{x/c, y/b\}, 0.65)$
- $e_{13} = (\{\}, \{x/c, y/b\}, 0.65)$

The initial situation is represented by state e_1 . The interpreter goes on working until it reaches state e_5 corresponding to the computations carried out to traverse the leftmost path of the search space. Since no solution can be found it backtracks to state e_4 and generates state e_6 , but as this new path is also a failed one it backtracks to state e_2 and generates state e_7 . As the truth value of state e_7 is smaller than 0.5 it backtracks again to state e_2 and generates states e_8 and e_9 . State e_9 represents a failed path and then the interpreter backtracks to e_8 and generates state e_{10} which is also a failed path and the interpreter backtracks to e_2 . From e_2 are generated states e_{11} , e_{12} and e_{13} . State e_{13} is a successful path since the goal is empty. Thus the computed answer for the initial goal is $\{x/c, y/b\}$ and its truth value is 0.65.

The previous interpreter has been defined as an adaptation of the Prolog procedure, that is, based on unification steps, depth first search and chronological backtracking. This interpreter is particularly inefficient when it has to run on multiple-valued logic programs instead of classical first order logic programs.

In the former case is not enough to find a branch ending with the empty clause, it must also have a truth degree α greater than a given threshold β . Therefore, the search space that has to be examined is always bigger or equal than the search space explored in the classical case.

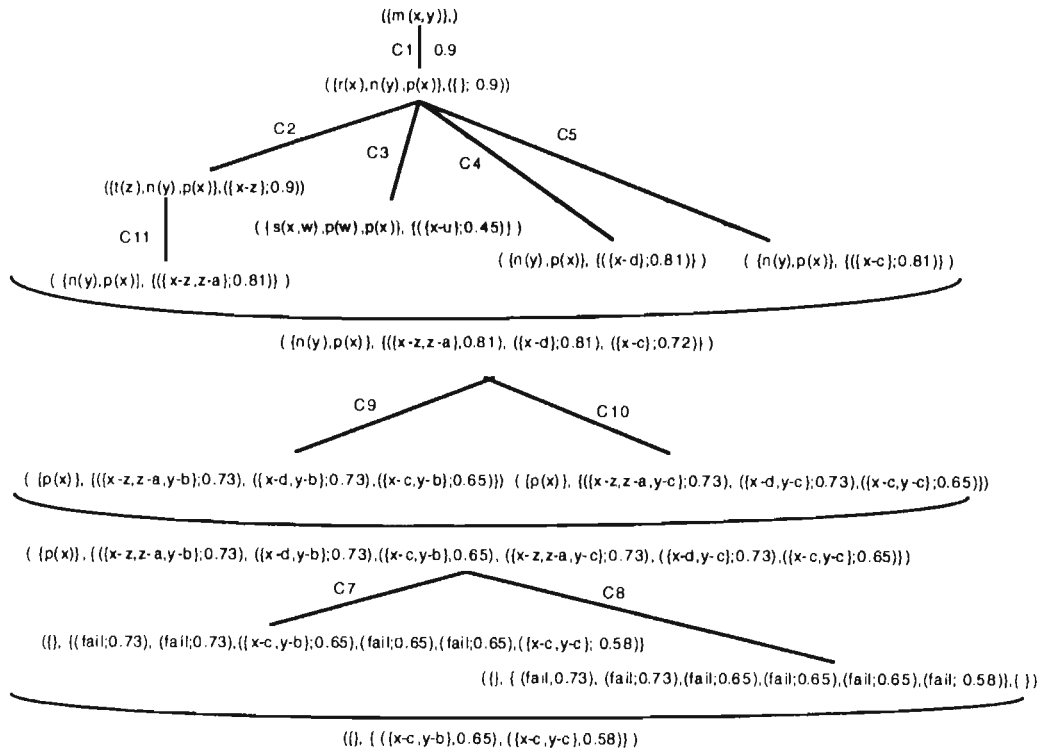
We can see that the previous interpreter has several drawbacks:

1. A high computational cost is paid handling inappropriately the substitutions as can be seen in the previous logic program. Thus, for each inference step: first, a unifier σ must be obtained, we call it inference unifier, between a head clause $\text{consequent}(C)$ and a literal q_1 in the current state; second, this unifier σ has to be applied to all the literals $\{q_2, \dots, q_n\}$ in the state at hand and in $\text{antecedents}(C)$; and third, a composition between the previous state unifier θ and the unifier inference σ must be performed in order to obtain the present state unifier. For instance, in the given problem, let us consider the state with unifier $\{z/x\}$ and literals $\{t(x), n(y), p(x)\}$ the unification between $t(x)$ and $t(a)$ leads to the unifier $\{x/a\}$. Then, this substitution is applied to the remaining literals of the goal which are instantiated in $\{n(y), p(a)\}$ (In this particular case $\text{antecedents}(C) = \{\}$). Finally, the composition between the actual state substitution $\{z/x\}$ and that just obtained $\{x/a\}$ is performed yielding to $\{z/a, x/a\}$. Moreover, it can be shown [5] that when first order terms (not only constants) are used the size of the terms of the substitutions could grow exponentially, i.e. the complexity of a unification step could increase exponentially with respect to the size of the involved literals in previous unifications up to the present one.
2. Each search space associated to the reduction of a given literal can be scanned many times. In the example, this is the case for literals p and n . It can be easily proved that the number of times that a search space might be examined can increase exponentially [5]. Thus backtracking based schemas [15, 13, 12] have expensive computational redundancies.

To eliminate these redundancies, we aim at designing a new interpreter based on a search strategy [5] completely different from the classical depth-first search and chronological backtracking, and with a different processing of unifiers. We will see that the proposed method avoids the redundancies of unfruitful backtrackings.

3.1 Efficient computing of unifiers

Concerning the computing of unifiers, we can note that the role of substitutions is to check that previous substitutions executed from the initial state until the present one represented by the state substitutions are compatible with the bindings of variables imposed by the present inference unification step. That is, previous inferences are compatible with the present one. Aiming at computing efficiently this compatibility we define a new operation *COMPAT* (equivalent

Figure 2: search space for $P \models (m(x, y); 0.5)$

to the rational infinite terms unification [2]). Hereforth, we will deal with no idempotent substitutions ρ and we will associate them with a system of bindings of variables $\{x-t : x/t \in \rho\}$. The *COMPAT* operation is applied to couples of systems of bindings of variables, one coming from the state substitution $S_\theta = \{x-t : x/t \in \theta\}$ and the other from the inference substitution $S_\sigma = \{y-t' : y/t' \in \sigma\}$ and it ends successfully iff the global system of bindings $S_\theta \cup S_\sigma$ can be unified. The proposed approach does not need to apply substitutions to literals.

Let us give a description of the operation *COMPAT*. In the literature unifiers are usually substitutions that verify the idempotent property, i.e. $\sigma \cdot \sigma = \sigma$. To compute this unifier any first order unification algorithm can be used. However, it has been shown [5] that it is not mandatory to employ unifiers with the idempotent property. Instead of handling idempotent substitutions, systems of bindings are used where a variable is bound at most to one term. The compatibility between the previous inference steps represented by the state bindings system S_θ and a present inference step represented by another inference bindings system S_σ is checked throughout the application of the *COMPAT* function over the two bindings systems. $COMPAT(S_\theta \cup S_\sigma)$ obtains another set of bindings $S'' = COMPAT(S_\theta \cup S_\sigma)$ such that $\{x/t : x-t \in S''\}$ is an (eventually not idempotent) mgu of the global set of bindings appearing in S_θ and S_σ . For example, if $S_\theta = \{x-y, y-a\}$ and

$S_\sigma = \{x-a\}$, $COMPAT(S_\theta \cup S_\sigma)$ is $\{x-y, y-a\}$. If $S' = \{x-b\}$ then $COMPAT(S_\theta \cup S_\sigma)$ fails to obtain an mgu. The operation *COMPAT* can be implemented by using the first phase of some first order unification algorithm [8, 4] which of course is computationally cheaper than obtaining an idempotent unifier with all the steps of the mentioned algorithms.

The mechanism of the interpreter concerning the processing of bindings is now the following:

A system of bindings variables S_θ is associated to each state of the interpreter. First, the inference system $S_\sigma = COMPAT(\{consequent(C)-q_1\})$ is calculated, such that $\sigma = \{x/t : x-t \in S_\sigma\}$ is an (eventually not idempotent) mgu of the pair of literals $\{consequent(C)-q_1\}$ involved in the inference step. Then, the bindings $S'' = COMPAT(S_\theta \cup S_\sigma)$ for each possible inference system S_σ such that $\theta = \{x/t : x-t \in S''\}$ is a (eventually not idempotent) mgu of the term pairs in $S_\theta \cup S_\sigma$ is obtained.

Now the actual approach of computing systems of bindings instead of idempotent unifiers is significantly faster than conventional computing of idempotent unifiers. It can be proved [5] that the operation of obtaining an idempotent unifier of a set of couples of terms is considerably harder than determining systems of bindings representing not idempotent unifiers.

Thus,

1. The conventional way of obtaining the idempotent mgu of $\{consequent(C)-q_1\}$ is computationally more expensive than

obtaining only the equivalent set of bindings of $\{antecedents(C) - q_1'\}$. Furthermore, literal q_1 is bigger in the conventional methods than the original literal of the goal q_1' because some previous mgu's (summarized in current state substitution θ) are successively applied to literal goals instantiating the variables of the original literal q_1' .

2. The composition of unifiers has a higher computational cost than the $COMPAT(S \cup S')$ operation because as mentioned previously the size of the unifiers can increase exponentially with respect to the size of the literals implicated in inference steps (when first order terms are used) whereas the computing of $COMPAT(S \cup S')$ only increases proportionally.
3. Note that no application of substitutions over literals is carried out in the proposed approach.

3.2 Search strategy

Concerning the search strategy, a new strategy is designed which is oriented towards the search for the total set of existing solutions [5], instead of being oriented towards the search for one solution as backtracking methods do. This strategy has the merit of avoiding most of the redundancies of previous schemas. Moreover, conversely to backtracking-based methods, the search space associated to the resolution of each literal in a clause is run over only once. Indeed, in the previous example and using chronological backtracking, the search space corresponding to the solution of literals n and p are obtained more than once (3 for n and 6 for p).

The idea behind the proposed search strategy is as follows. Being in a state e with literals $\{q_1, q_2, \dots, q_n\}$, all the possible paths beginning from this state and ending in states e' with literals $\{q_2, \dots, q_n\}$ are scanned. Then, the states of kind $(\{q_2, \dots, q_n\}, (S_i, \alpha_i))$ are obtained and grouped in a state $(\{q_2, \dots, q_n\}, \{(S_1, \alpha_1), \dots, (S_k, \alpha_k)\})$. Thus, a state in the present schema is formed by a set of literals and a set of couples formed by a binding system and a truth degree. The meaning associated to each state $(\{q_2, \dots, q_n\}, \{(S_1, \alpha_1), \dots, (S_k, \alpha_k)\})$ is that there exist k different paths from the initial state e with literal $\{q_2, \dots, q_n\}$ to a state e' with literals $\{q_2, \dots, q_n\}$ where (S_i, α_i) , $1 \leq i \leq k$, being respectively the systems and the truth value associated to the paths. The truth degree is calculated throughout the computing of the T-norm function over the program clauses used in the paths.

Thus, in a general case, a state $(\{q_1, q_2, \dots, q_n\}, \{(S_1, \alpha_1), \dots, (S_k, \alpha_k)\})$ contains k couples (systems, truth degrees). To verify which of these systems (or paths from the initial state) are compatible with a possible inference applicable to the present state, the $COMPAT$ function is applied for each head clause $consequent(C_j)$ with the same symbol predicate that q_1 . Thus, for the successful inference systems $S^j = COMPAT(\{consequent(C_j) - q_1\})$, the couple of bindings S^j and S_i are put together and the new set

of bindings $S_i^j = COMPAT(S_j \cup S_i)$ is obtained. Then the next states visited are $(antecedents(C_j) \cup \{q_2, \dots, q_n\}, \{(S_1^j, \alpha_1^j), \dots, (S_{n_j}^j, \alpha_{n_j}^j)\})$, $1 \leq j \leq m, n_j \leq k$. These states are visited and developed in a depth-first manner until reaching states $e_1, \dots, e_w$ with set of literals $\{q_2, \dots, q_n\}$. Then these sets are joined in a state whose literals are $\{q_2, \dots, q_n\}$ and whose set of couples (system, truth degree) is the union of the sets of couples in the states $e_1, e_2, \dots, e_w$. Notice that the paths between states with literals $\{q_1, q_2, \dots, q_n\}$ and $\{q_2, \dots, q_n\}$ are explored only once. Hence, a major difference in efficiency with respect to conventional methods is achieved since, as has been mentioned, the methods based on backtracking could execute this search space an exponential number of times.

The following algorithm is a more efficient first order interpreter with a more efficient computation of unifiers and exploration of the search space.

INTERPRETER(g, Set_I)

$Set_O \leftarrow \emptyset$

while $g \neq \emptyset$ do:

for each $(C, \alpha) \in P$ such that

$S_\sigma = COMPAT(consequent(C), head(g)) = S_\sigma \neq fail$ do:

$Set \leftarrow \{(COMPAT(S_\sigma \cup S_i), \alpha') :$

$(S_i, \alpha) \in Set_I, COMPAT(S_\sigma \cup S_i) \neq fail,$

$\alpha' = T(\alpha, val(C)) \geq \beta\}$

If $Set \neq \emptyset$ then do:

If $antecedents(C) = \{\}$ then $Set_O \leftarrow Set_O \cup Set$

Else $Set_O \leftarrow Set_O \cup INTERPRETER(antecedents(C), Set)$

$g \leftarrow tail(g)$

$Set_I \leftarrow Set_O$

return(Set_O)

Example

Let us consider the previous multiple-valued logic program problem. The interpreter begins with a state e_1 (see below). In this state, only an inference could be performed to reduce the literal $m(x, y)$. Applying this inference leads to state e_2 . Now, to resolve $p(x)$, 4 possible inferences with clauses C_2, C_3, C_4 and C_5 can be applied. Then the proposed schema reaches successively the states e_3 and e_4 applying clauses C_2 and C_{11} . Once a state with literals $\{n, p\}$ is obtained, this path is left and the search goes on by the next possible inference applying clause C_3 . But, as the truth value of this state e_5 is smaller than the imposed threshold 0.5 the search space following the reached state is skipped. Afterwards, clause C_4 is employed to reach a new state e_6 with literals $\{n, p\}$. After that, making inference with clause C_5 , another state e_7 with same literals $\{n, p\}$ is found. As no more inferences are applied to the state e_2 , the three states obtained with literals $\{n, p\}$ are merged in a state e_8 . The clauses C_9 and C_{10} to reduce the next literal n are considered. Throughout clause C_9 , state e_9 is reached after obtaining:

$COMPAT(\{n(y) - n(b)\}) = \{y - b\},$

$COMPAT(\{x - z, z - a\} \cup \{y - b\}) = \{x - z, z - a, y - b\}$

$COMPAT(\{x - d\} \cup \{y - b\}) = \{v - y, x - d, y - b\}$

$COMPAT(\{x - c\} \cup \{y - b\}) = \{x - c, y - b\}$

The search following this path is stopped and the other possible reduction literal with clause C_{10} is attempted yielding to state e_{10} . Then these two last states are merged in e_{11} . Finally, applying reducing literals with clauses C_7 and C_8 the final solution state e_{13} is determined. Notice that with this proposed schema there are neither fail states nor backtracking operations (notice that the search space corresponding to each literal is scanned at most once).

$$\begin{aligned}
 e_1 &= (\{m(x, y)\}, \{\{\}, 1\}) \\
 e_2 &= (\{r(x), n(y), p(x)\}, \{\{\{\}, 0.9\}\}) \\
 e_3 &= (\{t(z), n(y), p(x)\}, \{\{\{(x-z), 0.9\}\}\}) \\
 e_4 &= (\{t(z), n(y), p(x)\}, \{\{6\{x-z, z-a\}, 0.81\}\}) \\
 e_5 &= (\{s(u, w), p(w), n(y), p(x)\}, \{\{\{x-u\}, 0.45\}\}) \\
 e_6 &= (\{n(y), p(x)\}, \{\{\{x-d\}, 0.81\}\}) \\
 e_7 &= (\{n(y), p(x)\}, \{\{\{x-c\}, 0.72\}\}) \\
 e_8 &= (\{n(y), p(x)\}, \{\{\{x-z, z-a\}, 0.81\}, \{\{x-d\}, 0.81\}, \\
 &\quad \{\{x-c\}, 0.72\}\}) \\
 e_9 &= (\{p(x)\}, \{\{\{x-z, z-a, y-b\}, 0.72\}, \{\{x-d, y-b\}, \\
 &\quad 0.72\}, \{\{x-c, y-b\}, 0.64\}\}) \\
 e_{10} &= (\{p(x)\}, \{\{\{x-z, z-a, y-c\}, 0.72\}, \{\{x-d, y-c\}, \\
 &\quad 0.72\}, \{\{x-c, y-c\}, 0.64\}\}) \\
 e_{11} &= (\{p(x)\}, \{\{\{x-z, z-a, y-b\}, 0.72\}, \{\{x-d, y-b\}, \\
 &\quad 0.72\}, \{\{x-c, y-b\}, 0.64\}, \{\{x-z, z-a, \\
 &\quad y-c\}, 0.72\}, \{\{x-d, y-c\}, 0.72\}, \\
 &\quad \{\{x-c, y-c\}, 0.64\}\}) \\
 e_{12} &= (\{\}, \{\{fail, 0.72\}, \{fail, 0.72\}, \{\{x-c, y-b\}, 0.64\}, \\
 &\quad \{fail, 0.72\}, \{fail, 0.72\}, \{\{x-c, y-c\}, 0.64\}\}) \\
 e_{13} &= (\{\}, \{\{fail, 0.72\}, \{fail, 0.72\}, \{fail, 0.64\}, \{fail, 0.72\}, \\
 &\quad \{fail, 0.72\}, \{fail, 0.64\}, \{\}\}) \\
 e_{14} &= (\{\}, \{\{\{x-c, y-b\}, 0.64\}, \{\{x-c, y-c\}, \\
 &\quad 0.64\}\})
 \end{aligned}$$

4 Concluding remarks

One of the restrictions of classical logic programming languages is that each predicate can take only two values. In this paper, we have presented a way of overcoming this drawback by giving an interpreter for first-order multiple-valued logic programming.

The interpreter has been designed in such a way that the computation of unifiers and exploration of the search space are carried out efficiently. Another feature of this interpreter is that it can deal with a parametric family of multiple-valued logics since different interpretation functions for conjunction and implication can be used. In our proposal terms are only constants and variables and we consider as future work to extend the interpreter to handle functions too.

We also plan to adapt the negation as failure rule and the cut operator to the multiple-valued case.

Bibliography

- [1] BALDWIN, J. Evidential support logic programming. *Fuzzy Sets and Systems* 24 (1987), 1-26.
- [2] COLMERAUER, A. Equations and inequations on finite and infinite trees. In *Proceedings of the International Conference on Fifth Generation Computer Systems* (Tokio, 1984).
- [3] DUBOIS, D., LANG, J., AND PRADE, H. Poslog, an inference system based on possibilistic logic. *Proceedings North American Fuzzy Information Processing Society Congress* (1990), 177-180.
- [4] ESCALADA-IMAZ, G. *Optimisation d'Algorithmes d'Inference Monotone en Logique des Propositions et du Premier Ordre*. Université Paul Sabatier, Toulouse, 1989. (PhD Thesis).
- [5] ESCALADA-IMAZ, G. Un schema de contrôle en logique de premier ordre vers la recherche simultanée de toutes les solutions. *Revue d'Intelligence Artificielle* 4, 4 (1990), 49-107.
- [6] ESCALADA-IMAZ, G., AND MANYÀ, F. A linear interpreter for logic programming in multiple-valued propositional logic. In *Information Processing and Management of Uncertainty in Knowledge-Based Systems. IPMU'94* (1994), vol. 2, pp. 943-949.
- [7] LI, D., AND LIU, G. *A Fuzzy Prolog Database System*. Research Studies Press and John Wiley and Sons, 1990.
- [8] MARTELLI, A., AND MONTANARI, U. An efficient unification algorithm. *ACM Trans. Progr. Lang. Syst.* 4, 2 (1982), 258-282.
- [9] MARTIN, T., BALDWIN, J. F., AND PILSWORTH, B. W. The implementation of fpprolog: A fuzzy prolog interpreter. *Fuzzy Sets and Systems* 23 (1987), 119-129.
- [10] MUKAIDONO, M. Fundamentals of fuzzy prolog. *International Journal of Approximate Reasoning* 3 (1989), 179-193.
- [11] PATERSON, M. S., AND WEGMAN, M. N. Linear unification. *J. of Computer Science* 16, 2 (1978), 158-167.
- [12] PEREIRA, L. M., AND PORTO, A. Selective backtracking. In *Logic Programming*. Clark and Turnland, 1982.
- [13] PIETRZYKOWSKY, T., AND MATWIN, S. Exponential improvement of efficient backtracking. In *6th. International Conference on Automated Deduction* (1982), LNCS 138.
- [14] WEIGERT, T. J., TSAI, J., AND LIU, X. Fuzzy operator logic and fuzzy resolution. *Journal of Automated Reasoning* 10 (1993), 59-78.
- [15] WOLFRAN, D. A. Intractable unifiability problems and backtracking. In *3rd. International Conference on Logic Programming* (1986), LNCS 225.

Síntesis de funciones de pertenencia en la determinación de radios de influencia

Vicenç Torra
Dpt. d'enginyeria informàtica
Escola Tècnica Superior d'Enginyeria (ETSE)
Universitat Rovira i Virgili
Carretera de Salou, s/n
E-43006 Tarragona (Spain)
E-mail: vtorra@etse.urv.es

Resumen

Presentamos en este trabajo la síntesis de funciones de pertenencia cuando estas se han calculado a partir de un conjunto de observaciones de acuerdo con la definición de (Zhang, 1993). Se han aplicado los resultados en la determinación del radio de influencia de una trampa para estimar la densidad de una población de Podarcis lilfordi en la isla de Cabrera (Islas Baleares).

Palabras clave: Consenso, funciones de pertenencia, radios de influencia

1. Preliminares

En (Zhang, 1993) se presenta la definición formal de los conjuntos difusos mediante las llamadas particiones de franja. Se supone que para describir una categoría difusa $\hat{a}$ sobre un conjunto de referencia X existe una partición en tres conjuntos llamados X_0 , X_f y X_1 . El primer subconjunto está formado por todos aquellos elementos de X que no pertenecen a la categoría $\hat{a}$, X_1 por los que sin ninguna duda pertenecen a $\hat{a}$ y X_f consta de los elementos con pertenencia dudosa a la categoría $\hat{a}$. Llamamos a la partición de X formada por los tres conjuntos (X_0, X_f, X_1) partición en franja y a los subconjuntos de esta partición, respectivamente, de acuerdo con Zhang, subconjunto-0, franja y subconjunto-1.

Se supone además una ordenación parcial $\geq_{\hat{a}}$ en la franja donde $x \geq_{\hat{a}} y$ significa que x es como mínimo tan cualificado para ser miembro de $\hat{a}$ como y . Tenemos también una función μ

llamada medida de referencia que aplicada a un conjunto X corresponde al número de elementos que encontramos en el conjunto X .

Definición 1. La partición (X_0, X_f, X_1) referente a la categoría $\hat{a}$ junto con la función de ordenación parcial $\leq_{\hat{a}}$ y la medida de referencia μ permiten construir una función de pertenencia $\hat{A}_t$. Definimos, de acuerdo con (Zhang, 1993), dada una cuádrupla $(X, \mu, (X_0, X_f, X_1), \geq_{\hat{a}})$ la función de pertenencia $\hat{A}_t$ correspondiente a la categoría $\hat{a}$ como:

$$\hat{A}_t(x) := \begin{cases} 0 & \text{if } x \in X_0 \\ [\alpha(x), \alpha^+(x)] & \text{if } x \in X_f \\ 1 & \text{if } x \in X_1 \end{cases}$$

donde $\alpha(x)$ y $\alpha^+(x)$ se definen como:

$$\begin{aligned} \alpha(x) &:= \mu(\{y \in X_f | x >_{\hat{a}} y\}) / \mu(X_f) \\ \alpha^+(x) &:= \mu(\{y \in X_f | x \geq_{\hat{a}} y\}) / \mu(X_f) \end{aligned}$$

La relación $>_{\hat{a}}$ se construye a partir de $\geq_{\hat{a}}$ de la forma habitual.

En dominios en los que el conjunto de referencia tiene un orden total, por ejemplo la recta real, y en los que se conoce la monotonía de la función de pertenencia es posible construir a partir de un conjunto de datos experimentales un conjunto difuso asociado que aproximará la función de (Zhang, 1993). Definimos a continuación esta construcción.

Definición 2. Sea X un conjunto de referencia totalmente ordenado; sea $\geq$ la función de ordenación total. Sea Ξ un conjunto de observaciones de elementos de X y sea $\Xi=(\xi_0, \xi_f, \xi_1)$ una partición de Ξ en la que los elementos de cada ξ_i son elementos de la partición correspondiente X_i . Esto es, después de realizar la observación conocemos aquellos elementos observados que no pertenecen a la categoría $\hat{a}$, aquellos que pertenecen a ella con un cierto grado y aquellos que pertenecen completamente a ella. Decimos que la partición $\Xi=(\xi_0, \xi_f, \xi_1)$ es consistente con la función de ordenación si y solo si se satisfacen las condiciones siguientes:

$$\begin{aligned} \sup\{x: x \in \xi_0\} &\leq \inf\{x: x \in \xi_f\} \\ \sup\{x: x \in \xi_f\} &\leq \inf\{x: x \in \xi_1\} \end{aligned}$$

Definición 3. Sea X un conjunto de referencia totalmente ordenado; sea $\geq$ la función de ordenación total. Sea Ξ un conjunto de observaciones y sea $\Xi=(\xi_0, \xi_f, \xi_1)$ una partición consistente de X . Sea η una función de medida relativa (por ejemplo la función cardinalidad). En estas condiciones la función de pertenencia de la categoría $\hat{a}$ asociada a la cuádrupla $(X, \eta, (\xi_0, \xi_f, \xi_1), \geq)$ se define como:

$$\hat{A}(x) := \begin{cases} 0 & \text{if exists } a, b \in \xi_0 \text{ } a \leq x \leq b \\ 1 & \text{if exists } a, b \in \xi_1 \text{ } a \leq x \leq b \\ [\alpha(x), \alpha^+(x)] & \text{otherwise} \end{cases}$$

donde $\alpha(x)$ y $\alpha^+(x)$ se definen como:

$$\begin{aligned} \alpha(x) &:= \eta(\{y \in \xi_f | x > y\}) / \eta(\xi_f) \\ \alpha^+(x) &:= \eta(\{y \in \xi_f | x \geq y\}) / \eta(\xi_f) \end{aligned}$$

Esta función $\hat{A}$, que depende de la cuádrupla $(X, \eta, (\xi_0, \xi_f, \xi_1), \geq)$ aproxima a la función $\hat{A}_c$.

2. Resultados

Presentamos en esta sección diversos resultados sobre consenso de funciones de

pertenencia cuando estas se han construido mediante tuplas de la forma $(X, \eta, (\xi_0, \xi_f, \xi_1), \geq)$.

Definición 4. Dados dos conjuntos de observaciones y sus particiones $\Xi=(\xi_0, \xi_f, \xi_1)$ y $\Xi'=(\xi'_0, \xi'_f, \xi'_1)$ sobre un mismo dominio X , ambas consistentes respecto a una función de ordenación $\geq$, decimos que las particiones son compatibles cuando satisfacen:

$$\begin{aligned} \sup\{x: x \in \xi_0\} &< \inf\{x: x \in \xi'_f\} \\ \sup\{x: x \in \xi'_0\} &< \inf\{x: x \in \xi_f\} \\ \sup\{x: x \in \xi_f\} &< \inf\{x: x \in \xi'_1\} \\ \sup\{x: x \in \xi'_f\} &< \inf\{x: x \in \xi_1\} \end{aligned} \quad (r1)$$

Esto es, las particiones son compatibles cuando no existen solapamientos entre los subconjuntos de las dos particiones.

Definición 5. Sean $\Xi=(\xi_0, \xi_f, \xi_1)$ y $\Xi'=(\xi'_0, \xi'_f, \xi'_1)$ dos conjuntos de observaciones compatibles. El conjunto de observaciones unión $\Xi \cup$ se define como $\Xi \cup = (\xi_0 \cup \xi'_0, \xi_f \cup \xi'_f, \xi_1 \cup \xi'_1)$.

A continuación consideramos el consenso de funciones de pertenencia mediante medias aritméticas y comparamos el resultado con la función de pertenencia construida mediante el conjunto de observaciones unión. Presentamos primero un resultado que considera la función de pertenencia consenso cuando operamos con una media aritmética sin ponderación, y a continuación la media aritmética con ponderación.

Proposición 1. Sean $\hat{A}$ y $\hat{A}'$ dos funciones de pertenencia construidas mediante la definición 3 a partir de las cuádruplas $(X, \eta, \Xi=(\xi_0, \xi_f, \xi_1), \geq)$ y $(X, \eta, \Xi'=(\xi'_0, \xi'_f, \xi'_1), \geq)$. Sean los dos conjuntos de observaciones compatibles y con intersección de ξ_f y ξ'_f nula. En estas condiciones la función de pertenencia consenso definida como $A_c(x) = (\hat{A}(x) + \hat{A}'(x))/2$ es equivalente a la

función de pertenencia $\hat{A}_U$ construida a partir del conjunto de observaciones unión Ξ_U si la medida de ξ_f coincide con la medida de ξ'_f

Demostración. A partir de la igualdad

$$\hat{A}_c(x) = (\hat{A}_1(x) + \hat{A}_2(x)) / 2 \quad (1)$$

desarrollando $\hat{A}_i(x)$ para x en el conjunto $\xi_f \cup \xi'_f$, sustituyendo $\alpha_i(x)$ y $\alpha'_i(x)$ por sus definiciones

$$\alpha(x) = \eta(\{y \in \xi_f | x > y\}) / \eta(\xi_f)$$

$$\alpha'(x) = \eta(\{y \in \xi'_f | x > y\}) / \eta(\xi'_f)$$

$$\alpha_U(x) = \eta(\{y \in \xi_f \cup \xi'_f | x > y\}) / \eta(\xi_f \cup \xi'_f)$$

y definiendo para un cierto x de $\xi_f \cup \xi'_f$ las constantes α , β , a y b :

$$\alpha = \eta(\{y \in \xi_f | x > y\}) \quad a = \mu(\xi_f)$$

$$\beta = \eta(\{y \in \xi'_f | x > y\}) \quad b = \mu(\xi'_f)$$

escribimos la ecuación equivalente:

$$2(\alpha + \beta) / (a + b) = \alpha / a + \beta / b$$

A su vez la podemos reescribir como:

$$(a - b)(\alpha b - \beta a) = 0$$

que tiene por soluciones

$$a = b$$

$$\alpha / a = \beta / b$$

De estas condiciones solamente la primera, que equivale a

$$\eta(\xi_f) = \eta(\xi'_f) \quad (2)$$

tiene significado en el marco de las particiones en franja. La segunda equivale a:

$$\eta(\{y \in \xi_f | x > y\}) = \eta(\{y \in \xi'_f | x > y\}) \eta(\xi_f) / \eta(\xi'_f)$$

Haciendo en esta igualdad $x = x_0$ para un cierto $x_0 \in \xi_f$ obtenemos un cierto valor para $\eta(\{y \in \xi_f | x > y\})$; haciendo después $x = x_1$ con $x_1 > x_0$ de forma que $\eta(\{y \in \xi_f | x > y\})$ no varíe obtenemos para $\eta(\{y \in \xi'_f | x > y\})$ un valor diferente al anterior con lo que llegamos a una contradicción. Por tanto solamente (2) es una solución válida. $\square$

En las figuras 1 y 2 se presentan 2 ejemplos de funciones de pertenencia consenso. En estos ejemplos se utiliza como función η la función

cardinalidad. En el primero se presentan dos funciones que satisfacen $\eta(\xi_f) = \eta(\xi'_f)$ y su función consensuada. En la figura 2 se muestra el caso con $\eta(\xi_f) \neq \eta(\xi'_f)$ y se muestra tanto la función de pertenencia consenso mediante la media aritmética como la función de pertenencia asociada al conjunto unión de Ξ_U . En la figura 1 tenemos $\eta(\xi_1) = \eta(\xi_2) = 3$ y en la figura 2 $\eta(\xi_1) = 4$, $\eta(\xi_2) = 2$.

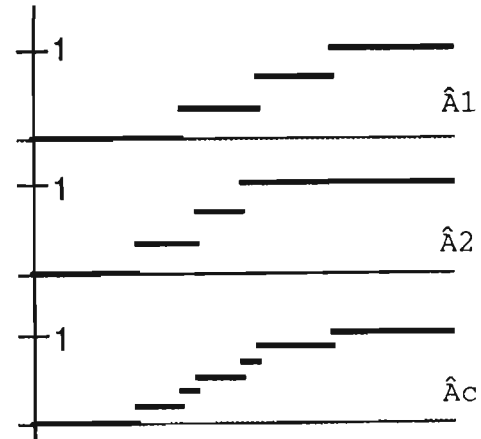


Figura 1

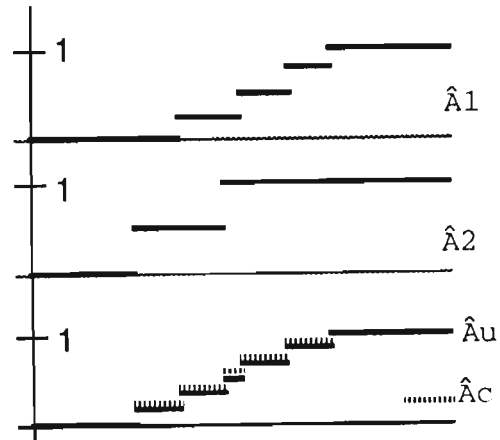


Figura 2

Proposición 2. Sean $\hat{A}$ y $\hat{A}'$ dos funciones de pertenencia construidas mediante la definición 3 a partir de las cuádruplas $(X, \eta, \Xi = (\xi_0, \xi_f, \xi'_f), \geq)$ y $(X, \eta, \Xi' = (\xi'_0, \xi'_f, \xi_f), \geq)$. Sean los dos conjuntos de observaciones compatibles y con la intersección de ξ_f y ξ'_f nula. En estas condiciones la función de pertenencia consenso definida como

$A_c(x) = (p\hat{A}(x) + q\hat{A}'(x))/2$ con $p+q=1$ es equivalente a la función de pertenencia $\hat{A}_\cup$ construida a partir del conjunto de observaciones unión $\Xi_\cup$ si y solo si

$$p = \eta(\xi_f) / (\eta(\xi_f) + \eta(\xi'_f)).$$

En las proposiciones precedentes hemos considerado el caso en que la intersección de ξ_f y ξ'_f es nula y por tanto de medida cero. Consideramos a continuación una intersección no vacía.

Proposición 3. Sean $\hat{A}$ y $\hat{A}'$ dos funciones de pertenencia construidas mediante la definición 3 a partir de las cuádruplas $(X, \eta, \Xi = (\xi_0, \xi_f, \xi_1), \geq)$ y $(X, \eta, \Xi' = (\xi'_0, \xi'_f, \xi'_1), \geq)$. Sean los dos conjuntos de observaciones compatibles. En estas condiciones la función de pertenencia consenso definida como $A_c(x) = (\hat{A}(x) + \hat{A}'(x))/2$ es equivalente a la función de pertenencia $\hat{A}_\cup$ construida a partir del conjunto de observaciones unión $\Xi_\cup$ si y solo si.

$$(a-b)(\alpha b - \beta a) + (a+c)(\alpha c - \gamma a) + (b+c)(\beta c - \gamma b) = 0$$

donde

$$\begin{aligned} c &= \eta(\xi_f \cap \xi'_f) & \gamma &= \eta(\{y \in \xi_f \cap \xi'_f \mid x > y\}) \\ a &= \eta(\xi_f) - c & \alpha &= \eta(\{y \in \xi_f \mid x > y\}) - \gamma \\ b &= \eta(\xi'_f) - c & \beta &= \eta(\{y \in \xi'_f \mid x > y\}) - \gamma \end{aligned}$$

Corolario 1. En el caso que $\eta(\xi_f) = \eta(\xi'_f)$ tenemos que las condiciones anteriores se cumplan si,

$$\eta(\xi_f) - \eta(\xi_f \cap \xi'_f) = \eta(\xi'_f) - \eta(\xi_f \cap \xi'_f) = 0$$

o bien,

$$\eta(\xi_f \cap \xi'_f) = 0$$

Esto es, o bien la medida de la intersección es cero y reencontramos los casos estudiados previamente o bien los conjuntos ξ_f y ξ'_f tienen la misma medida.

3. Aplicación

En esta sección aplicamos los resultados conseguidos en las secciones precedentes en un entorno con funciones de pertenencia monótonas definidas sobre el eje de los reales, por tanto con ordenación total. En esta aplicación, en lugar de tener un único conjunto de observaciones tenemos dos. El primero, llamado de observaciones positivas y representado Ξ^+ , consiste en aquellos x_i que apoyan a la construcción del conjunto difuso R , esto es que si $x \leq x_i$ entonces $R(x) \geq R(x_i)$; el segundo, de observaciones negativas (Ξ^-), formado por los que apoyan a la construcción del conjunto difuso negación de R , representado R_- , afirmando que si $x \geq x_i$ entonces $R_-(x) \geq R_-(x_i)$.

A partir de las particiones en franja de ambos conjuntos podemos construir las funciones de pertenencia R_{Ξ^+} y R_{Ξ^-} que aproximan R y R_- . Estos conjuntos difusos al ser complementarios deben satisfacer:

$$R_-(x) = N(R_+(x)) \quad (3)$$

igualdad que nos permite escribir R como la media ponderada de R y $N(R_-)$, esto es:

$$R = (p R + q N(R_-)) \quad \text{con } p+q=1 \quad (4)$$

En general, las aproximaciones R_{Ξ^+} y R_{Ξ^-} no satisfacen las igualdades (3) y (4). Sin embargo, puede emplearse (4) para conseguir una nueva aproximación de la función de pertenencia de R (R_s) más correcta a partir de los conjuntos de observaciones Ξ^+ y Ξ^- , de los resultados de síntesis y de el conjunto unión conseguido en la sección 3. Esto es:

$$R_s = (p R_{\Xi^+} + q N(R_{\Xi^-})) \quad \text{con } p+q=1 \quad (5)$$

Hemos aplicado los resultados descritos en las secciones precedentes en la determinación del radio de influencia de una trampa de tipo "Pitfall" (Telleria, 1986) para la captura de saurios como parte de un procedimiento (método de censo) para estimar la densidad de población de *Podarcis lilfordi* en la isla de Cabrera (Islas Baleares) (Sáez, 1993a, 1993b).

El método de censo utilizado fue el de Petersen (Krebs, 1989), consistente en una primera fase de captura y marcaje de individuos de la población y una segunda de captura (o

recaptura). El tamaño de la población se estima a partir de la fracción de individuos capturados que llevan marcas.

Para la captura se instalan un conjunto de trampas dispuestas en una malla que recubre toda la zona de muestreo (ver figura 1).

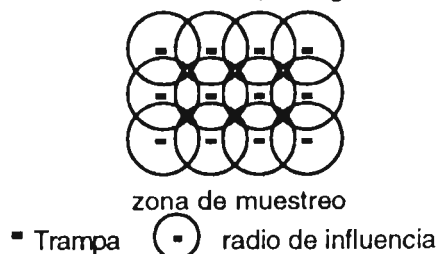


Figura 1. Zona de muestreo y radio de influencia

Este método supone que la población asentada en la parcela de muestreo es cerrada (esto es, no experimenta variaciones por natalidad, mortalidad o migración). Cada individuo ocupa un área particular, en la que el animal se mueve y realiza sus actividades rutinarias llamada dominio vital, que puede solaparse en cierto grado con otros dominios vitales de otros individuos.

Dos de las mayores dificultades que presenta esta técnica de censo son (Skalski, Robson, 1992):

1. Satisfacer el presupuesto de igual capturabilidad de los individuos muestreados. Esta condición implica una disposición adecuada de las trampas, de forma que todos los individuos que ocupan la superficie de la malla de trampas sean potencialmente capturables. Esto significa elegir una distancia entre trampas lo suficientemente pequeña como para que no se capturen tan sólo los individuos del entorno de las trampas dejando áreas sin muestrear en el interior de la malla. A su vez para conseguir mejores resultados con menos observaciones es adecuado maximizar la distancia entre trampas para cubrir el máximo de superficie muestreada.

2. Conocer la superficie real que ocupan los individuos muestreados para una estimación correcta de la densidad y por tanto de la población. La superficie real es

mayor que la ocupada por la malla de muestreo, puesto que existen individuos con dominios vitales que se solapan sólo parcialmente con la malla y que por tanto también son susceptibles de ser capturados.

Ninguno de los trabajos publicados sobre experiencias de marcaje-recaptura en saurios tiene en cuenta ambos problemas. En cambio, se han desarrollado varias formas de solventarlos en poblaciones de micromamíferos (Tellería, 1986; Skalski, Robson 1992), pero resultan inapropiadas para poblaciones de elevada densidad como *Podarcis lilfordi*, ya que requerirían la inversión de un enorme esfuerzo de muestreo.

Por ello se ideó una experiencia de marcaje-reobservación para conocer el radio de influencia de una trampa, es decir de qué distancia proceden los individuos capturados en ella. El valor resultante se utilizó como una aproximación a la distancia a la que se separaron las trampas y como la anchura de la banda que hay que añadir a la superficie de la malla para hallar el área realmente muestreada. Dicha experiencia consistió en el marcaje de los individuos capturados en una sola trampa a lo largo de un día (5 de agosto de 1993). Al día siguiente se recorrieron 8 transectos (itinerarios) radiales, aproximadamente equidistantes y convergentes en el lugar donde se situó la trampa. A lo largo de los recorridos se registraron las distancias a las que se observaron individuos, tanto marcados como no marcados.

A partir de las observaciones se calculó el radio de influencia aplicando los resultados de las secciones anteriores puesto que el radio de influencia es una cantidad difusa que puede representarse mediante una función monótona decreciente (suponiendo que la trampa se encuentra en el origen de coordenadas) ya que si a una cierta distancia de la trampa un animal se encuentra dentro del radio de influencia entonces cualquier animal a una distancia más pequeña se encuentra también.

Las observaciones se agruparon en dos conjuntos; el primero de observaciones positivas formado por los individuos marcados

(apoyan el concepto de radio de influencia), el segundo de observaciones negativas formado por los individuos no marcados (apoyan al concepto complementario "no influencia de la trampa").

Hemos definido a partir de estos conjuntos las particiones en franja del conjunto X de la forma siguiente (ver figura 2)

$$X_{R,1} = X_{\neg R,0} = [0, \inf_X \{X \in \Xi^-\})$$

$$X_{R,0} = X_{\neg R,1} = (\sup_X \{X \in \Xi^+\}, \infty)$$

$$X_{R,0} = X_{\neg R,1} =$$

$$= (\inf_X \{X \in \Xi^-\}, \sup_X \{X \in \Xi^+\})$$

y por consiguiente las particiones en franja de Ξ^+ y Ξ^- como:

$$\Xi^+ = (\Xi^+ \cap X_{R,0}, \Xi^+ \cap X_{R,f}, \Xi^+ \cap X_{R,1})$$

$$\Xi^- = (\Xi^- \cap X_{R,0}, \Xi^- \cap X_{R,f}, \Xi^- \cap X_{R,1})$$

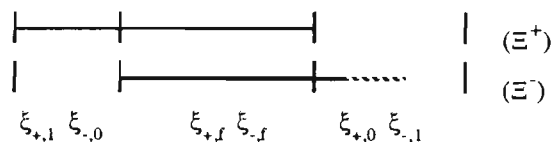


Figura 2. Particiones en franja. Individuos marcados (Ξ^+) y no marcados (Ξ^-).

Hemos construido a partir de estos conjuntos las funciones de pertenencia R_{Ξ^+} y R_{Ξ^-} y la función de pertenencia de síntesis R_s de acuerdo con (2) con p y q de forma que se satisficieran las condiciones de la proposición 2. A partir de la función de pertenencia R_s se escogió un radio de influencia con el que se realizaron 3 sesiones de captura durante 3 días consecutivos (10, 11 y 12 de agosto de 1993) mediante la instalación de 12 trampas. Los resultados del censo se encuentran en (Sáez, 1993b).

AGADECIMIENTOS

Agradezco a Encarna Sáez (Institut d'estudis avançats de les Illes Balears) la ayuda recibida y los datos suministrados, y también al Profesor Claudi Alsina (UPC, Barcelona).

Bibliografía

- [1] Dubois, D., Prade, H., "Fuzzy Sets and Systems: Theory and Applications", Academic press, New York, 1980.
- [2] Krebs, C.L., "Ecological methodology". Harper Collins, New York, 1989.
- [3] Sáez, E., "Population density of *Podarcis lilfordi* from the archipelago of Cabrera and some ecological factors affecting them", 7th ordinary general meeting Societas Europaeae Herpetologica, Barcelona, Setiembre, 1993.
- [4] Sáez, E., "Censo de las poblaciones insulares" en "Estudio de la fauna endémica y singular del archipelago de Cabrera", documento interno de ICONA, 1993.
- [5] Skalski, J.R., Robson, D.S., "Techniques for wildlife investigations. Design and analysis of capture data". Academic Press, San Diego, 1992.
- [6] Tellería, J.L., "Manual para el censo de vertebrados terrestres", Ed. Raices, Madrid, 1986.
- [7] Torra, V., "Contribució a l'estudi de funcions de síntesi per a la intel·ligència artificial", Tesis Doctoral, Universitat Politècnica de Catalunya, 1994.
- [8] Zhang, L., "Structural and functional quantization of vagueness", Fuzzy Sets and Systems, Vol. 55, (1993), pp. 51-60.

Hardware

Realizaciones Hardware de Sistemas Basados en Lógica Difusa

S. Sánchez Solano, C. J. Jiménez, A. Barriga, J. L. Huertas

Dpto. de Diseño de Circuitos Analógicos, Centro Nacional de Microelectrónica
Avda. Reina Mercedes s/n. Universidad de Sevilla. 41012 Sevilla
Tel.: 95 423 99 23. Fax: 95 423 18 32 - 462 45 06. E-mail: santiago@cnm.us.es

Resumen

Se propone una taxonomía que clasifica los sistemas de inferencia basados en lógica difusa de acuerdo con el carácter de la información que manejan en las diferentes etapas del proceso. Atendiendo a dicha taxonomía se revisan algunas de las implementaciones en hardware digital de sistemas difusos reportados en la literatura en la última década. Finalmente se proponen una serie de alternativas que permiten la realización eficiente de sistemas para aplicaciones de control en tiempo real.

1 Introducción

La lógica difusa aporta un marco conceptual y algorítmico para tratar aquellos problemas en los que la incompleta definición de las variables de estado y/o la vaguedad de la estrategia de resolución a aplicar requieren la utilización de mecanismos de inferencia aproximados. A diferencia del esquema de razonamiento humano, que nos permite manipular información imprecisa con relativa facilidad, el carácter marcadamente numérico de los sistemas de cálculo convencional los hace poco eficientes al tratar con este tipo de problemas. Cuando las técnicas de inferencia fuzzy se aplican a situaciones que requieren una elevada velocidad de respuesta (toma de decisiones y control en tiempo real) surge la necesidad de implementar en hardware nuevas estructuras de cálculo que respondan al paradigma fuzzy.

En un sistema de inferencia basado en lógica difusa la base de conocimientos está estructurada en un conjunto de reglas de descripción lingüística del tipo IF-THEN:

$$\text{si } X_1 \text{ es } A_1^f \text{ y } X_2 \text{ es } A_2^f \text{ y } \dots X_I \text{ es } A_I^f \Rightarrow Y \text{ es } B^f \quad (1)$$

donde X_i e Y corresponden a las entradas y la salida del sistema y A_i^f y B^f son etiquetas lingüísticas que representan funciones de pertenencia descritas mediante conjuntos difusos.

Cada una de las reglas del sistema induce una relación difusa (R^f) definida por:

$$\mu_{R^f}(u_1, u_2, \dots, u_I, v) = \Phi \{ T_c [\mu_{A_1^f}(u_1), \mu_{A_2^f}(u_2), \dots, \mu_{A_I^f}(u_I)], \mu_{B^f}(v) \} \quad (2)$$

donde u_i y v son elementos de los universos de discurso de las entradas y la salida del sistema, T_c es el conectivo de conjunción de los antecedentes (con propiedades de una T-norma) y Φ es una de las funciones de implicación que aparecen reportadas en la literatura [1] [2].

La regla composicional de inferencia de Zadeh establece que, dadas una regla y una observación (A'), la conclusión (B') puede obtenerse como la composición Sup-* ($\circ$) entre la observación y la relación difusa inducida por la regla. El operador $*$ es también una T-norma.

$$B' = A' \circ R^f \quad (3)$$

La conclusión global inferida al aplicar las diferentes reglas que constituyen la base de conocimientos se obtendrá por agregación de las soluciones parciales aportadas por cada regla:

$$B' = \bigcup_r B'^r \quad (4)$$

Si definimos la matriz de implicación global como la combinación de las matrices de implicación que representan las distintas relaciones difusas:

$$R = \bigcup_r (R^1, R^2, \dots, R^R) \quad (5)$$

la ecuación (4) puede expresarse como:

$$B' = A' \circ R \quad (6)$$

2 Sistemas de Inferencia Fuzzy

Según se desprende de la formulación anterior, un sistema de inferencia basado en lógica difusa viene determinado por las opciones elegidas para implementar los conectivos de conjunción de antecedentes (T_c) y de agregación de reglas ($\cup$), la función de implicación (Φ) y la T-norma de la regla composicional de inferencia ($*$). Sin embargo, a la hora de describir las distintas realizaciones en hardware de los sistemas fuzzy es preferible introducir una clasificación de los mismos atendiendo al tipo de información que manipulan en las entradas al sistema y en la descripción de los consecuentes de las reglas.

2.1 FIFC (Entrada Fuzzy - Consecuente Fuzzy)

El primer tipo de sistema que consideraremos es aquel en el cual las entradas vienen descritas mediante conjuntos difusos. Esto incluye, tanto a sistemas de toma de decisiones donde las entradas son etiquetas lingüísticas asociadas a distribuciones de posibilidades, como a sistemas complejos de control en los que las entradas se generan a partir de una distribución estadística. En este caso, la distribución de posibilidades obtenida como conclusión al aplicar una observación a una regla viene dada por:

$$\mu_{B^r}(v) = \bigvee_u \{ T_c [\mu_{A^r_1}(u_1) \dots \mu_{A^r_I}(u_I)] * R^r \} \quad (7)$$

Si sustituimos la ecuación (2) en (7):

$$\mu_{B^r}(v) = \bigvee_u \{ [T_c (\mu_{A^r_1}(u_1) \dots \mu_{A^r_I}(u_I))] * [\Phi (T_c (\mu_{A^r_1}(u_1) \dots \mu_{A^r_I}(u_I)) , \mu_{B^r}(v))] \} \quad (8)$$

La realización en hardware de la expresión anterior puede facilitarse considerablemente si la función de implicación es una T-norma (T_Φ), ya que en este caso, aplicando las propiedades asociativa y distributiva de las T-normas respecto al operador $\bigvee$, (8) puede reescribirse como:

$$\mu_{B^r}(v) = T_\Phi \{ T_c [\bigvee_{u_1} (\mu_{A^r_1}(u_1) * \mu_{A^r_I}(u_I) \dots \bigvee_{u_I} (\mu_{A^r_1}(u_1) * \mu_{A^r_I}(u_I))] , \mu_{B^r}(v) \} \quad (9)$$

Cada uno de los términos que aparecen entre corchetes en la expresión (9) definen el grado de similitud entre una entrada difusa y su correspondiente antecedente en la regla. La composición de dichos grados de similitud mediante el operador T_c determina el grado de activación (GDA) de la regla r -ésima. De esta forma, bajo la suposición de que la función de implicación sea una T-norma, la ecuación (4) puede reescribirse como:

$$B^r = \bigcup_r [T_\Phi (GDA^r , B^r)] \quad (10)$$

2.2 SIFC (Entrada Singleton - Consecuente Fuzzy)

En las aplicaciones de control las entradas provienen normalmente de sensores que proporcionan valores concretos. En lugar de "fuzzificar" las entradas para implementar un sistema FIFC, dichos valores pueden interpretarse como singularidades difusas. Un fuzzy-singleton es un conjunto difuso que incluye un único elemento con grado de pertenencia 1:

$$\mu_{A^r}(u_i) = 0 \text{ si } u_i \neq \bar{u}_i, \mu_{A^r}(\bar{u}_i) = 1 \quad (11)$$

Mientras que esta circunstancia no modifica en absoluto el cálculo de la matriz de implicación global de la ecuación (6), sí simplifica considerablemente el cálculo del grado de activación en la expresión (10), ya que dicho cálculo se reduce a la obtención de los grados de pertenencia de cada antecedente para el valor de su entrada y su composición mediante el operador de conjunción. El resultado de la inferencia puede expresarse como:

$$\mu_{B^r}(v) = \bigcup_r \{ T_\Phi [\alpha^r , \mu_{B^r}(v)] \} \quad (12)$$

donde

$$\alpha^r = GDA^r = T_c [\mu_{A^r_1}(\bar{u}_1), \dots, \mu_{A^r_I}(\bar{u}_I)] \quad (13)$$

Desde el punto de vista del coste computacional del algoritmo de inferencia, mientras que la obtención de la conclusión en los sistemas FIFC implica dos operaciones vectoriales (para recorrer el universo de discurso de los antecedentes y del consecuente), en los SIFC el cálculo se reduce a una operación escalar y otra vectorial.

2.3 SISC (Entrada Singleton - Consecuente Singleton)

Considerando de nuevo las aplicaciones de la lógica difusa al control, la conclusión del sistema de inferencia debe ser sometida a un proceso de concretización antes de ser aplicada a los actuadores del sistema controlado. Por este motivo los sistemas FIFC y SIFC incluyen, habitualmente, una etapa de procesamiento adicional que calcula el valor más significativo de la conclusión difusa. La complejidad de esta etapa va a depender del método de defuzzificación utilizado y de las características de las funciones de pertenencia de los consecuentes de las reglas (simetría, solapamiento, etc), pero en general representa la realización de una nueva operación vectorial. Por ejemplo, la técnica normalmente utilizada del centro de gravedad viene expresada por la ecuación:

$$\bar{y} = \sum_v \mu_B(v) \cdot v / \sum_v \mu_B(v) \quad (14)$$

Para acortar el tiempo de procesamiento de esta etapa del controlador suele recurrirse a la utilización de métodos simplificados que permitan manipular información "pre-defuzzificada". Desde el punto de vista de implementación, el más eficiente de estos métodos es el que sustituye la información difusa de los consecuentes de las reglas por un fuzzy-singleton localizado en el lugar de máximo grado de pertenencia del conjunto difuso original.

Un sistema de inferencia SISC puede interpretarse como aquel en que cada regla propone una conclusión concreta, con una determinada "fuerza" definida por su correspondiente grado de activación. El cálculo de la acción combinada de varias reglas requiere un mecanismo para evaluar la media de las diferentes conclusiones, ponderada por el grado de activación:

$$\bar{y} = \sum_r \alpha^r \cdot B^r / \sum_r \alpha^r \quad (15)$$

Los sumatorios del cálculo del centroide en la ecuación (14) se extienden a todo el universo de discurso del consecuente. Por el contrario en la ecuación (15) el índice de suma es el número de reglas. Si consideramos además que el número de reglas activas es normalmente muy inferior al número total de reglas del sistema, puede deducirse que la velocidad de inferencia de los sistemas SISC será superior a la de los otros tipos de sistemas fuzzy descritos previamente.

3 Realizaciones de Sistemas de Inferencia Fuzzy

Al considerar la implementación electrónica de un sistema de inferencia fuzzy pueden seguirse tres estrategias diferentes. La primera de ellas consiste en resolver off-line la ecuación (6) y, posteriormente, implementarla mediante memoria o lógica programable [3] [4]. Si bien esta aproximación es la que permite un mayor grado de libertad a la hora de elegir las diferentes opciones del sistema fuzzy, su realizabilidad práctica se limita a casos donde el número de entradas y el grado de discretización es reducido. Una segunda aproximación consiste en calcular off-line la matriz de implicación global (5) y realizar la inferencia on-line de acuerdo con (6) [5]. Este procedimiento permite una total libertad a la hora de elegir la función de implicación, pero presenta problemas de implementación similares a los de la alternativa anterior. Finalmente, algunos mecanismos de inferencia como los de Mamdani y Larsen (que utilizan, respectivamente, el mínimo y el producto como funciones de implicación) admiten una realización completamente on-line donde, siguiendo la ecuación (10), se evalúa regla a regla la conclusión del sistema para el valor actual de las entradas. Esta estrategia permite minimizar la cantidad de memoria necesaria en las aproximaciones anteriores.

Una de las principales características de un sistema fuzzy es su elevado nivel de paralelismo. Dicho paralelismo surge como consecuencia de dos factores. Por un lado, la presencia de operaciones vectoriales que afectan a todos los elementos de los conjuntos difusos. Por otra parte, en un sistema basado en lógica difusa, a diferencia de un sistema experto convencional, varias reglas pueden estar activas simultáneamente cooperando con diferentes niveles de activación a la solución final. Como discutiremos en los siguientes apartados, la forma de tratar ambos niveles de paralelismo va a condicionar la arquitectura de los diferentes tipos de sistemas fuzzy.

3.1 Procesado de Reglas en Paralelo

La inmensa mayoría de las realizaciones analógicas y las primeras realizaciones digitales de sistemas FIFC [6] y SIFC [7] recurren a una arquitectura en la que se proporciona un data-path para cada regla (Fig. 1a). En este caso las funciones de pertenencia de antecedentes y consecuentes están asociadas a las reglas, de forma que la operación de cálculo de grados de activación de las diferentes reglas se lleva a cabo en paralelo. Sin embargo, mientras que en algunas de las implementaciones analógicas se realiza en paralelo también el procesamiento de los consecuentes, el elevado coste en hardware que ello supone hace que, en las realizaciones digitales, se recurra a la serialización de la operación del sistema.

Incluso en este último supuesto es preciso proveer al sistema de operadores de entradas múltiples que realicen la agregación de las diferentes reglas. Dichos operadores se implementan mediante conexión en cascada de operadores de dos entradas, lo que ralentiza considerablemente la operación del sistema. Además, el espacio de memoria necesario para almacenar los conjuntos difusos que representan a los antecedentes y consecuentes es proporcional al número de reglas, lo que puede imponer serias restricciones a la realización de sistemas con un elevado número de reglas.

Las distintas operaciones de procesado en un sistema FIFC corresponden a: 1) cálculo de los grados de activación de las diferentes reglas; 2) cómputo de las conclusiones parciales; 3) agregación de la conclusión; y 4) defuzzificación. La Fig. 1b muestra el número de ciclos de reloj necesarios para realizar cada una de estas etapas, donde E_i y E_o son los grados de discretización de los universos de discurso de las entradas y la salida y $N_o = \log_2 E_o$ coincide con el número de bits de precisión que se requieran en la obtención del resultado final de la inferencia. Una vez realizada la operación 1, las operaciones 2 y 3 y el cálculo del numerador y denominador del centroide pueden realizarse simul-

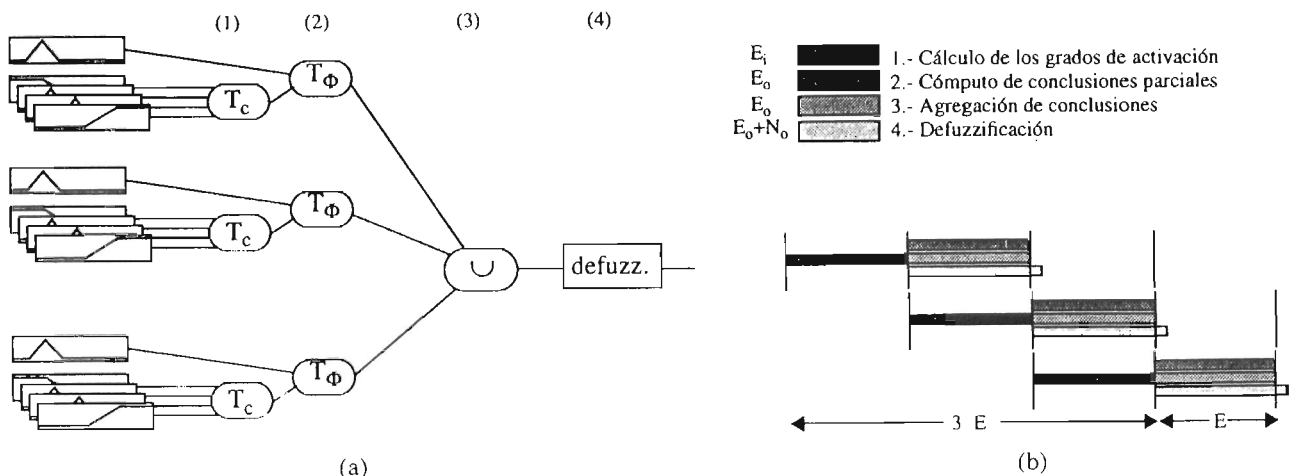


Fig. 1.- Procesado de reglas en paralelo: a) arquitectura; b) temporización.

táneamente. De esta forma, aplicando técnicas de segmentación con tres etapas de pipeline, y suponiendo que el número de elementos de los universos de discurso de entradas y salida es E , el número de ciclos de reloj necesarios para realizar una inferencia es $3E$, pero el sistema es capaz de proporcionar una nueva salida cada E ciclos.

En un sistema SIFC, por el contrario, el esquema de segmentación queda descompensado porque el cálculo del grado de activación de cada regla se reduce a una operación escalar. Este hecho, junto con los inconvenientes de área que comentamos antes, llevó a plantear nuevas arquitecturas que describiremos a continuación.

3.2 Procesado de Reglas en Serie

Una segunda estrategia de realización, ampliamente utilizada en implementaciones digitales, es aquella que estructura la base de conocimiento en dos zonas claramente diferenciadas. Una base de reglas almacenada en una memoria, donde las reglas vienen expresadas en forma simbólica mediante las etiquetas lingüísticas de los antecedentes y consecuentes. Y unos bloques de generación de las funciones de pertenencia de antecedentes y consecuentes, que son en este caso comunes a toda la base de reglas y obtenidas como una partición difusa del correspondiente universo de discurso (Fig. 2a).

Puesto que la utilización de una memoria sólo permite acceder a una regla en cada ciclo de reloj, el procesamiento de la base de reglas debe realizarse secuencialmente. Los grados de activación de las reglas que comparten un mismo consecuente son adecuadamente combinados y almacenados en elementos de memoria [8] [9].

Con esta aproximación el ahorro de memoria es considerable. Asimismo, el coste de implementación de la operación 3 depende del número de funciones de pertenencia del consecuente en lugar de depender del número de reglas. El tiempo de respuesta puede mante-

nerse igual que en el caso anterior (una inferencia cada E ciclos) siempre que el número de reglas sea menor que el número de elementos del universo de discurso del consecuente (Fig. 2b).

Con objeto de poder aumentar el número de reglas sin degradar la respuesta temporal del sistema algunos autores proponen otras alternativas. Sasaki et al. [10] [11] utilizan un esquema de reagrupamiento de reglas donde la base de reglas original es particionada en varias bases de reglas asociadas a cada uno de los posibles consecuentes. Esta aproximación da lugar a una estructura SIMD donde cada elemento de procesamiento realiza la inferencia sobre una de las sub-bases de reglas. Por otro lado, Ikeda introduce el concepto de reglas activas, como aquellas en que el grado de activación es distinto de cero. Puesto que la aportación a la salida de las reglas no activas es nula, no tienen por qué contemplarse en el proceso de inferencia. No obstante, las implementaciones propuestas en [12] y [13] plantean el inconveniente de necesitar una estructura de cálculo de regla activa que ocupa un área excesiva.

En las arquitecturas que hemos revisado hasta ahora la velocidad máxima de inferencia viene limitada por el hecho de tener que procesar secuencialmente la información difusa del consecuente, mientras que uno de los factores que en mayor medida contribuyen al área es el espacio necesario para almacenar o generar las funciones de pertenencia.

Diferentes autores han propuesto mecanismos para mejorar la velocidad de respuesta de los sistemas fuzzy. [14] y [15] proponen esquemas de pre-defuzzificación que aceleran el cálculo. [9] y [11], por su parte, utilizan consecuentes tipo singleton aunque procesan las reglas en serie.

En relación con el área, si se limita a priori el grado de solapamiento de las funciones de pertenencia, es posible reducir mucho el tamaño del sistema [14]. Además, en este caso el número de reglas potencialmente activas queda también limitado como el grado de solapamiento elevado al número de entradas [16]. Combi-

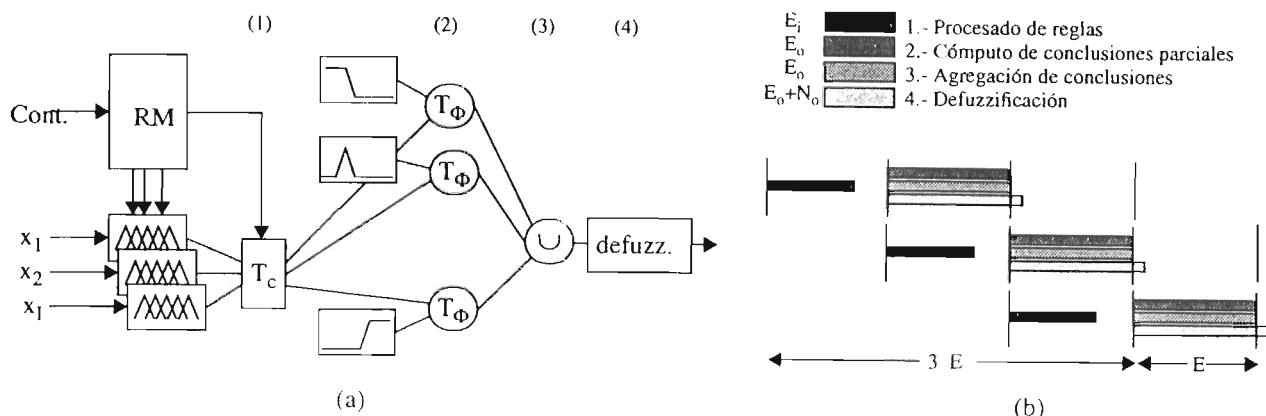


Fig. 2.- Procesado de reglas en serie. a) arquitectura, b) temporización.

nando adecuadamente ambos conceptos es posible llevar al silicio de forma muy eficiente el tipo particular de sistema difuso que describiremos en el siguiente apartado.

4 Sistema SISC con Procesado de Reglas Activas

En los sistemas SISC el número de ciclos de reloj necesarios para evaluar la conclusión es igual a la precisión. Por este motivo no es eficiente analizar secuencialmente la base de reglas completa. La velocidad del sistema puede incrementarse considerablemente si recurrimos a la utilización de un mecanismo de inferencia conducido por las reglas activas. El esquema de base de reglas simbólicas almacenadas en una memoria convencional debe, en este caso, ser sustituido por un esquema de memoria asociativa donde las reglas activas busquen su conclusión.

La Fig. 3a ilustra simbólicamente el procedimiento que proponemos. Fijado un conjunto de valores de entrada, los distintos bloques de generación de funciones de pertenencia proporcionan tantos pares (etiqueta - grado de pertenencia) como grado de solapamiento se haya previsto en el sistema. Un conjunto de multiplexores controlados por un contador permite recorrer secuencialmente las distintas combinaciones de reglas potencialmente activas. En cada ciclo del contador los grados de pertenencia son procesados a través del operador de conjunción para calcular el grado de activación de la regla, mientras que las etiquetas de los antecedentes direccionan la posición de memoria que contiene su correspondiente consecuente. Por último, la suma de los sucesivos productos de grado de activación por consecuente corresponde al numerador de la ecuación (15), cuyo denominador es la suma de los grados de activación de las diferentes reglas.

En este tipo de sistema las distintas operaciones de procesamiento son: 1) cálculo de los grados de activación de las reglas activas, 2) cómputo de las sumas parciales del numerador y denominador, y 3) división. El número de ciclos que se requieren para realizar la inferencia es

igual al número de reglas potencialmente activas más el empleado en la división (típicamente uno o dos ciclos más que la precisión que deseamos obtener). Como se ilustra en la Fig. 3b, la elección de un esquema de segmentación adecuado permite reducir el número de ciclos entre dos inferencias consecutivas al mayor de los dos valores anteriores.

Aunque las prestaciones de los sistemas SISC con procesamiento de reglas activas son mucho mejores que las de otros tipos de realizaciones, su velocidad de operación y el consumo de área pueden incluso mejorarse si se utilizan determinadas técnicas de implementación. Así, para reducir el tamaño de la memoria de consecuentes cuando hay reglas inespecificadas o cuando el número de etiquetas no es potencia de dos, puede recurrirse a la utilización de un circuito combinacional (PLA o estructuras lógicas programables) que establezca un mapeado adecuado entre las entradas y la memoria. Por otra parte, cuando el número de reglas potencialmente activas es superior a la precisión (número de entradas o grado de solapamiento elevado), es posible utilizar un contador con entradas de inhibición para controlar los multiplexores, de forma que sólo se procesen las reglas activas.

5 Conclusiones

El coste de implementación y la velocidad de operación de un sistema de inferencia basado en lógica difusa dependen del tipo de sistema considerado y de la estrategia seguida en la evaluación de las reglas. En esta comunicación se han clasificado los sistemas fuzzy de acuerdo con la naturaleza difusa o concreta de la información que manejan en las diferentes etapas del proceso y se han descrito las tres estrategias de evaluación de reglas que mejor se adaptan a cada tipo de sistema.

Los principales resultados obtenidos se resumen en la tabla 1, donde E representa el número de elementos del universo de entrada, R el número de reglas, R_a el número de reglas activas, I el número de entradas, S el grado de solapamiento, F_0 el número de etiquetas lin-

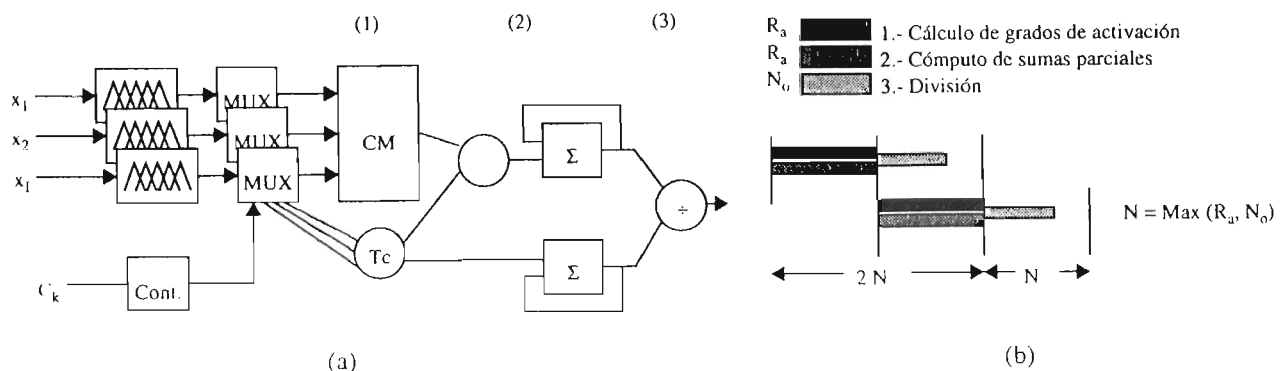


Fig. 3.- Procesado de reglas activas: a) arquitectura; b) temporización.

Tipo	Procesado de Reglas	Nº Ciclos / Inferencia	Unidades de Memoria
FIFC	En Paralelo	E	$R * (I + 1)$
SIFC	En Serie	Max (E, R)	$I * S + F_0$
SISC	Activas	Max (R_a , N_0)	$I * S$

Tabla 1: Comparación de los distintos tipos de sistemas fuzzy.

güísticas de la salida y N_0 el número de bits de la salida. En todos los casos se está suponiendo la existencia de una única salida y no está siendo tenida en cuenta la memoria de almacenamiento de reglas. Cada unidad de memoria referida en la columna "unidades de memoria" es la memoria necesaria para almacenar una función de pertenencia (elementos del universo de discurso x nº de bits de los grados de pertenencia).

Bibliografía

- [1] M. Mizumoto and H. Zimmermann, "Comparison of fuzzy reasoning methods", *Fuzzy Sets and Systems*, vol. 8 pp. 253-283, 1982.
- [2] Chuen Chien Lee, "Fuzzy Logic in Control Systems: Fuzzy Logic Controller", Partes I y II, *IEEE Transactions on Systems, Man and Cybernetics*, Vol. 20 N. 2 (1990) pp. 404-432.
- [3] M. A. Manzoul and D. Jayabharathi, "Fuzzy Controller on FPGA Chip", *Proc. of 1st IEEE ICFS'92*, San Diego, pp. 1389-1316.
- [4] J.Y. Leong, M.H. Lim and K.T. Lau, "A General Approach to Encoding Heuristics on Programmable Logic Devices", *Proc. fifth IFSA World Congress 1993*, Seoul, pp. 917-920.
- [5] A. Bugarin, S. Barro y R. Ruiz, "Soluciones Hardware en Control Borroso", editores "A. Sobrino y S. Barro", "Estudios de Lógica borrosa y sus aplicaciones", Universidad de Santiago de Compostela, 1993.
- [6] H. Watanabe, W. Dettloff and K. Young, "A VLSI Fuzzy Logic Controller with Reconfigurable, Cascadable Architecture", *IEEE J. Solid State Circuits*, Vol. 25 No 2 Apr 1990 pp 376-382.
- [7] M. Togai and S. Chiu, "A Fuzzy Logic Chip and Fuzzy Inference Accelerator for Real-Time Approximate Reasoning", *Proc. of 17th. International Symposium of Multiple Valued Logic*, Mayo 1987, pp. 25-29.
- [8] A.P. Ungering, K. Thuener and K. Goser, "Architecture of a PDM VLSI Fuzzy Logic Controller with Pipelining and Optimized Chip area", *Proc. IEEE ICFS'93*, San Francisco, pp. 447-452.
- [9] T. Katashiro, "A Fuzzy Microprocessor for Real-time Control Applications", *Proc. fifth IFSA World Congress 1993*, Seoul, pp. 1394-1397.
- [10] M. Sasaki, F. Ueno and T. Inoue, "An 8-bit Resolution 140 kFLIPS Fuzzy Microprocessor", *Proc. fifth IFSA World Congress 1993*, Seoul, pp. 921-924.
- [11] M. Sasaki, F. Ueno and T. Inoue, "7.5 MFLIPS Fuzzy Microprocessor Using SIMD and Logic-in-Memory Structure", *Proc. IEEE ICFS'93*, San Francisco, pp. 527-534.
- [12] H. Ikeda, Y. Hiramoto and N. Kisu, "A Fuzzy Inference Processor with an "Active-Rule-Driven" Architecture", *Proc. of Symposium on VLSI circuits*, 1991, pp. 25-26.
- [13] H. Ikeda, N. Kisu, Y. Hiramoto and S. Nakamura, "A Fuzzy Inference Coprocessor Using a Flexible Active-Rule-Driven Architecture", *Proc. IEEE ICFS'92*, San Diego, pp. 537-544.
- [14] Tzi-cker Chiueh, "Optimization of Fuzzy Logic Inference Architecture", *IEEE Computer*, Vol. 25 N. 5 (1992), pp. 67-71.
- [15] D. J. Ostrowski, P.Y.K. Cheung and K. Roubaud, "An Outline of the Intuitive Design of Fuzzy Logic and its Efficient Implementation", *Proc. IEEE ICFS'93*, San Francisco, pp 184-189.
- [16] Herbert Eichfeld, Michael Löhner and Michael Müller, "Architecture of a CMOS Fuzzy Logic Controller with optimized memory organisation and operator design" *Proc. IEEE ICFS'92*, San Diego, pp 1317-1323

ARQUITECTURAS SISTÓLICAS DE PROCESAMIENTO DIFUSO Y SU SIMULACIÓN

Luis de Salvador⁽¹⁾, Marcos García⁽¹⁾, Julio Gutiérrez^(1,2)

⁽¹⁾ Laboratorio de Hardware y Control Avanzado
Instituto Nacional de Técnica Aeroespacial
Ctra. Ajalvir km. 4 - 28850

Torrejón de Ardoz - Madrid - España
e-mail: salvadorla@inta.es

⁽²⁾ Dept. Tecnología Fotónica - Univ. Politécnica de Madrid
Campus de Montegancedo - 28660- Madrid

Resumen

El presente trabajo describe el estudio de arquitecturas sistólicas de procesamiento de reglas difusas realizado en el Laboratorio de Hardware y Control Avanzado-INTA. Las bases teóricas de dichas arquitecturas son desarrolladas y analizadas. Asimismo, los esquemáticos resultantes son simulados utilizando un lenguaje de descripción hardware (VHDL) sobre librerías "standard cells" de ES2. Esto permite una evaluación muy aproximada de sus rendimientos reales. De esta forma se detectan las deficiencias inherentes a este tipo de sistemas y se esbozan distintos caminos para superarlas.

Palabras clave: Fuzzy, Control Fuzzy, Sistólicos, Simulación, VHDL.

1 Introducción

Uno de los campos de investigación en la lógica difusa es el desarrollo de "hardware" específico que permita obtener todas las ventajas de trabajar con este tipo de modelos. Existen, fundamentalmente, dos ramas de investigación en controladores borrosos: los diseños analógicos y los diseños digitales. Dentro de estos últimos destacan las soluciones basadas en arquitecturas sistólicas.

Sistolizar un problema consiste en dividir este en operaciones elementales e idénticas (en su mayor parte) que pueden ejecutarse sincrónicamente en paralelo [1],[2]. Todas estas operaciones elementales se representan como un conjunto de unidades de proceso homogéneas, interconectadas formando arrays de "n" dimensiones. Dependiendo del diseño del array, las entradas y salidas se realizan a través de

todas o parte de las unidades que se encuentran en los límites del array

Las bases teóricas de la sistolización del proceso de inferencia difusa han sido definidas por diversos autores [3],[4],[5],[6]. Estas han conducido al diseño de múltiples arquitecturas sistólicas de procesamiento.

En esta etapa es adecuado el uso de herramientas de simulación. Estas herramientas permiten:

1. Formalizar la especificación de los diseños. De forma que no solo se formalice la teoría sobre la que estos se basan.
2. Estudiar y comprobar los diseños previamente a su implementación.
3. Testear la ejecución del sistema para comprobar fallos en la teoría de diseño.

La herramienta de simulación utilizada es un simulador de VHDL¹. La adecuación del VHDL como lenguaje de especificación y simulación se justifica en su propia definición [7],[8],[9],[10].

2 Fundamentos Teóricos

A continuación se desarrollan los principios teóricos de los procesadores sistólicos de reglas difusas. Por defecto se consideran conocidas las bases de la lógica difusa [11], [12].

Para desarrollar una arquitectura sistólica la mayoría de los autores [3], [4], [5], [6] realizan una serie de pasos que a continuación se indican:

1. - Definición del conjunto de reglas del modelo a desarrollar. Dichas reglas serán de la forma
IF A_1 is EA_{1y} and/or ... A_n is EA_{ny}
THEN B_1 is EC_{1y} and/or ... B_m is EC_{my}
siendo n el número de antecedentes, EA_{xy} una etiqueta de la variable A_x , m el

número de consecuentes y EC_{xy} una etiqueta de la variable B_x . En ellas se cumple que $\{A_x\} \cap \{B_y\}$ es vacío, es decir, el sistema de reglas es no encadenado [6]. Las palabras clave "and" y "or" representan, respectivamente, la norma y co-norma utilizada [11], [12].

2. - Modificación del conjunto de reglas para que cada una implique una sola variable consecuente. Las reglas del conjunto quedarán

IF A_1 is EA_{1y} and/or ... A_n is EA_{ny}
THEN B_x is EC_{xy}

3. - Normalización de las partes antecedentes para eliminar los operadores "or" y expandir la regla, de forma que intervengan todos los antecedentes. Esto, como se demostrará más adelante, resulta un grave error desde el punto de vista teórico, dando un conjunto de reglas que no es equivalente del inicial.
4. - Agrupación en un mismo conjunto de todas las reglas en las que el consecuente se refiere a la misma variable.

A partir de ahora todo el resto del proceso se referirá al mismo conjunto, realizándose las mismas operaciones para cada conjunto diferente.

5. - Creación de una "base de reglas". Esta base de reglas es una matriz que tiene toda la información sobre las condiciones a cumplir por los antecedentes para que la variable B_x dé como resultado un valor u otro. La creación de esta matriz se realiza de la siguiente forma:

5.1.- La dimensión será el número de antecedentes + 1.

5.2.- La longitud de cada dimensión es el número de etiquetas de la variable asociada.

5.3.- La posición $(x_1, x_2, \dots, x_n, x_{n+1})$ tendrá un 1 o 0 dependiendo de que la regla IF A_1 is EA_{1,x_1} and ... A_n is EA_{n,x_n}

THEN B_x is $EC_{x,x_{n+1}}$

pertenezca o no al conjunto. También es posible considerar valores entre 0 y 1, para considerar la fuzzificación de la conjunción de antecedentes o la introducción de modificadores.

Cabe fijarse con especial atención en el punto 5.2, pues presenta un cambio conceptual importante con respecto al resto de los autores [3], [4], [5], [6], que consideran que la longitud

de cada dimensión es igual al tamaño del universo de discurso de la variable en cuestión. Tras estudios realizados en el laboratorio nos dimos cuenta que es suficiente con utilizar únicamente el número de etiquetas de dicha variable si a la entrada de datos se le incorpora una unidad fuzzificadora. Esto reduce el tamaño de la base de reglas en dos órdenes de magnitud, haciéndola bastante más manejable, además de reducir el número de operaciones a ejecutar de forma considerable, por lo que se acelera la velocidad de la arquitectura diseñada. Tal y como se ha diseñado la arquitectura, la unidad fuzzificadora, no genera ningún retardo que obligue a reducir la velocidad de proceso, por lo que la ganancia obtenida al utilizar esta técnica es elevada.

6. - Generación de la matriz de datos de entrada. La matriz se genera de la siguiente forma:

6.1.- La dimensión será el número de antecedentes.

6.2.- La longitud de cada dimensión es el número de etiquetas de la variable asociada.

6.3.- La posición $(x_1, x_2, \dots, x_n)$ tendrá como valor el mínimo de los grados de pertenencia de los datos de entrada a las etiquetas $EA_{1,x_1} \dots EA_{n,x_n}$.

7. - Multiplicación de ambas matrices para obtener el resultado de la inferencia. Este paso es el que realmente se sistoliza, pudiendo efectuarse de muchas maneras.

8. - Defuzzificación del resultado. Este paso se puede realizar siguiendo cualquiera de los métodos existentes

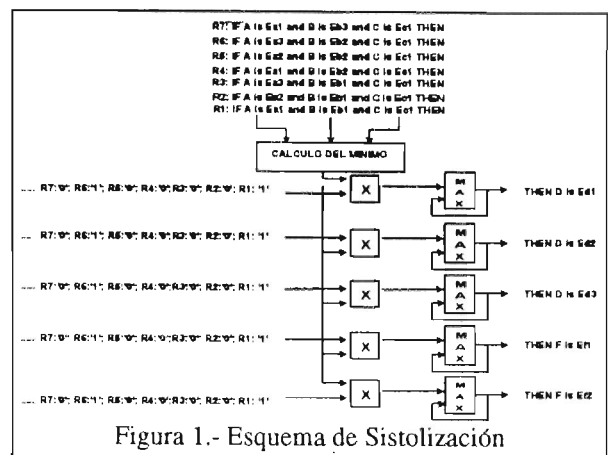


Figura 1.- Esquema de Sistolización

2.1 Sistolización

En nuestro modelo, en vez de multiplicar las dos matrices anteriormente descritas, se reduce en una dimensión la de la base de reglas, generando una matriz con la información de todas las reglas referentes a una sola etiqueta de una variable consecuente. De este modo, el resultado de la multiplicación de ambas matrices da el grado de pertenencia del consecuente a dicha etiqueta.

Para realizar la multiplicación se utilizará una unidad max-min que se describe más adelante, precisándose tantas unidades como etiquetas de variables consecuentes haya.

El esquema que describe el funcionamiento del sistema sistólico se expone en la figura 1.

Dicha figura describe el funcionamiento, de forma general, del sistema sistólico descrito en la parte teórica. Brevemente, el conjunto de reglas posibles es evaluado en forma secuencial. El resultado se introduce en unidades "mínimo" (representadas por X) sincronizadamente con un vector binario. Este vector se corresponde con las reglas que afectan a cada etiqueta consecuente. Los valores resultantes, una vez evaluadas todas las reglas, se defuzzifican.

3 Arquitectura

La arquitectura que se deriva de la sistolización realizada en el apartado anterior, está formada por tres elementos principales:

- 1.- Fuzzificador
- 2.- La unidad de inferencia difusa
- 3.- Defuzzificador

El sistema, a este nivel, aparece como un pipeline de tres segmentos. Esto implica que la ejecución de los tres elementos anteriores se encuentra sincronizada: el retardo máximo de cualquiera de ellas determina el rendimiento del sistema.

En este caso, la limitación de rendimiento viene determinada por la unidad intermedia: la unidad de inferencia difusa.

3.1 Fuzzificador

En la arquitectura desarrollada, el fuzzificador ha sido implementado como un banco de memorias: una memoria por cada variable antecedente. La entrada a la memoria dedicada a un antecedente es un índice que realiza la correspondencia entre el valor del universo exterior en el universo de discurso: un

valor de 8 bits en un universo de discurso de 256 posiciones. Cada entrada referencia a una palabra de longitud igual al número de etiquetas posibles para el antecedente por la precisión en bits utilizada en la evaluación de la función de pertenencia:

Entrada_x:[Etiq.1][Etiq.2]....[Etiq.7]

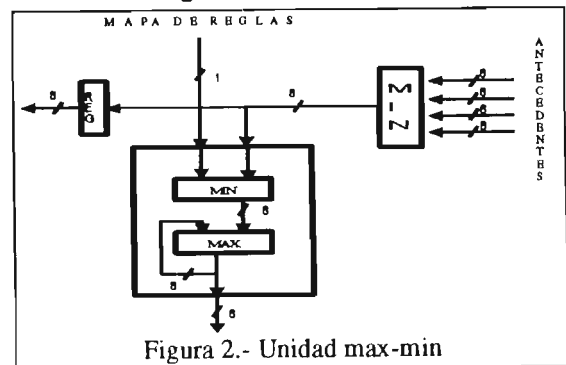
En la simulación de este diseño se utilizan un máximo de 7 etiquetas con tres bits de precisión en la función de pertenencia. De esta forma, el tipo de función de pertenencia que se puede configurar es totalmente libre, sin condicionamiento alguno. Por lo tanto, cada memoria tiene tantas salidas como etiquetas se han definido en el universo de discurso.

3.2 Unidad de Inferencia

La unidad de inferencia difusa aquí desarrollada sigue las pautas marcadas por otros autores [3], [4], [5], [6], pero con una excepción que es esencial en nuestra arquitectura, que es la representación del mapa de reglas tal como se indicaba en la sección 2. Dividimos la unidad de inferencia en tres partes principales:

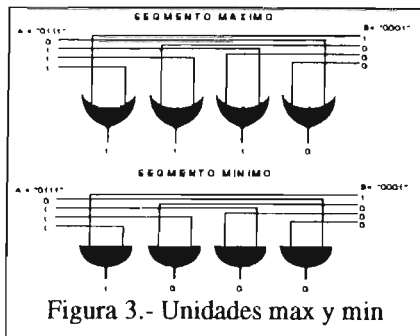
- a.- Array sistólico
- b.- Producto de antecedentes
- c.- Mapa de reglas

El array sistólico está formado por un encadenamiento de unidades "max-min" detalladas en la figura 2.



Cada una de las unidades "max-min" se corresponde con un operador elemental de los descritos en la anterior sección, con la siguiente optimización: en primer lugar, se realiza el mínimo del producto lógico de antecedentes, parte que es común a todas las unidades "max-min". Si bien la función de pertenencia a cada etiqueta se evalúa con tres bits, estos se descodifican en 8 bits, (codificación unaria), lo que permite evaluar

máximos y mínimos utilizando un único nivel de puertas (figura 3) y, por lo tanto, construir



unidades de alta velocidad.

La unidad de producto de antecedentes se construye a partir de un conjunto de registros de desplazamiento. Cada registro tiene la longitud del máximo número de etiquetas que pueden aparecer en una variable antecedente, con máximo establecido en la síntesis del circuito. La dimensión de la salida de los registros es igual a la precisión del la función de pertenencia, en este caso 3 bits. Una circuitería auxiliar de contadores ordena la salida de etiquetas antecedentes así como optimiza el rendimiento del circuito en el caso de que el número de antecedentes sea menor que el máximo admitido.

El mapa de reglas está formado por bancos de memoria. Cada memoria corresponde a una variable consecuente y a una etiqueta, y tiene que proporcionar el flujo de códigos de pertenencia (flujo de 1's y 0's) a la misma velocidad a la que procesan las unidades "max-min". Al ser estas muy veloces, es necesario utilizar un sistema de pipeline que proporcione la necesaria latencia de acceso memoria.

3.3 Defuzzificador

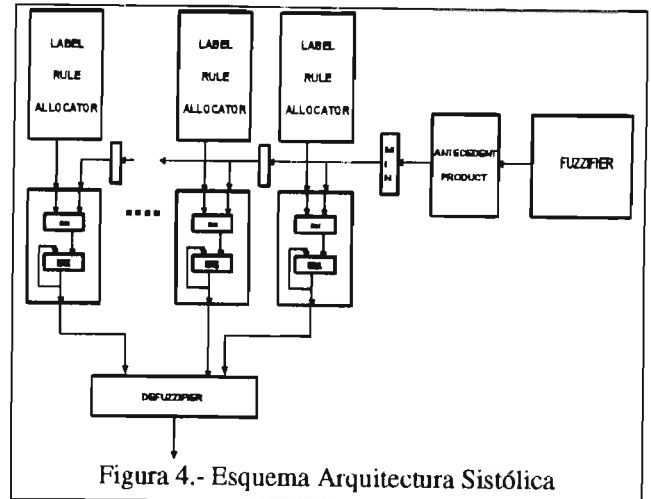
La unidad defuzzificadora permite procesar a la misma velocidad en la que son generados los valores finales en el sistema de inferencia. El método utilizado es de centro de masas. El esquema que compone esta unidad está formado por:

1. Unidad multiplicadora en paralelo [1][2].
2. Divisor en memoria.

Al utilizar ocho niveles de codificación, es posible utilizar una unidad divisora implementada en memoria. Tal implementación incorpora sencillez y velocidad al diseño. Así mismo, el "through-output" del sistema de

defuzzificación se equipara al de la unidad fuzzificadora.

El esquema completo de la arquitectura del sistema sistólico es el que se muestra en la figura 4.



4 Simulación

Entre la descripción de los diseños sobre papel y su implementación efectiva hay un paso imprescindible: la simulación en herramientas CAD. En este punto es cuando afloran los principales problemas ocultos en la concepción de la arquitectura, se determinan tiempos estimados de proceso, de "through-output", etc.

La herramienta utilizada es el lenguaje de descripción hardware VHDL¹. Sus principales ventajas son: su estandarización [7] [8], modularidad, capacidad de síntesis sobre "standard cells" o FPGA's y su parametrización. Esta última propiedad permite definir de forma flexible distintos parámetros de un circuito. En el caso que nos ocupa, los parámetros establecidos son:

1. Número de antecedentes.
2. Número de consecuentes.
3. Etiquetas por antecedente.
4. Etiquetas por consecuente.
5. Máximo número de etiquetas por antecedente.

Las librerías utilizadas, que han sido desarrolladas en el mismo laboratorio, corresponden a las células de 1um de ES2².

Al depender el tiempo de ejecución en la unidad de inferencia del número de etiquetas por antecedente, la circuitería se ha adaptado para especificar el número exacto de etiquetas

utilizadas. De esta forma obtener las máximas prestaciones para cada configuración.

La simulación ha sido realizada generando un circuito con los siguientes parámetros fijados:

1. Cuatro antecedentes de hasta siete etiquetas.
2. Siete etiquetas consecuentes.
3. Cada universo de discurso con 256 muestras.
4. Tres bits de precisión en la evaluación de pertenencia.
5. Hasta 8 Kbytes de memoria

Dicho circuito ha sido programado para mantener cuatro antecedentes y tres etiquetas por antecedente. El rendimiento del mismo se encuentra a continuación:

1. Posibilidad de ejecutar el sistema con un periodo de reloj de 10 ns.
2. Ejecución de la totalidad de la base de reglas en 810 ns.
3. Rendimiento total del sistema de 1.2 MFLIPS³.
4. Frecuencia de ejecución de cada regla de 100 MHz.
5. Superficie de silicio ocupada de aproximadamente 40-60 mm².
6. Tiempo de programación de la base de reglas y fuzzificador de 46.4 us.

Por lo tanto, la arquitectura mantiene rendimientos muy elevados si se compara con otros desarrollos, incluso no sistólicos [13], [14], [15] y teniendo en cuenta que este no es el rendimiento máximo del sistema. Así mismo, la utilización de memoria es pequeña, al contrario que en los primeros sistemas sistólicos, y la superficie total del circuito no es muy grande.

5 Problemas Asociados a los Sistemas Sistólicos

En el desarrollo del presente trabajo en el laboratorio han aparecido disfunciones de carácter conceptual y ciertos puntos críticos en simulación de las arquitecturas definidas. Por lo tanto, a continuación se analizan divididos en: problemas teóricos (los determinantes) y problemas de implementación.

5.1 Problemas Teóricos

Una de las características y ventajas que se consideran propiedad de los sistemas sistólicos es su capacidad para ejecutar todas las reglas posibles del sistema difuso [3],[4],[5],[6].

Aunque este aspecto es considerado como obvio por todos los autores⁴, tropieza con el siguiente problema:

- Supongamos un sistema de control con dos variables de entrada [A, B] y una de salida [C].
- Supongamos solamente dos etiquetas por variable [EA₁ EA₂],[EB₁ EB₂],[EC₁ EC₂].
- Supongamos un sistema sistólico.

En este supuesto no es posible representar la regla:

IF A is EA_y THEN C is EC_y

Pues al ser ortogonal el sistema de representación de reglas, la regla más aproximada que puede ser utilizada es:

IF A is EA_y and (B is EB₁ or B is EB₂) THEN C is EC_y

Si bien en lógica clásica (o "crisp") ambas proposiciones son equivalentes, no así en lógica fuzzy. Esto es debido a que la propiedad de complementariedad no existe en el álgebra difusa [11]. De esta forma:

$u_{EB1}(x)$ or $u_{EB2}(x)$ no implica igual a 1

$u_{EB1}(x)$ and $u_{EB2}(x)$ no implica igual a 0

Un caso similar ocurre con una regla en la que existe la negación de un antecedente.

De ahí que lleguemos a la conclusión que los sistemas sistólicos, tradicionalmente definidos o incluso con las mejoras propuestas, ejecutan únicamente un conjunto limitado de reglas: el formado por aquellas reglas que tienen en su parte antecedente el producto lógico de todas las etiquetas antecedentes que forma parte del sistema.

5.2 Problemas de Implementación

Entre los puntos críticos, que a la hora de implementación, surgen en el diseño y simulación de sistemas sistólicos de procesamiento de reglas destacan:

1. El tamaño de las memorias a utilizar crece exponencialmente con el número de antecedentes. Esto limita la utilización de arquitecturas de muchos antecedentes por limitaciones en la superficie de silicio.
2. La velocidad de proceso es altamente dependiente del número de antecedentes y del número de etiquetas por antecedente.
3. Obliga a procesar un gran número de reglas, precisamente en un sistema difuso, uno de cuyos objetivos es la economía de reglas.

4. La utilización de unidades max-min de codificación binaria es extremadamente lenta para hacer efectiva la utilización de sistemas sistólicos como el descrito.
5. La utilización de unidades max-min de codificación en unario, si bien permite mantener un rendimiento del sistema lo suficientemente elevado.

Todos estos aspectos afectan a sistemas de control difuso muy complejos, aquellos con un gran número de antecedentes y de etiquetas por antecedente. En cambio, para las aplicaciones habituales, constituye una opción con alto rendimiento y sin un alto coste de superficie de silicio.

Conclusiones

En resumen, este trabajo presenta dos aspectos fundamentales de las arquitecturas sistólicas de procesamiento difuso.

En primer lugar, la optimización que supone utilizar el producto de las etiquetas de cada variable antecedente, en vez de el producto de los universos de discurso. Esta modificación supone colocar a los sistemas sistólicos en una posición realmente competitiva a nivel de rendimiento, así como un significativo ahorro en superficie de silicio utilizado, respecto a los anteriores desarrollos.

El segundo aspecto ha sido descubrir una limitación intrínseca en la capacidad de proceso de los sistemas sistólicos. Contrariamente a lo comúnmente establecido, estos sistemas sistólicos no realizan la ejecución de todas las posibles reglas del sistema. Muy al contrario, solo es posible la ejecución de un subconjunto limitado de ellas: aquellas que su parte antecedente esta formada por conjunción de todos los antecedentes. Esto es debido a que el álgebra fuzzy carece de la propiedad de complementariedad, que sí existe en lógica clásica.

Independientemente de la oportunidad de usar procesadores como los descritos, queda un importante trabajo de desarrollo sobre estos diseños. Las optimizaciones han de estar orientadas a la expansión del conjunto de reglas utilizado y a la mejora de la relación rendimiento - número de antecedentes. Este trabajo se está realizando actualmente en el Laboratorio de Hardware y Control Avanzado-INTA.

Notas

- (1) El V/System de Model Technology
- (2) European Silicon Structures, librería con especificaciones industriales.
- (3) Fuzzy Logic Inferences Per Second.
- (4) De hecho no es justificado de modo formal en ninguno de ellos.

Referencias

- [1] Hwang, Kai. Computers Arithmetic. *John Wiley & Sons*
- [2] Hwang, Kai. Advanced Computer Architectures. *McGraw-Hill*
- [3] M.A. Manzoul and S. Tayal. Systolic VLSI Array for Muti-variable Fuzzy Control Systems. *Cybernetics and Systems: An International Journal*, 21
- [4] M.A. Manzoul. Fuzzy Inference on a Systolic Array. *Proceedings of 18th Modeling and Simulation Conference*
- [5] F. Fernandez, A. Ruiz, J. Gutierrez. Diseño Sistemático de un Procesador Sistólico de Inferencias Difusas. *II Congreso Español sobre Tecnologías y Lógica Fuzzy*
- [6] Bugarín, A. y Barro, S. y Ruiz, R. Soluciones Sistólicas en Control Borroso. *II Congreso Español sobre Tecnologías y Lógica Fuzzy*
- [7] IEEE Std 1076-1987. IEEE Standard VHDL Language Reference Manual. *IEEE Inc.*
- [8] IEEE Standards Interpretations. IEEE Standard VHDL Language Reference Manual. *IEEE Inc.*
- [9] Navabi, Z. VHDL, Analysis and Modeling of Digital Systems. *McGraw Hill*
- [10] Coelho, D. The VHDL handbook.
- [11] G.J. Klir and T. Foller. Fuzzy Sets, Uncertainty and Information. *Prentice-Hall*
- [12] B. Kosko. Neural Networks and Fuzzy Systems. *Prentice-Hall*
- [13] H. J. Zimmermann. Fuzzy Technologien. Prinzipien, Werkzeuge, Potentiale. *VDI Verlag.*
- [14] M. Sasaki, F. Ueno, T. Inoe. 7.5MFlips Fuzzy Microprocessor Using SIMD and Logic-in-Memory Structure. *2nd IEEE International Conference on Fuzzy Systems. 1993*
- [15] V. Catania, A. Puliafito, M. Russo, L. Vita. A VLSI Fuzzy Inference Processor Based on a Discrete Analog Approach. *IEEE Transactions on Fuzzy Systems- May 1994*

Controlador Difuso Tipo "Singleton" en Modo de Corriente

I. Baturone, S. Sánchez Solano, A. Barriga, J. L. Huertas

Dpto. de Diseño de Circuitos Analógicos

Centro Nacional de Microelectrónica - Universidad de Sevilla

Edificio CICA. Avda. Reina Mercedes S/N. 41012-Sevilla, SPAIN

Tf: (95) 462 38 11 FAX: (95) 423 18 32

e-mail: lumi@cnm.us.es

Resumen

Se describe la realización de un controlador difuso que implementa el mecanismo de inferencia simplificado utilizando circuitos analógicos CMOS en los que las señales son representadas mediante intensidades. No obstante, el sistema completo se comunica con el exterior mediante tensiones, con objeto de permitir una interfaz sencilla con la circuitería de control convencional. Se discute el método de inferencia empleado, así como distintas alternativas para la realización de los circuitos de generación de los antecedentes, los operadores Min/Max y la etapa de ponderación de reglas. La elección de las soluciones más eficaces para cada bloque del controlador permite construir un sistema óptimo en términos de velocidad de operación y consumo de área y potencia.

1 Introducción

La realización eficiente de sistemas microelectrónicos para aplicaciones de control basadas en lógica difusa requiere una cuidadosa elección, tanto de los mecanismos de inferencia utilizados, como de las estructuras de circuito que implementan cada uno de los bloques del controlador. En relación con los mecanismos de inferencia es necesario asumir una serie de simplificaciones que faciliten la realización de los circuitos. Por este motivo, la inmensa mayoría de los sistemas reportados en la literatura utilizan como funciones de pertenencia triángulos y trapecios con parámetros programables, restringen los conectivos de conexión entre antecedentes a las operaciones mínimo y máximo, y utilizan funciones de implicación con propiedades de T-normas (que permiten una estrategia de resolución regla a regla de la inferencia evitando la necesidad de calcular las matrices de implicación inducidas por cada regla).

De la variedad de métodos de inferencia propuestos, el método *singleton* o método simplificado es el más adecuado para implementaciones *hardware*, a la vez que ofrece la suficiente libertad en la elección de antecedentes y consecuentes como para asegurar un control eficiente. En los siguientes apartados de esta comunicación analizaremos brevemente los distintos mecanismos de inferencia para comprobar cómo el método simplificado surge como caso particular en

todos ellos. A continuación identificaremos los bloques básicos de un controlador tipo *singleton*. Finalmente, utilizando técnicas CMOS analógicas en modo de corriente, describiremos las realizaciones elegidas para los distintos operadores, comparándolas con otras soluciones.

2 Mecanismos de Inferencia Fuzzy

De acuerdo con los mecanismos de inferencia que implementan, los distintos controladores difusos pueden clasificarse en tres grupos que difieren entre sí en el tipo de consecuentes utilizados:

Grupo 1.

Engloba los controladores que emplean consecuentes descritos por funciones de pertenencia trapezoidales o triangulares o con forma de campana. En este caso, la conclusión inferida al aplicar cada una de las reglas es un conjunto difuso C_i descrito mediante una función de pertenencia $\mu_{C_i}(x)$. Al considerar la acción conjunta de todas las reglas la solución global puede obtenerse aplicando dos estrategias diferentes: 1) calcular el valor más significativo de cada uno de los C_i y asignarle un peso para obtener el resultado como una media ponderada de las conclusiones parciales; 2) agregar todos los C_i mediante el operador máximo y proporcionar como salida el valor más significativo del conjunto difuso resultante. La ponderación de los C_i puede realizarse en base al área de la superficie que encierra $\mu_{C_i}(x)$ con el eje x (S_i), al valor máximo que toma $\mu_{C_i}(x)$, o al grado de activación de la regla a la que corresponden (H_i). Como valores más significativos de un conjunto difuso, los más empleados son el centro de gravedad (cdg) y el punto de máxima pertenencia (pmp).

El coste de implementación de las dos estrategias anteriores es diferente. Si la *defuzzificación* se lleva a cabo eligiendo los cdg para representar los conjuntos difusos, la primera estrategia requiere una operación de división por cada regla, aparte de la operación final de división para ponderar todas las reglas. La aplicación de la segunda estrategia da lugar a los controladores de tipo Mamdani, en los que sólo se calcula el cdg del conjunto C resultante. La salida de este tipo de controladores viene dada por la expresión:

$$y_o = \frac{\sum_x x \mu_{C_i}(x)}{\sum_x \mu_{C_i}(x)} \quad (1)$$

Grupo 2.

A este grupo pertenecen los controladores cuyos consecuentes se describen con funciones de pertenencia monótonas no decrecientes. El método de inferencia que aplican, denominado método de Tsukamoto, emplea el parámetro H_i como peso para cada regla y, como punto representativo de cada C_i , aquél para el que μ_{C_i} vale H_i . Si denotamos este punto por α_i , la expresión para la salida del controlador es:

$$y_o = \frac{\sum_i H_i \alpha_i}{\sum_i H_i} \quad (2)$$

Los cálculos requeridos son similares a los del método anterior, sin embargo, para obtener α_i es preciso conocer la función inversa de μ_i .

Grupo 3.

Los controladores de este tipo fueron propuestos por Takagi y Sugeno y se caracterizan por utilizar como consecuentes funciones de las variables de entrada. Las conclusiones proporcionadas por las diferentes reglas coinciden con los valores que toman dichas funciones para un determinado valor de las entradas. Como peso de cada regla también usan el parámetro H_i . La salida del controlador se calcula como:

$$y_o = \frac{\sum_i H_i f_i(\hat{x})}{\sum_i H_i} \quad (3)$$

3 Método Simplificado

En los controladores asociados al grupo 1, cuando los consecuentes vienen descritos por funciones de pertenencia simétricas y se utilizan funciones de implicación que mantienen la simetría (implicaciones de

Mamdani o Larsen) el centro de gravedad del consecuente C_i coincide con el punto en que μ_{C_i} alcanza el grado de pertenencia máximo (c_i). Una vez establecidos los c_i , el método más eficiente para realizar la inferencia es ponderarlos por el grado de activación de la correspondiente regla. La salida del controlador responde a la expresión:

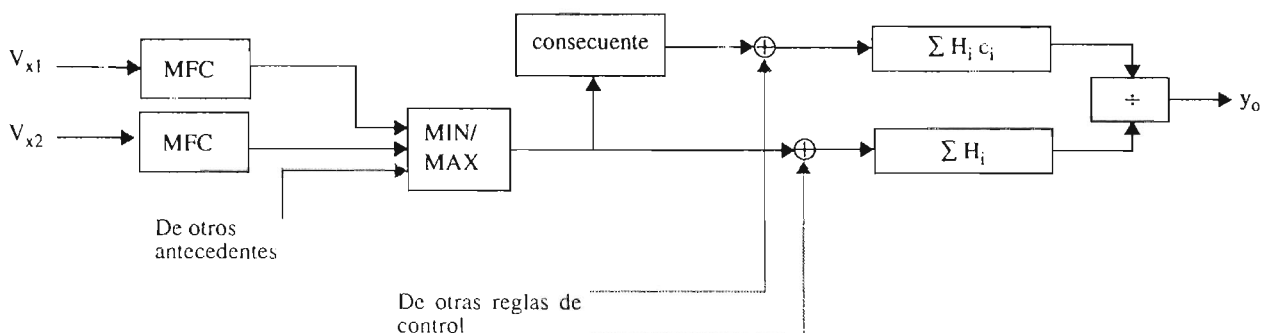
$$y_o = \frac{\sum_i H_i c_i}{\sum_i H_i} \quad (4)$$

El resultado inferido por el método simplificado coincide con el obtenido por el mecanismo de Mamdani cuando en el cálculo de la expresión (1) se tienen en cuenta las zonas de solapamiento provocadas por la agregación de consecuentes. Puede comprobarse asimismo que las expresiones (2) y (3) coinciden con la (4) cuando, tanto las funciones monótonas no decrecientes usadas por Tsukamoto, como las funciones de las variables de entrada utilizadas por Takagi y Sugeno se reducen a una constante de valor c_i .

El mecanismo de inferencia simplificado puede por tanto considerarse como caso particular de los tres tipos de controladores descritos. Sin embargo, su coste de implementación es muy inferior ya que, como se muestra en el diagrama de la Fig. 1, los únicos bloques necesarios son circuitos generadores de funciones de pertenencia que generan los antecedentes de las distintas reglas, operadores Min o Max para implementar los conectivos entre dichos antecedentes, y un bloque de ponderación de reglas que realiza la expresión (4). En los siguientes apartados discutiremos con detalle las distintas opciones de realización de cada uno de estos bloques.

4 Circuitos Generadores de Funciones de Pertenencia (MFCs)

En los estudios teóricos sobre controladores difusos suelen considerarse funciones de pertenencia no lineales del tipo gaussianas o con forma de campana. Sin embargo, por motivos análogos a los que conducen a la adopción del método simplificado, en las aplicaciones prácticas se aproximan habitualmente por funciones



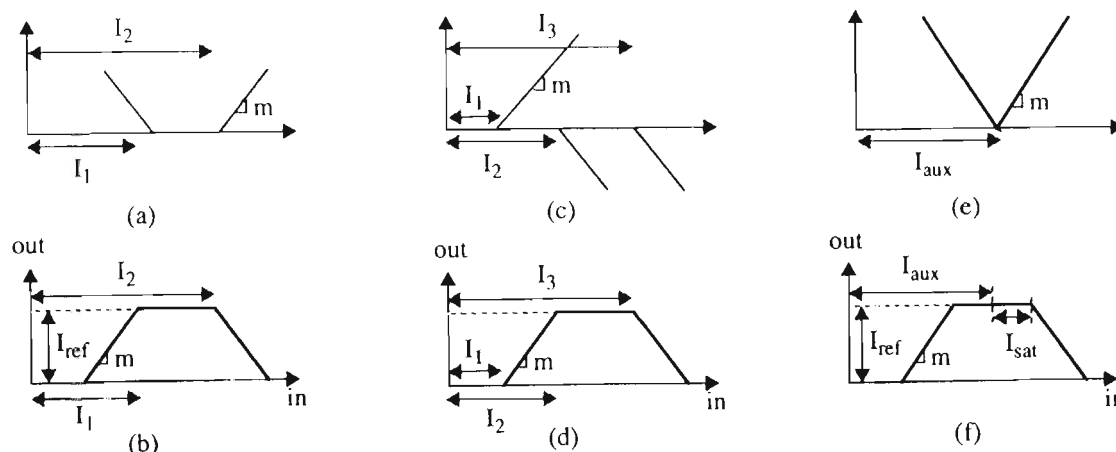


Fig. 2.- Técnicas de generación de funciones de pertenencia triangulares o trapezoidales.

lineales a tramos. Esto se traduce en la elección de triángulos o trapecios cuyas bases (o pendientes) sean fácilmente programables.

Las operaciones básicas para conseguir funciones lineales son la suma, la resta y la rectificación o diferencia acotada. Esta última se define como:

$$a \ominus b = \begin{cases} a - b & \text{si } a > b \\ 0 & \text{en otro caso} \end{cases}$$

El trabajar con intensidades reduce las operaciones de suma y resta a simples conexiones de cables. En cuanto a la rectificación, la mayoría de los circuitos difusos la realizan directamente a través de un MOS flotante en configuración de diodo o mediante los transistores de entrada de un espejo de corriente [1].

En los circuitos MFCs reportados en la literatura se han adoptado tres metodologías para la generación de funciones lineales a tramos tipo trapecio o triángulo. La primera de ellas consiste en sumar dos funciones rectificadas como las indicadas en la Fig. 2a y calcular la diferencia acotada entre el resultado y un valor máximo (I_{ref} en la Fig. 2b) [1]. Matemáticamente, la transformación realizada es:

$$I_{out} = I_{ref} \ominus [m(I_{in} \ominus I_2) - m(I_1 \ominus I_{in})] \quad (5)$$

La técnica usada en [2] combina tres funciones rectificadas (una por cada punto de ruptura) como se indica en las Fig. 2c y 2d. Esto se corresponde con;

$$I_{out} = m(I_{in} \ominus I_1) - m(I_{in} \ominus I_2) - m(I_{in} \ominus I_3) \quad (6)$$

Por último, la metodología propuesta por los autores e ilustrada en las Fig. 2e y 2f emplea un operador valor absoluto (resultado de sumar una señal rectificada con la opuesta de su rectificación) y otras dos rectificaciones posteriores (con I_{sat} e I_{ref}) [3]. De acuerdo con la expresión:

$$I_{out} = I_{ref} \ominus m(| I_{in} - I_{aux} | \ominus I_{sat}) \quad (7)$$

Los circuitos basados en esta última técnica proporcionan mayor precisión, ya que no dependen de la com-

binación de dos o tres funciones como ocurre en las otras dos aproximaciones. Además son más eficientes en cuanto a ocupación de área, consumo de potencia y velocidad, y es posible incluir programabilidad de forma sencilla [3]. Por todo esto son los circuitos empleados para representar los antecedentes en el controlador que aquí se describe.

Para dotar a este MFC de una entrada en tensión, se requiere un convertidor V-I. El esquema completo se ilustra en la Fig. 3, en la que los espejos de corriente tipo N y P se han representado por los bloques cN y cP respectivamente. El conjunto formado por M_1 , M_2 y el bloque cN actúa como un operador valor absoluto, proporcionando una intensidad de salida como la mostrada en la Fig. 2e. El bloque cP realiza la diferencia acotada con I_{sat} y fija el valor de la pendiente m . La intensidad de salida I_{out} representa el complemento del valor de pertenencia de la señal V_{in} . Para obtener la función de pertenencia de la Fig. 2f falta por realizar la diferencia acotada con I_{ref} . No obstante, como veremos en el siguiente apartado, las implementaciones en modo de corriente de circuitos máximo son menos costosas que las de mínimo (generalmente usados como conectivos de los antecedentes en las reglas). Por ese motivo, resulta más eficiente combinar los complementos de los valores de las funciones de pertenencia en un operador máximo y posteriormente proceder a la rectificación con I_{ref} .

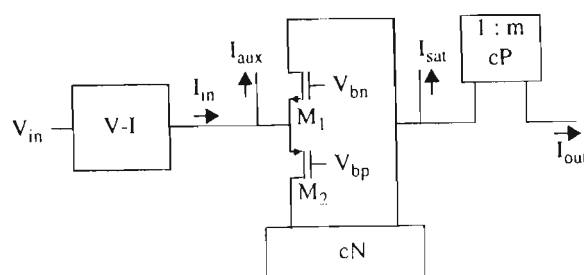


Fig. 3.- Circuito MFC con entrada en tensión y procesado en intensidad.

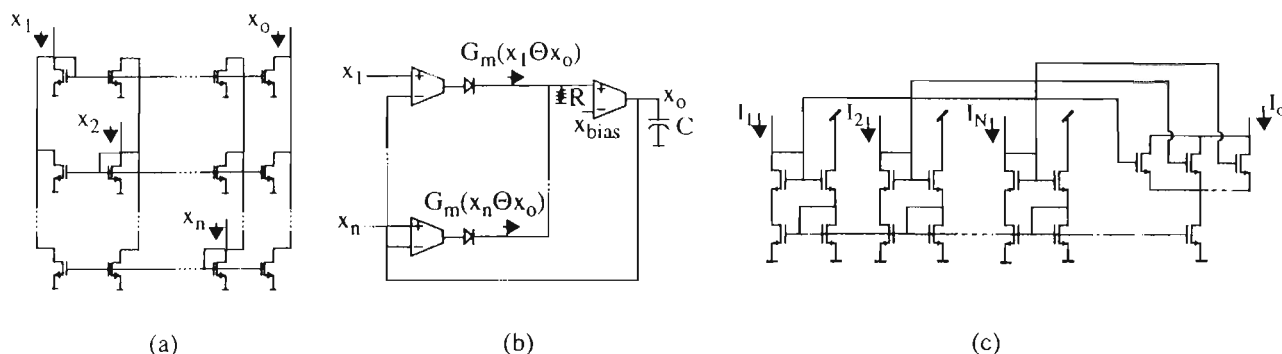


Fig. 4.- Técnicas de generación de máximos de múltiples entradas.

5 Circuitos Conectivos (MIN/MAX)

Los operadores Min/Max se presentan conjuntamente en este apartado porque sus realizaciones son complementarias o bien uno de ellos puede obtenerse a partir del otro mediante la relación de De Morgan:

$$\text{Min}(I_i) = \overline{\text{Max}(\bar{I}_i)} = I_{\text{ref}} - \text{Max}(I_i)$$

donde I_{ref} representa el máximo grado de pertenencia.

Los primeros sistemas difusos CMOS emplearon operadores de dos entradas, muy sencillos de conseguir con circuitos en modo de corriente, pero cuya adaptación al caso de múltiples entradas requiere su conexión en cascada (formando un árbol binario), con la consiguiente acumulación de errores y retrasos.

En la búsqueda de operadores máximo multi-entrada, se han propuesto los siguientes algoritmos:

$$x_o = \sum_i (x_i \ominus (x_o - \xi_i)) \quad \text{con} \quad \xi_i = x_i \ominus \sum_j i \xi_j \quad (8)$$

$$x_o = A \sum_i (x_i \ominus x_o) \quad (9)$$

$$x_o = \sum_i x_i H(x_i \ominus x_o) \quad (10)$$

donde x_i son las entradas, x_o la salida, A un número grande y H la función de Heaviside definida como:

$$H(x) = \begin{cases} 1 & \text{si } x > 0 \\ 0 & \text{en otro caso} \end{cases}$$

Desde un punto de vista microelectrónico, la simplicidad de estos algoritmos se evalúa en base a los circuitos que los implementan.

En [4] se analiza el primero de estos algoritmos y se propone una realización en modo de corriente (Fig 4a). Puesto que cada intensidad de entrada debe ser replicada tantas veces como entradas tenga el sistema, resulta una estructura de complejidad n^2 .

El segundo de los algoritmos se propuso en [5], junto con un circuito a base de transconductores (Fig 4b). Realizaciones en modo de corriente han sido discutidas en [6-7]. Son circuitos más simples que los del caso anterior (de complejidad n). Su desventaja es que el error relativo viene dado por $A/(1+A)$, de modo que A debe ser grande. Si la amplificación se consigue con un espejo de corriente [7] se necesitan transistores de geometrías muy grandes.

El tercer algoritmo [8] permite mayor precisión que el anterior, pues la expresión (10) se obtiene a partir de la (9) cuando A tiende a infinito. Una típica realización en modo de tensión es la clásica puerta de diodos, empleada para la operación OR. Para evitar los errores que surgen como consecuencia de las caídas de tensión en los diodos se emplean esquemas realimentados, incluyendo un diodo en el lazo de realimentación de un amplificador operacional [9-10] o utilizando un amplificador diferencial y una llave en vez de un diodo [11]. Cuando este circuito actúa como máximo, su operación consiste en comparar cada tensión de entrada con la de salida, de modo que sólo la tensión mayor cierre su llave. La misma estrategia se sigue en el circuito máximo en modo de corriente de [12], que usa un comparador de intensidades en vez de tensiones.

El problema de los circuitos anteriores (salvo de los que usan amplificadores operacionales, que están compensados internamente) es su tendencia a la inestabilidad, como consecuencia de los numerosos polos parásitos en los lazos de realimentación. Para evitarlo hay que introducir un polo dominante a muy bajas frecuencias. Este problema no aparece en el circuito máximo en modo de corriente descrito en [13] porque en el camino de realimentación entre intensidades de entrada y salida sólo intervienen dos transistores, que son los que consiguen la ganancia A elevada. Como se ilustra en la Fig 4c, para n entradas, este circuito consta de $5n+1$ transistores.

El circuito que emplearemos para la realización de los conectivos de las reglas del controlador difuso descrito en esta comunicación consiste en una modificación del circuito anterior, propuesta recientemente por los autores [14], que utiliza sólo $3n+1$ transistores y consigue la misma precisión, pero pudiendo operar a más velocidad y consumiendo menos área y potencia. Como se muestra en la Fig. 5, consiste en la combinación de espejos de corriente Wilson que comparten un transistor conectado como diodo. Los transistores T_i actúan como seguidores de tensión, a la vez que conducen la intensidad desde un nudo de baja impedancia a otro de alta, evitando errores de carga.

6 Circuitos para Ponderación de Reglas

Desde la perspectiva de los circuitos MOS analógicos, la división ha sido tradicionalmente una operación

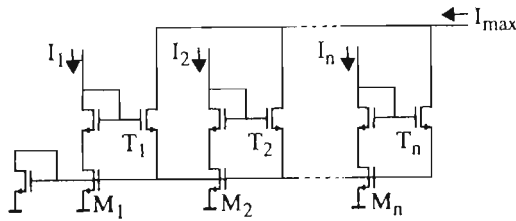


Fig. 5.- Circuito empleado para la operación máximo.

costosa en tiempo y área. Por este motivo, en mucho de los controladores fuzzy propuestos en la literatura se impone la condición de que el denominador de la expresión (4) tome el valor 1, o se calculan los valores normalizados de los grados de pertenencia de acuerdo con la expresión:

$$y_o = \frac{\sum_i H_i}{\sum_i \bar{H}_i} c_i = \sum_i \bar{H}_i c_i \quad (11)$$

Mientras que la operación de normalización puede realizarse de forma eficiente en tecnologías bipolares, no ocurre lo mismo cuando se utilizan dispositivos MOS. Algunos autores recurren al empleo de lazos globales de normalización [13, 15] que presentan dos inconvenientes: 1) implementan sólo de forma aproximada la expresión (11); y 2) limitan la velocidad de inferencia del controlador.

La estrategia seguida en este trabajo se basa en abordar directamente la realización de la expresión (4) mediante un circuito divisor que acepte intensidades como entradas y proporcione una tensión a la salida. Para ello hemos empleado una técnica de transresistencia variable, mediante elementos que permiten una relación $V_o = I_n \cdot R$, con el valor de la resistencia R controlado por una intensidad I_d , de la forma $R = V_a / I_d$.

Cuando se analizan diferentes estructuras con transistores que operan en la región de inversión fuerte (en saturación y zona óhmica), encontramos que el operador divisor más adecuado para la realización del bloque final de controladores difusos es el que emplea transistores en zona óhmica, ya que permite un rango dinámico mayor que el alcanzado con estructuras basadas en transistores en saturación y requiere un consumo de potencia menor [16].

Las geometrías de los transistores que constituyen el elemento de transresistencia variable (mostrado en la Fig. 6a) pueden ser mínimas y no están sujetas a restricciones tan severas como en un transresistor en saturación, que tiene que permitir rangos de tensión suficientes para los espejos de corriente. Además, la combinación de los 4 transistores en zona óhmica puede trabajar a altas frecuencias, pues es una estructura autocompensada en cuanto a sus capacidades parásitas intrínsecas [17], y es más robusta frente a desapareamientos en los transistores. El circuito divisor completo empleado en el controlador que aquí describimos se ilustra en la Fig. 6b. Puede entenderse

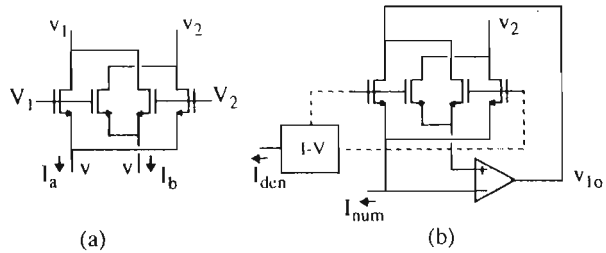


Fig. 6.- (a) Transresistor en región lineal. (b) Circuito divisor.

como un amplificador en configuración inversora realimentado por una resistencia de valor $1/\beta(V_1 - V_2)$, que es a su vez proporcional a $1/I_{den}$. Es decir, hemos aprovechado para realizar el divisor la estructura típica que se emplea para convertir una intensidad en una tensión de salida.

7 Arquitectura del Prototipo Propuesto

Los bloques elegidos para la realización de las funciones de pertenencia, conectivos entre antecedentes (Min/Max) y ponderación de reglas son totalmente compatibles entre sí. La arquitectura del controlador resultante se indica en la Fig. 7. Según apuntamos con anterioridad, su diseño para una aplicación concreta se reduce a la elección de las geometrías de determinados espejos de corriente (los marcados con letra negrilla en la Fig. 7).

La velocidad de operación de los bloques MFC y Min/Max viene limitada por el tiempo invertido en realizar la diferencia acotada (pues ésta conlleva el paso de transistores de corte a conducción). En general, estos tiempos están por debajo de los 300 ns. En cuanto al bloque divisor, su dinámica viene gobernada por un polo dominante cuya magnitud es proporcional al ancho de banda del amplificador y al valor de I_{den} . Para valores típicos el tiempo de respuesta de este bloque también es de unos cientos de nanosegundos. Por todo ello, el controlador propuesto puede alcanzar el rango de las MFLIPS.

8 Conclusiones

La utilización del método de inferencia simplificado y la adecuada elección de los circuitos que implementan los diferentes operadores facilita la realización en hardware de sistemas de control basados en lógica difusa. Las principales características de la arquitectura que proponemos son el reducido coste de área, la elevada velocidad de operación, la facilidad de interfaz con la circuitería de control convencional y la posibilidad de incluir técnicas de programación digital.

Bibliografía

- [1] T. Yamakawa, T. Miki y F. Ueno, "The Design and Fabrication of the Current Mode Fuzzy Logic Semi-Custom IC in the Standard CMOS IC Technology", Proc. IEEE Int. Symp. Multiple-Valued Logic, May 1985, pp. 76-82

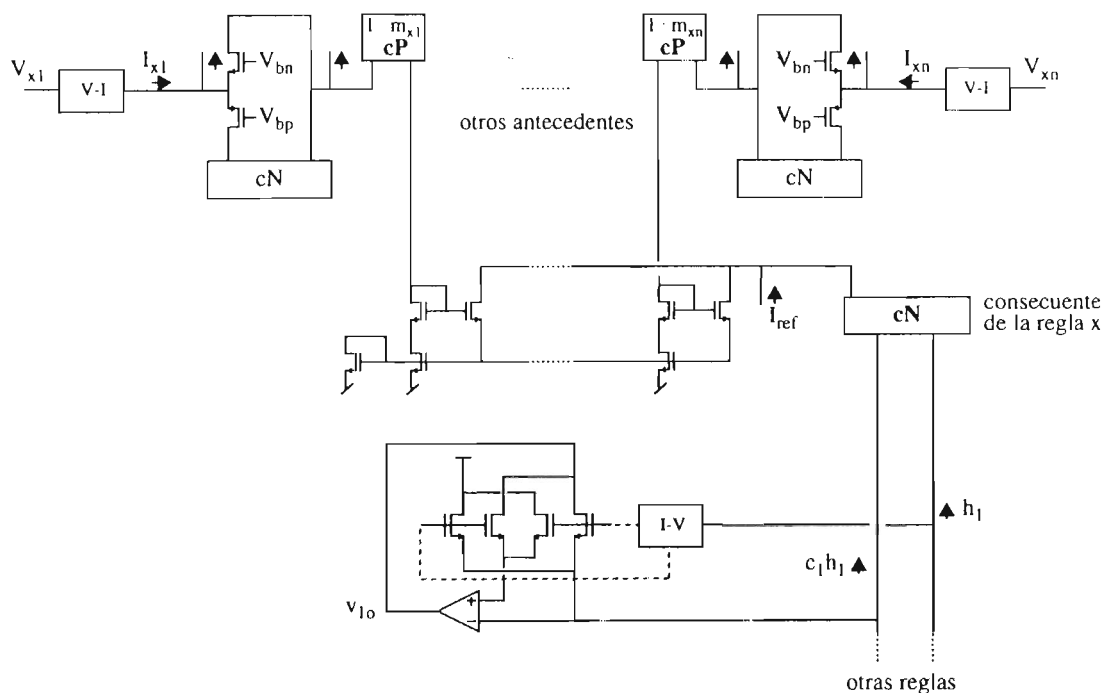


Fig. 7.- Esquemático del controlador propuesto.

- [2] T. Kettner, C. Heite y K. Schumacher, "Analog CMOS Realization of Fuzzy Logic Membership Functions", IEEE Trans., Julio 1993, SC- 28, (7), pp. 857-861.
- [3] J. L. Huertas, S. Sánchez-Solano, A. Barriga e I. Baturone, "Serial Architecture For Fuzzy Controllers: Hardware Implementation Using Analog/Digital VLSI Techniques", 2nd. Int. Conf. on Fuzzy Logic and Neural Networks, Iizuka, 1992, pp. 535-538.
- [4] M. Sasaki, T. Inoue, Y. Shirai y F. Ueno, "Fuzzy multiple-input maximum and minimum circuits in current mode and their analysis using bounded difference equations", IEEE Trans. on Computer, Vol. 39 (6), June 1990, pp. 764-774.
- [5] T. Inoue, F. Ueno y T. Motomura, "Analysis and Design of Analog CMOS Building Blocks for Integrated Fuzzy Inference Circuits", Proc. ISCAS, Singapore 1991, pp. 2024-2027.
- [6] J. L. Huertas, S. Sánchez-Solano, I. Baturone y A. Barriga, "Building Blocks for Current-Mode Implementation of VLSI Fuzzy Microcontrollers", Proc. IFSA, Seúl 1993, pp. 929-932.
- [7] K. Tsukano, T. Inoue y F. Ueno, "A design of current-mode analog circuits for fuzzy inference hardware systems", Proc. ISCAS, May 1993, pp. 1385-1388.
- [8] B. W. Suter y M. Kabrisky, "On a Magnitude Preserving Iterative MAXnet Algorithm", Neural Computation, Vol. 4, pp. 224-233, 1992.
- [9] J. G. Graeme, "Designing with Operational Amplifier", Mc Graw Hill, 1977
- [10] M. Hassoun y A. Nabha, "Implementation of $O(n)$ complexity max/min circuits for fuzzy and connectionist computing", Proc. IEEE Int. Conf. on Neural Networks, Marzo 1993, pp. 998-1003
- [11] T. Kettner, K. Schumacher y K. Goser, "Realization of a Monolithic Fuzzy Logic Controller", Proc. ESSCIRC, pp. 66-69, Sevilla, 1993.
- [12] F. V. Verdú, B. Linares Barranco, A. Gago y A. Rodríguez Vázquez, "Concepto y Diseño CMOS de Operadores Máximo Multi-Entrada en Modo de Corriente", Proc. VIII Congreso de Diseño de Circuitos Integrados, Málaga, 1993, pp. 83-87.
- [13] M. Sasaki, N. Ishikawa, F. Ueno y T. Inoue, "Current Mode Analog Fuzzy Hardware with Voltage Input Interface and Normalization Locked Loop". IEICE Trans. Fundamentals, Vol. E75-A (6), June 1992, pp. 650-654.
- [14] I. Baturone, J. L. Huertas, A. Barriga y S. Sánchez Solano, "Current-Mode Multiple-Input Maximum Circuit", Electronics Letters, Vol. 30 (9), April 1994, pp. 678-680.
- [15] T. Yamakawa, "A Fuzzy Inference Engine in Non-linear Analog Mode and its Application to a Fuzzy Logic Control", IEEE Trans. on Neural Networks, Vol. 4, No. 3, pp. 496-522, Mayo 1993.
- [16] I. Baturone, S. Sánchez Solano, A. Barriga y J. L. Huertas, "Estudio de Circuitos Divisores para Controladores Difusos", pendiente de aceptación en IX Congreso de Diseño de Circuitos Integrados, Las Palmas de Gran Canaria, 1994.
- [17] M. Ismail, "Four-Transistor Continuous-Time MOS Transconductor", Electronics Letters, Vol. 23, pp. 1099-1100, Sept. 1987.

A Design Approach for Neuro/Fuzzy Chips

Fernando Vidal-Verdú

Dept. de Arquitectura y Tecnología de
Computadores y Electrónica. Universidad
de Málaga. Plaza El Ejido s/n
29013 Málaga, Spain.
e-mail: vidal@ctima.uma.es

Angel Rodríguez-Vázquez

Dept. of Analog Design, Centro Nacional
de Microelectrónica. Edificio CICA,
Avda. Reina Mercedes s/n
41012 Sevilla, Spain
e-mail: angel@cnm.us.es

Abstract

This paper outlines a systematic approach to design fuzzy inference systems using analog integrated circuits in CMOS, standard VLSI technologies. Inference is performed by using Takagi and Sugeno's if-then rules, particularly where the rule's output contain only a constant term-- a singleton. All proposed circuits emphasize simplicity at the circuit level-- a prerequisite to increasing system level complexity and operation speed. A three-input, four-rule controller has been designed for demonstration purposes in a 1.6 μ m CMOS single-poly, double-metal technology. We include measurements from prototypes and detailed HSPICE simulations. These results operation speed in the range of 5MFlips with systematic errors below 1%. A hybrid learning algorithm, suitable for hardware implementation, is also proposed, as long as a modified membership function circuit that allows electrical tuning of all parameters, as desired for learning purposes.

Key words: neural networks, fuzzy theory, analog circuits, current-mode circuits, normalization open loop.

1 Introduction

Software implementations of fuzzy inference systems typically operate below Kflip rate (flip stands for fuzzy logic inferences per second), not fast enough for many high-speed control problems, like those related to automotive engines [1]. During the last few years different authors

have focused on the development of dedicated hardware, using IC technology [1], [2], [3], to overcome this drawback. In particular, *analog* circuits are worth considering for this application due to the intrinsically higher *speed* and lower *power* consumption than their digital counterparts [4]. Functional *efficiency* (measured as the device count for a given operation) of analog circuits is also much larger than for digital, due to the possibility of versatile exploitation of small analog devices (formed by a few transistors) for a wide variety of low-level linear and non-linear processing required for fuzzy inference. Finally, the intrinsically lower accuracy of analog circuits does not seem to be a major limitation for most fuzzy system applications.

Previous proposals for analog fuzzy circuitry use bipolar transistors and/or linear resistors [2],[3] -- not readily implementable in the standard VLSI technology. Consequently, these ICs are expensive to produce, and not fully compatible to other conventional digital circuitry that may be needed to integrate together with the fuzzy circuitry, for complex control tasks. To overcome this drawback, all circuits proposed in this paper use MOS transistors as the only circuit primitive, and thus are fully compatible with the cheapest CMOS single- poly scaled technologies.

Since one of the major problems encountered in fuzzy systems is how to capture expertise, our approach focuses on algorithms and circuits that enable the incorporation of learning.

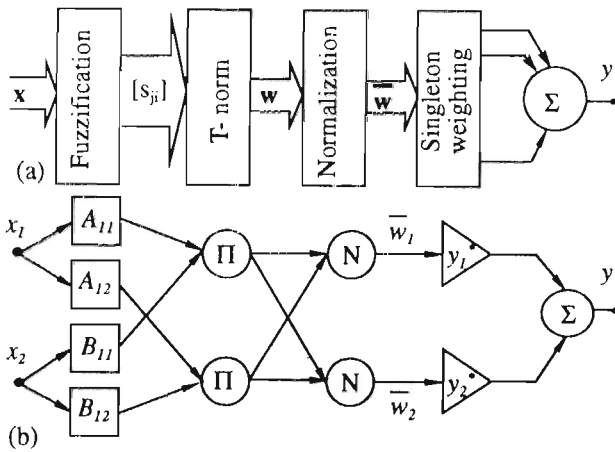


Fig.1: (a) Singleton neuro-fuzzy architecture; (b) Exemplary architecture for two input and two rules.

2 A Neuro-Fuzzy Architecture.

2.1 Multilayer Architecture

Circuit building blocks are arranged in the *layered* architecture of Fig.1 [4], where it is implicitly assumed that inference is performed using Takagi and Sugeno's *singleton* algorithm [5]. Referring to Fig.1 we can identify the catalog of analog functions needed to implement neuro/fuzzy controllers:

Layer 1: It contains a node per each fuzzy label of the input variables. Layer input is $\mathbf{x}^T = [x_1, x_2, \dots, x_M]$, and layer output is the matrix of matching degrees,

$$[s_{ji}] = [\mu_{ji}(x_j)] \quad 1 \leq j \leq M, 1 \leq i \leq N \quad (1)$$

where $\mu_{ij}(x)$ is the membership function associated to the i -th linguistic label of the j -th input variable. Hence, each node in this layer realizes a nonlinear transformation.

Layer 2: It maps this matrix of matching degrees into the vector of firing rule activities, $\mathbf{w}^T = [w_1, w_2, \dots, w_N]$. Each vector component is calculated by a corresponding processing node as follows,

$$w_i = \min(s_{1i}, s_{2i}, \dots, s_{Mi}) \quad (2)$$

Layer 3: Each node in this layer calculates the averaged firing activity of its corresponding

rule, starting from the vector of firing activities, as follows,

$$\bar{w}_i = w_i / \left(\sum_{i=1, N} w_i \right) \quad 1 \leq i \leq N \quad (3)$$

Layer 4: This layer multiplies each component of $\bar{\mathbf{w}}$ by its corresponding singleton, thus obtaining the vector of graded rule's consequents,

$$z_i = \bar{w}_i y_i^* \quad 1 \leq i \leq N \quad (4)$$

Layer 5: This contains a single node which aggregates the consequent outcome of the individual rules, to obtain the inferred output.

$$y(\mathbf{x}) = \sum_{i=1, N} z_i \quad (5)$$

Nodes in layers 1 and 4 are *adaptive*, while the remaining have a fixed function. By setting the parameters that control the shapes and locations (inside the universe of discourse intervals) of the membership functions and the singletons, Fig.1 can learn a prescribed input-output mapping [4].

2.2 Learning Strategy

Jang proposes a hybrid algorithm to train Fig.1 by using *back propagation* for layer 1 and *LMS* for layer 4. However, it is very costly for hardware implementation, which motivates looking for better suited alternatives. A detailed viability analysis taking into account the specific nature of the circuitry and its function leads us to explore *weight perturbation* [6] for layer 1 and *outstar* for layer 4 [7]. This algorithm substitutes the derivatives of *back propagation* by finite difference equations, and tries to keep the complete algorithm feedforward through computation of the influence of each parameter on the global error. Assuming that ω is the learning parameter and E is the global error at the output, changes in ω are given by

$$\Delta\omega = \frac{\eta [E(\omega) - E(\omega + \text{pert})]}{\text{pert}} \quad (6)$$

where *pert* is a small perturbation and η the

learning rate -- both constants. Note that weight update hardware involves evaluation of the error with perturbed and unperturbed weight, and then multiplication by a constant.

On the other hand, functional equivalence between layer 4 and a Grossberg neuron in the *counter propagation* network can provide a good alternative for *LMS*, the *outstar* rule. Singletons are then updated as follows,

$$y_{i_{new}}^* = y_{i_{old}}^* + \mu [T - y(x)] \quad (7)$$

where y_i^* is the singleton whose rule antecedent is maximum, that is $w_i(\mathbf{x}) = \max\{w_1(\mathbf{x}), w_2(\mathbf{x}), \dots, w_N(\mathbf{x})\}$. In this case, update hardware involves a multi-input maximum circuit apart from the multiplication by a constant.

3 CMOS Premise Building Blocks

3.1 Membership Function Circuits (Layer 1)

CMOS PWL membership function circuits have been proposed in [8]. Herein we will consider *bell-like* functions. The exact bell shape [6],

$$\mu(x) = \frac{1}{1 + \left(\frac{x-E}{\Delta}\right)^{2B}} \quad (8)$$

where Δ , E , and B are the adaptation parameters, uses involved analog circuitry [9]. However, membership functions with bell-like shape and continuous derivative are realized in a simple manner by the transconductance mode circuit of Fig.2(a), consisting of two interconnected source-coupled MOS differential pairs. This structure is similar to that used in high-speed folding ADCs [10] and exploits the operation of the MOS differential amplifier as a current switch with soft transition region.

Fig.2(a) obtains a bell-like membership function through a linear, KCL combination of two soft-limiter characteristics: one with positive slope and the other with negative slope as Fig.2(c) illustrates. A square-law model of the MOS transistor, obtains the following expressions for the transition regions in Fig.2(c),

$$\mu(x) = \begin{cases} \sqrt{\beta I_u} x_1 \sqrt{1 - (\beta x_1^2 / I_u)} & -F < x_1 < F \\ (-\sqrt{\beta I_u}) x_2 \sqrt{1 - (\beta x_2^2 / I_u)} & -F < x_2 < F \end{cases} \quad (9)$$

where $x_1 = x - E_1$, $x_2 = x - E_2$, $\beta = kW/L$, $F = (I_u/\beta)^{1/2}$, k is a technological parameter (whose value for NMOS transistors in a typical technology is about $50\mu A/v^2$) and W and L are the transistor width and length. The unitary current I_u in (9) is a normalization value which corresponds to the largest matching degree value ($m_i=1$).

By making $E_1 = E - \Delta$ and $E_2 = E + \Delta$, (9) provides a first-order approximation to the bell-shape of (8) with the slope at the *crossover* [6] points (parametrized by B in (6)) given by:

$$g_1 = \sqrt{\beta_1 I_u} \quad (10)$$

Thus, the membership function width and position can be tuned by proper setting of the reference voltages E_1 and E_2 , while membership function slope is determined by transistor geometries. Electrical tuning of slope, which is desirable for learning purposes, is achieved by using the transconductor 2 of Fig.2b. In this case, the slope at the crossover points is given by:

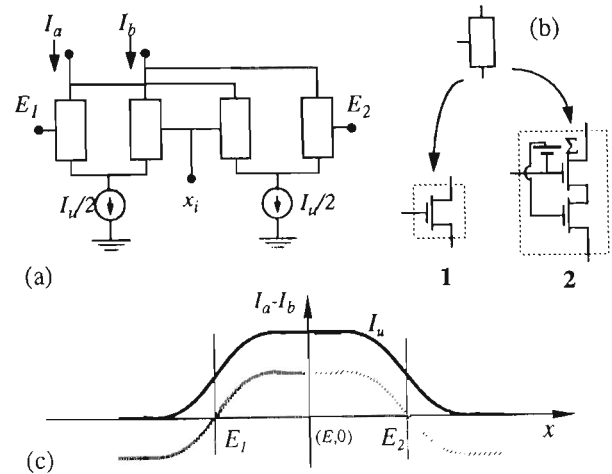


Fig.2: Bell-like CMOS membership function: (a) Circuit structure; (b) Two possible transconductors: simple--1, and composite--2; (c) Membership function shape.

$$g_2 = \frac{\left((\beta_1 \beta_2 / (\beta_1 + \beta_2)) \sqrt{(I_{ss}/\beta_2) + (\sqrt{I_{ss}/\beta_1} + \Sigma)^2} \right) (2 - 2\Sigma \sqrt{\beta_2 / (\beta_1 + \beta_2)})}{\sqrt{(I_{ss}/\beta_2) + (I_{ss}/\beta_1) + 2\Sigma \sqrt{I_{ss}/\beta_1} + \Sigma^2 \beta_2 / (\beta_1 + \beta_2)}} \quad (11)$$

where Σ is the tuning parameter and $I_{ss} = I_u/2$.

3.2 T-Norm Nodes (Layer 2)

The calculation of the minimum among the matching degrees in fuzzy inference rules is functionally equivalent to obtaining the complement of the maximum among the complements of these matching degrees,

$$w = \min(s_1, s_2, \dots, s_M) = \overline{\max(\overline{s_1}, \overline{s_2}, \dots, \overline{s_M})} \quad (12)$$

where the upper bar denotes complement, calculated from the original variable as follows,

$$\overline{z} \equiv 1 - z \quad (13)$$

Note that the complement operator is easily realized in current-mode, by using KCL, with 1 in (13) being the unitary current, I_u . Let us now consider the implementation of the maximum operator. The classical approach used in analog computation for voltage-mode circuits is based on the following steady-state equation,

$$-i_o + \sum_{k=1, M} A u_{-1}(i_{ik} - i_o) = 0 \quad (14)$$

where i_o is the output, i_{ik} are inputs, and $u_{-1}(\cdot)$ denotes *rectification operator*. Fig.3(a) illustrates this concept while Fig.3(b) shows a conceptual CMOS current-mode schematic for it. This circuit is similar to that presented in [11] for the *winner-take-all* operation. However, contrary to the winner-take-all, the circuit of Fig.3(b) is designed not only to select the maximum among a set of input currents, but also to propagate that maximum current to the output node. Fig.3(c) illustrates circuit operation. The maximum current determines the value of the common gate voltage, V_g . The only input transistor that operates in saturation region is that which is driven by maximum input current. All the rest operate in ohmic region.

Fig.3 requires careful analog design to reduce errors due to channel length modulation effects, which appear if transistor output nodes are not equipotential. In particular, our design approach yields 0.3% error for 15 μ A current.

4. CMOS Consequent Building Blocks

4.1 Normalization Circuitry (Layer 3)

Using analog dividers to evaluate (3) is impractical -- analog dividers are costly and inaccurate. A convenient alternative uses feedback to maintain constant a sum of vector components [2], [3], [12]. Unfortunately, transient response of this normalization scheme is rather poor -- a negative consequence of feedback. In particular, it obtains times around 1 μ s (90%) when used in the CMOS 1.6 μ m 3-input 4-rule controller of Section VI.

On the other hand, the normalization circuit operation can be summarized as follows

$$\overline{w}_i = F(w, \overline{w}) \quad (15)$$

and

$$|\overline{w}| = \sum_{i=1, 4} \overline{w}_i = A \quad (16)$$

where $F(\cdot)$ is an increasing monotonic function of w_i and A is a real constant. Fig.5 (illustrated for a case with 4 inputs) presents a circuit which realizes this function without feedback, and

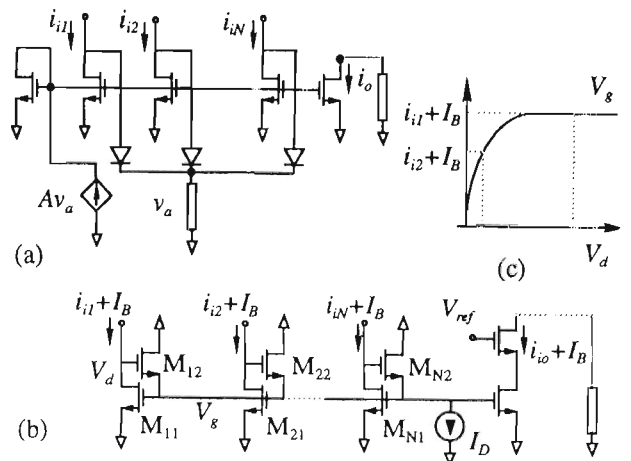
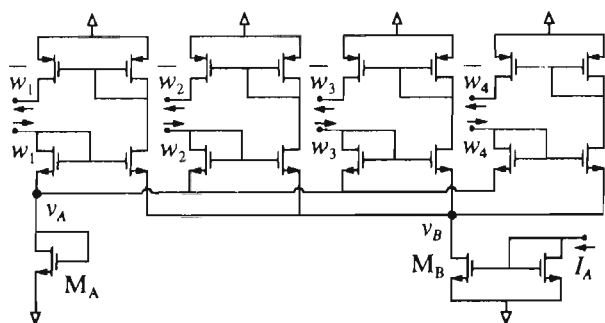


Fig.3: CMOS current-mode maximum/propagate circuit: (a) Concept; (b) Basic schematic; (c) Illustrating operation principle.



hence yields much better transient response than previous proposals. The proposed circuit consists of two source coupled NMOS arrays; the one at the bottom implements a non-linear I/V conversion, and produces a voltage input for each output transistor of the top array, whose drain current is finally replicated by a PMOS current mirror. Square-law calculations on this circuit give

This expression fulfills (13), while (14) is forced by KCL.

Main error sources in Fig.5 are channel length modulation and common mode rejection. Proper design yields 0.8% error using our design approach.

4.2 Output Circuitry (Layers 4 and 5)

building blocks like those proposed for layer 1. Electrical tunability of the membership function circuit is demonstrated by measurements from real prototypes (Fig.8a,b) and from HSPICE simulations using level 6 transistor models (Fig.8c).

6 Conclusions

A set of innovative building blocks for analog fuzzy controllers design has been presented. The simplicity of these blocks, as well as the compact design they achieve (less than 3 mm² for a three input-four rules controller), permit operation speeds of up to 5MFlips with low power consumption. Besides this, the great modularity of this approach enables increased controller complexity, by adding rules and/or input, with no extra design effort. Computer simulations have also demonstrated the viability of on-chip learning in a neuro-fuzzy system, which justifies the interest in adaptive building blocks, like those presented here for layer 1.

Acknowledgments:

To Manuel Delgado-Restituto for fruitful discussions and for providing the results in Fig.8a,b. To Manuel Caracul Urbano for providing computer simulations in Fig.7

References

- [1] K. Namakura et al.: "Fuzzy Inference and Fuzzy Inference Processor". *IEEE Micro*, pp. 37-48, Oct. 1993.
- [2] T. Miki et al.: "Silicon Implementation for a Novel High-Speed Fuzzy Inference Engine: MFlips Analog Fuzzy Processor". *Journal of Intelligent and Fuzzy Systems*, Vol. 1, pp. 27-42, 1993.
- [3] T. Yamakawa: "A Fuzzy Inference Engine in Nonlinear Analog Mode and Its Application to a Fuzzy Logic Control". *IEEE Trans. on Neural Networks*, Vol. 4, pp. 496-522, May 1993.
- [4] J.S.R. Jang and C.T. Sun: "ANFIS: Adaptive-Neuro-Based Fuzzy Inference System". *IEEE Trans. on Systems, Man and Cybernetics*, Vol. 23, pp. 665-685, May 1992.
- [5] T. Takagi and Sugeno: "Derivation of Fuzzy Control Rules from Human Operator's Control Action". *Proc. of the IFAC Symp. on Fuzzy Information, Knowledge Representation and Decision Analysis*, pp. 55-60, July 1989.
- [6] M. Jabri and B. Flower: "Weight Perturbation: An Optimal Architecture and Learning Technique for Analog VLSI Feedforward and Recurrent Multilayer Networks". *IEEE TRANSACTIONS ON NEURAL NETWORKS*, Vol. 3, No. 1, pp. 154-157, January 1992.
- [7] E. Sánchez-Sinencio and Lau (editors): "Artificial Neural Networks", IEEE Press 1992.
- [8] A. Rodríguez-Vázquez and M. Delgado-Restituto: "Generation of Chaotic Signals using Current-Mode Techniques". *Journal of Intelligent and Fuzzy Systems*, Vol.2, January 1994 (to appear).
- [9] C. Toumazou et al. (editors): "Analog IC Design: The Current Mode Approach". Peter Peregrinus 1990.
- [10] J. Van Valburg and R.J. Van de Plassche: "An 8-b 650Mhz Folding ADC". *IEEE Journal of Solid-State Circuits*, Vol.39, Dec. 1992.
- [11] J. Lazzaro, R. Ryckebusch, M. A. Mahowald, and C. A. Mead, "Winner-take-all Networks of O(n) Complexity", *Advances in Neural Information Processing Systems*, Vol. 1, D. S. Touretzky, Ed. Los Altos, CA: Morgan Kaufmann, 1989.
- [12] M. Sasaki et al.: "Current-Mode Analog Fuzzy Hardware with Voltage Input Interface and Normalization Locked Loop". *IEICE Trans.-Fundamentals*, Vol. E75-A, pp. 650-654, June 1992.

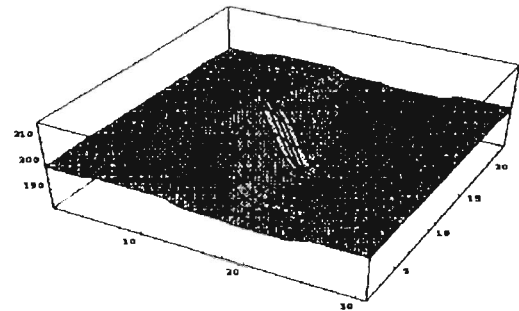


Fig.6 Controller output for two active input and different singleton values.

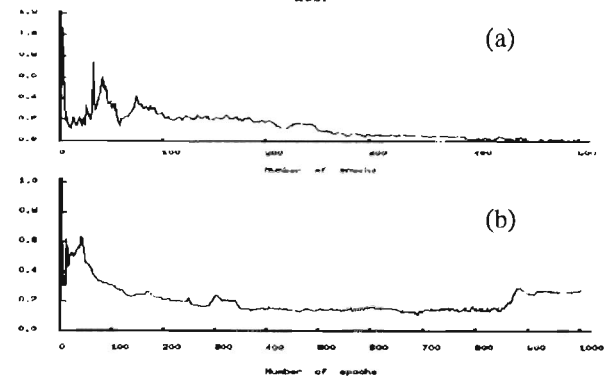


Fig.7 Learning Algorithms: (a) RMSE curve for weight perturbation algorithm with $\eta/pert=7.5$ and $pert=0.001$; (b) RMSE curve for hybrid algorithm with $\eta/pert=7.5$, $pert=0.001$, and $\mu=0.01$.

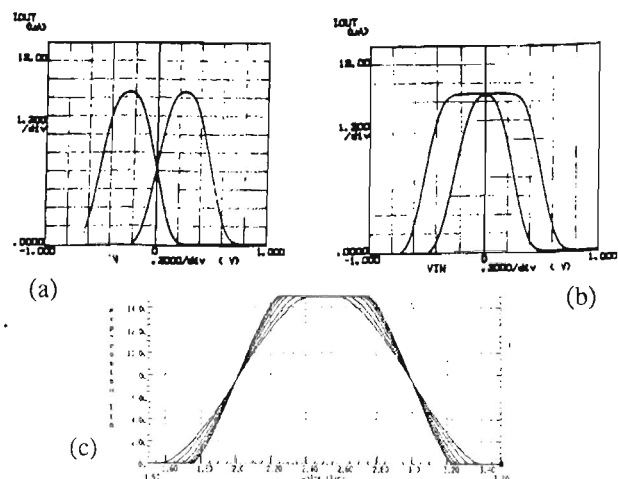


Fig.8 Membership function tunability: (a) Location tunability; (b) Width tunability; (c) Slope tunability with transconductor 2 in Fig.2b.