



**CENTRO DE CIENCIAS APLICADAS Y
DESARROLLO TECNOLÓGICO**
UNIVERSIDAD NACIONAL AUTÓNOMA DE MÉXICO



REPORTE TECNICO

IT-TINF-2009-01

LOGÍCA BORROSA DINÁMICA

**José Luis Pérez Silva
Antonio Garcés Madrigal
Francisco Caviedes Contreras
Andrea Miranda Vitela**

OCTUBRE DEL 2008

INDICE

CAPÍTULO 1.

1.1. INTRODUCCION	5
1.2. DUBITADORES LEIBNIZIANOS	9
1.3. DUBITADORES KANTIANOS	10
1.3. DUBITADORES DE HENGEL	10
1.4. RAZONAMIENTO POR APROXIMACIÓN	11
1.5. MODELOS CAUSALES	12
1.6. LÓGICA FORMAL	13
1.7. TEORÍAS DE PROBABILIDAD	15
1.8. TEORÍA DE LA EVIDENCIA.	16
1.9. TEORÍA DE LOS CONJUNTOS BORROSOS	16
1.10. CONJUNTOS BORROSOS DINÁMICOS	17

CAPÍTULO 2

2.1. LOGICA BORROSA	20
2.2. BORROSIDAD	25
2.3. LÓGICA BORROSA DE PRIMER ORDEN	26
2.4. FORMULISMO	27
2.5. TEORÍAS DE LÓGICA BORROSA DE PRIMER ORDEN	29

CAPÍTULO 3

3.1. RAZONAMIENTO APROXIMADO	33
3.2. REGLAS DE TRASLACIÓN DEL RAZONAMIENTO APROXIMADO	33
3.3. REGLAS DE INFERENCIA EN RAZONAMIENTO APROXIMADO	36

CAPÍTULO 4

4.1. CONJUNTOS BORROSOS DINAMICOS	41
4.2. NOTACIÓN, TERMINOLOGÍA Y OPERACIONES BÁSICAS.	41
4.3. OPERACIONES ALGEBRAICAS.	47
4.4. RELACIÓN BORROSA DINÁMICA.	49
4.5. PROYECCIÓN Y CONJUNTOS BORROSOS DINÁMICOS CILÍNDRICOS	50
4.6. EL PRINCIPIO DE EXTENSIÓN	50
4.8. CONJUNTOS DINÁMICOS BORROSOS DEL TIPO II	52

CAPÍTULO 5

5.1. LÓGICA BORROSA DINÁMICA	54
5.2. EL EFECTO DE LA BORROSIDAD EN HECHOS Y REGLAS.	68

CAPÍTULO 6

6.1. INFERENCIA EN LOGICA BORROSA DINAMICA.	71
6.2. TIPOS DE PROPOSICIONES.	72
6.3. DEDUCCION.	75
6.4. INFERENCIA DE PROPOSICIONES CUANTIFICADAS.	78
6.5. CARDINALIDAD DE UN CONJUNTO BORROSO.	80
. EL SILOGISMO INTERSECCION/PRODUCTO	80
6.8. CADENAS DE PROPOSICIONES.	81
6.9. SEMANTICA Y PRAGMATICA DE LOS SISTEMAS BORROSOS.	82
6.10. RESTRICCIONES BORROSAS DINÁMICAS	86
6.11. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO I	86
6.12. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO II	87
6.13. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO III	88
6.14. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO IV	88

6.15. RAZONAMIENTO BORROSO DINÁMICO	88
6.16. EL CONCEPTO DE DISTRIBUCIÓN DINÁMICA DE POSIBILIDADES	89
6.17. ECUACIONES DE ASIGNACIÓN DE POSIBILIDAD DINÁMICA	90
6.18. PROYECCIÓN Y PARTICULARIZACIÓN	91
6.19. EL CONCEPTO DE UNA VARIABLE LINGÜÍSTICA DINÁMICA	92
 CAPÍTULO 7	
7.1. CONECTIVOS LOGICOS DINAMICOS Y FORMAS SILOGISTICAS DINAMICAS	95
7.2. EL PRINCIPIO DE EXTENSIÓN Y LOS OPERADORES BORROSOS DINÁMICOS	100
7.3. SUBCONJUNTOS BORROSOS DE TIPO II	100
7.4. CONJUNTO DE OPERADORES NO MONÓTONOS	105
7.5. CUANTIFICADORES LINGÜÍSTICOS	107
7.6. RAZONAMIENTO SILOGÍSTICO CON CUANTIFICADORES	109
 CAPÍTULO 8	
8.1. RAZONAMIENTO APROXIMADO DINAMICO	114
8.2. VARIABLES LINGÜÍSTICAS DINÁMICAS	118
8.3. COMPARADORES BORROSOS DINÁMICOS	120
8.4. DECLARACIONES CONDICIONALES BORROSAS	122
8.5. REPRESENTACIÓN DE LA IMPLICACIÓN DINÁMICA MULTIVALUADA	122
8.6. IMPLICACIONES VALUADAS EN UN INTERVALO	124
8.7. RESULTADOS EN CONDICIONAMIENTO	125
8.8. CALIFICACIÓN DE VERDAD	127
8.9. CALIFICACIÓN DE CERTIDUMBRE Y POSIBILIDAD. EL CASO PRECISO.	129
8.10. REGLAS DE CUALIFICACIÓN DE VERDAD	131
8.11. REGLAS DE CUALIFICACIÓN DE INCERTIDUMBRE	131
8.12. EL MODUS PONENS GENERALIZADO	132
8.13. APROXIMACION LOGICA DEL RAZONAMIENTO APROXIMADO DINAMICO	134
8.14. FACTORES DE CERTIDUMBRE Y LÓGICA MULTIVALUADA	135
8.15. EL PUNTO DE VISTA TEÓRICO POSIBILÍSTICO DE LOS FACTORES DE CERTEZA	137
8.16. UN PARADIGMA DE LOS SISTEMAS DE INFORMACIÓN	138
8.17. EL PROBLEMA DE COMPOSICIONALIDAD	139
8.18. SISTEMAS FORMALES	141
8.19. TIPOS DE ECUACIONES RELACIONALES DINÁMICAS	141
 BIBLIOGRAFÍA	 144

CAPÍTULO 1.

- 1.1. INTRODUCCION
- 1.2. DUBITADORES LEIBNIZIANOS
- 1.3. DUBITADORES KANTIANOS
- 1.4. DUBITADORES DE HENGEL
- 1.5. RAZONAMIENTO POR APROXIMACIÓN
- 1.6. MODELOS CAUSALES
- 1.7. LÓGICA FORMAL
- 1.8. TEORÍAS DE PROBABILIDAD
- 1.9. TEORÍA DE LA EVIDENCIA.
- 1.10. TEORÍA DE LOS CONJUNTOS BORROSOS
- 1.11. CONJUNTOS BORROSOS DINÁMICOS

1.1. INTRODUCCION

En 1895 Cantor desarrolla la teoría tradicional de conjuntos y en ella define un conjunto como una colección M de m objetos distintos definidos por nuestra percepción o por nuestro conocimiento. La idea de conjunto se usa para construir los números, el álgebra abstracta, y la mayoría de las matemáticas modernas que conocemos actualmente.

Filósofos y matemáticos como cuestionaron las raíces matemáticas de la teoría de conjuntos. Bertrand Russell apuntó que los conjuntos no podían ser definidos en forma arbitraria sin producir paradojas. El problema se observa como si viniera de situaciones del mundo real donde el axioma de especificación no se puede usar para construir el conjunto en forma consistente. Esto es, hay propiedades significativas que no se pueden usar para únicamente definir el conjunto en forma razonable. Por ejemplo consideremos el universo de objetos, se puede definir un conjunto C_p como una colección de elementos $x \in U$ que satisfacen alguna propiedad $p(x)$. La propiedad se debe de satisfacer para todos o para ninguno, así se puede escribir la definición de C_p en forma simbólica:

$$C_p = \{x \in U \mid p(x)\},$$

que se lee como: *C_p es el conjunto de todas las x existentes en U tales que satisfagan la característica $p(x)$.*

Lo anterior funciona muy bien conforme los objetos sean abstractos y bien definidos. Un simple ejemplo donde lo anterior falla es cuando se intenta definir el conjunto de todos los números que son cercanos a otro número, por ejemplo el 100, o el conjunto de las personas ricas. Estos ejemplos están basados en las propiedades significativas, pero los conjuntos que a menudo definimos fallan al no poder describir la situación en forma apropiada. ¿Cómo se pueden definir los números cercanos a 100? ¿Cómo se pueden definir las personas ricas? La razón de esta falla radica en el uso de la lógica de sólo dos valores permitidos, cierto o falso.

La lógica borrosa, como una lógica multivaluada, ofrece una alternativa de solución ya que permite proposiciones que son parcialmente satisfechas en cierto grado en el intervalo $[0,1]$. Así podemos definir un conjunto borroso B por medio de una función $\mu_B: U \rightarrow [0,1]$ que mapea los elementos del universo U en $[0,1]$. Esta función se conoce como función de membresía de un conjunto borroso.

La teoría de los conjuntos borrosos está basada en conceptos fundamentales. Las primeras raíces de los conjuntos borrosos están radicadas en la segunda mitad del siglo XIX, principalmente en la controversia entre Cantor y Kronecker sobre el significado de los conjuntos infinitos. Kronecker planteaba el problema de la siguiente forma:

No se puede aplicar la formulación de los conceptos del sistema modular con un número infinito de elementos. Si nunca se ha intentado aprobar esta formación como lógica, luego sólo está permitido, con las reservas de que en vista de lo especial, la aplicación de las operaciones aritméticas a nociones que no están suficientemente especificadas como pruebas, y nos permite aplicarlas sólo en los casos especiales; luego esto significa que actualmente en cada caso la introducción de un sistema

modular con un número infinito de elementos se torne superflua.

La reacción de Dedekind fue:

Si un sistema S es un sujeto materia de nuestro razonamiento, si está completamente determinado, siempre se puede decir si un elemento es o no es parte del sistema S. De esta manera la determinación se acompaña de una forma para decidir en esta determinación, que es enteramente indiferente de las siguientes consideraciones... Menciono esto explícitamente porque el Sr. Kronecker ha intentado, recientemente, imponer ciertas restricciones a la formación libre de conceptos matemáticos, cosa que reconozco como sin fundamento.

El compromiso entre los puntos de vista de Kronecker y Dedekind puede ser el siguiente: Un conjunto S está completamente determinado sí y sólo si existe un procedimiento que especifique el grado en el cual un elemento es miembro de S. Usando la primigenia teoría de conjuntos en nuestro metalenguaje, esta aproximación nos lleva a la lógica binaria en el caso de las funciones características, como en el caso de la lógica multivaluada se llega al concepto de función de membresía de Zadeh. En este sentido la regresión de Kronecker de conjuntos infinitos, y la defensa de Dedekind de la teoría de conjuntos de Cantor, son las precursoras de los conjuntos borrosos.

Otra de las fuentes de la teoría de los conjuntos borrosos es la paradoja de Poincare en la interpretación de Menger en el concepto de microgeometría. Contraria a la geometría de Reimann la microgeometría se relaciona con los problemas de la geometría en pequeño, en concreto la descripción del continuo. Dice Menger:

Desde un punto de vista positivista la dificultad de describir el continuo no radica en el empleo de cierto tipo especial de números, sino que radica en la identificación de los elementos individuales, en particular, en su descripción por números de cualquier clase. Ninguno está tan agudamente consciente del problema conectado con la noción del continuo físico, cuyas dificultades han sido enfatizadas muchas veces por Poincare, quien ha caracterizado el continuo físico con la formulación siguiente:

$$A = B, \quad B = C, \quad A \neq C$$

Simbolizando el hecho de que un elemento de un continuo físico puede ser indistinguible de uno y distinguible de otro. El continuo matemático, de acuerdo con Poincare, es una construcción cuyo propósito principal es llegar a una fórmula que él llama contradictoria, y repugnante.

Dice Poincare:

Creo que la última solución del problema de la microgeometría puede caer en una teoría de las probabilidades que admita una etapa intermedia entre lo distinguible y lo indistinguible, y que acepte que se pueden superponer. Estas características esenciales, de dicha teoría, radican en que los cúmulos no son puntos de un conjunto, sino que describen figuras, como elipsoides. Más bien estarían en muchas relaciones probabilísticas de superposición, para las cuales se debe de desarrollar una teoría de medida.

Usando los conjuntos primarios en nuestro metalenguaje se pueden formalizar estas ideas de la siguiente manera: Un espacio microgeométrico es una pareja (X, M) , donde X es un conjunto llamado conjunto de cúmulos, y $M: X \times X \rightarrow [0,1]$ es un mapeo que satisface los siguientes axiomas:

1. $M(x,y) = 1$ si y sólo si $x = y$,
2. $M(x,y) = E(x,y)$,
3. $M(x,y) + M(y,z) - 1 \leq M(X,y)$.

$M(x,y)$ se interpreta como la extensión de la superposición entre x y y . Más aún, induce de una manera simple una relación R_M en X que es: reflexiva, simétrica, pero en general no transitiva, dada por:

$$R_M := \{(x,y) \mid 0 < M(X,y)\}.$$

Los espacios microgeométricos se pueden ver como modelos matemáticos que satisfacen la paradoja de Poincare. Nótese que el modelo de Von Newman de la física cuántica permite la construcción del espacio microgeométrico en el cual la extensión de la superposición se relaciona con la verificabilidad no simultánea. La verificabilidad es otra fuente de la borrosidad y la extensión de la superposición es formalmente una ecuación borrosa en los cúmulos.

Otra de las fuentes de los conjuntos borrosos es la imprecisión inherente a la toma de decisiones humana. Esta fue la principal motivación de Zadeh, cuando postula los conjuntos borrosos.

Zadeh escribe:

En efecto, existe una gran distancia entre lo que podemos llamar teóricos de sistemas animados, y teóricos de sistemas inanimados en el presente, y no se puede asegurar que esta distancia se estreche, mucho menos que se cierre en el futuro cercano. Hay quien piensa que esta distancia refleja la inadecuación fundamental de las convenciones matemáticas –la matemática de puntos precisamente definidos, funciones, conjuntos medidas de probabilidad, etc.– para copiar con el análisis los sistemas biológicos, y que para trabajar efectivamente con estos sistemas, que son más complejos que los sistemas hechos por el hombre, se necesita una clase radicalmente diferente de matemáticas, las matemáticas de lo borroso o de los cúmulos, que no se pueden describir en términos de distribuciones de probabilidad.

¿Qué nos puede ofrecer el estudio de los fundamentos de los conjuntos borrosos? Ciertamente que podemos colocar a los conjuntos borrosos en la tradición de la matemática multivaluada, pero también debemos de hablar de la aplicación práctica de los conjuntos borrosos.

La fundamentación para los conjuntos borrosos nos dará una base rigurosa para la práctica actual de los que aplican la teoría. Mucha de las personas que trabajan con conjuntos borrosos los desean usar con propósitos prácticos: el diseño de controladores borrosos para varios procesos en tiempo real, o para construir sistemas computacionales que lidien con información imprecisa y razonamiento aproximado. Los ingenieros requerimos una teoría borrosa que sea robusta, en el sentido que no sea particularmente sensible a los detalles del modelo y sus conectividades usadas, pero flexible, de forma

tal que el modelo se pueda sintonizar para ofrecer altos niveles de desempeño. Luego la fundamentación para los conjuntos borrosos nos debe de ofrecer una variedad de conectivos que clarifiquen los cúmulos de elecciones posibles. Estudios recientes han incluido un gran número de posibles conectivos. La fundamentación debe permitir esta variedad, pero hay poner orden en ella.

Un sistema lógico proposicional típicamente usa conectivos como “y”, “o”, “no”, y “sí esto luego aquello”. En la literatura aplicada reciente, han sido cuestionados algunos de estos conectivos. Una guía razonable es que queremos extender la lógica clásica de dos valores, de forma tal que cuando nos restrinjamos a los valores extremos, los conectivos se comporten en la forma de los de la lógica de dos valores. La robustez deseada se puede pensar como una condición de continuidad para los operadores en $[0,1]$. Cambios pequeños en el juicio subjetivo del grado en el cual un individuo pertenece a un conjunto borroso particular, puede no producir cambios en el resultado de la inferencia que involucre a tal individuo. Esto permite que los conectivos de Lukaciewicz en $[0,1]$ sean más aplicados que los conectivos intuitivos.

Cuando Zadeh postula los conjuntos borrosos, deseó generalizar la idea tradicional de un conjunto y de una proposición para obtener los grados de pertenencia y valores de verdad respectivamente. Estos esfuerzos se atribuyen a las complicaciones que aparecen durante el modelado físico del mundo real, y que son:

1. Las situaciones reales no son clásicas y resolutivas, por lo que es difícil describirlas precisamente.
2. La descripción completa de un sistema real a menudo requerirá datos más detallados que aquellos en los que el sujeto, el humano, pueda reconocer simultáneamente su proceso y comprensión.

Acerca de lo anterior nos dice Russell:

Todas las habilidades lógicas tradicionales asumen que se deban de emplear los símbolos precisos. Pero esto no es aplicable a esta vida terrestre, sino sólo a una existencia celestial imaginada.

Zadeh nos dice:

Como la complejidad de los sistemas se incrementa, nuestra habilidad de hacer declaraciones precisas y significativas acerca de su comportamiento disminuye hasta umbrales más allá de los cuales llegarán a ser características mutuamente exclusivas la precisión y el significado. Esta declaración es el principio de incompatibilidad. Lo que nos dice es que lo más cercano al mundo real es lo borroso.

Existen dos tipos básicos de incertidumbre que se pueden presentar en las situaciones del mundo real:

1. Incertidumbre estocástica debida a una falta de información de los estados futuros de un sistema ya que no se pueden conocer completamente. Este tipo de incertidumbre tiene carácter estocástico y se maneja por las teorías probabilísticas.
2. Borrosidad, que es la vaguedad implícita en la descripción semántica de un evento,

o fenómeno.

El resultado de un evento estocástico es siempre cierto o falso. En situaciones donde el evento no está bien definido, la salida puede ser una cantidad diferente a cierto o falso, que podemos cuantificar como el grado de creencia. Estos eventos se modelan por conjuntos borrosos ya que su función de membresía puede tomar valores entre uno y cero.

Cuando intentamos resolver un problema o tomar una decisión, a la parte esencial del esfuerzo lo llamamos razonamiento. Informalmente hablando podemos decir que razonar es el ejercicio de inferir información acerca de un aspecto no observable de un hecho, basados en la información de los aspectos observables del mismo. Este ejercicio de razonar lo podemos ver de una manera muy general como que dada la información acerca de los aspectos observables de un problema particular, y un conocimiento acerca del dominio al cual pertenece el problema, se tiene que determinar alguna información acerca de los aspectos inobservables de él.

Dependiendo de la naturaleza y tipo de la información observable, y de los objetivos del razonador, la tarea de razonar se puede establecer de una manera más refinada. Por ejemplo cuando usamos el cálculo de predicados de primer orden para representar conocimiento, el razonamiento se puede especificar como que dados los aspectos observados de un problema en términos de los valores de verdad de algunos predicados, se deben determinar los valores de verdad de otros predicados.

Muchas tareas de toma de decisiones demandan más que la mera inferencia acerca de los aspectos no observables del problema. La tarea demanda que el ejercicio de razonamiento genere hipótesis o una explicación, para justificar los eventos observados.

Se usan dos conceptos primitivos que son: las entradas y las creencias del observador. Las creencias son algunos principios o verdades que se presuponen válidas, y las entradas son las observaciones experimentales que el observador realiza. Una red de hechos es la interconexión de entradas y creencias que se construyen por medio de un conjunto de relaciones y operadores. Es por lo tanto una red de verdades contingentes. En este sentido podemos clasificar los modelos de dubitación como:

1.2. DUBITADORES LEIBNIZIANOS

Estos modelos de dubitadores aspiran a construir redes óptimas de sucesos para un problema. La optimalidad llega a ser importante porque para cualquier conjunto dado de entradas, es posible construir un número de diferentes redes de sucesos. Estos dubitadores asumen que existe una red óptima, y el proceso de construir la red convergirá con él, ya que las habilidades son inherentes al proceso de construcción. Es un método de razonamiento general. Cada nuevo resultado científico es una entrada que se puede ligar en la red con anteriores resultados, y en especial cuando el campo está constituido en una teoría. La teoría da las relaciones y los operadores amarran los resultados en la forma de una red de sucesos que corresponda a las creencias del dubitador. Un resultado que salga de la gran red se ignora, mientras que los resultados que encajan en la red son tomados en cuenta. Estos dubitadores asumen la existencia de

un modelo apriori del problema y tienden a configurar las entradas del problema acordes con el modelo.

1.3. DUBITADORES KANTIANOS

Estos dubitadores no presuponen la existencia a priori de modelos del problema. El único modelo apriori que se asume es lo básico que permita al dubitador observar las entradas. Un dubitador kantiano inicialmente contiene un conjunto de modelos. Cada modelo es un conjunto independiente de creencias y puede contener sus primitivos, axiomas y reglas de inferencia. Un dubitador selecciona un modelo del conjunto de modelos y realiza una red de Leibnitz usando las entradas y las creencias del modelo. El dubitador determina la extensión a la cual esta red de hechos, de acuerdo con algún criterio, es satisfactoria. El modelo que genera la red de hechos más satisfactoria es la solución a la duda.

1.4. DUBITADORES DE HENGEL

Estos dubitadores buscan desarrollar la habilidad de ver las mismas entradas desde diferentes puntos de vista. El dubitador posee un número de modelos, y cada uno de ellos tiene un conjunto independiente de creencias y puede contener primitivos, axiomas y reglas de inferencia. Denotaremos estos modelos por M . Sea I un conjunto de proposiciones que pueden ser ciertas en el problema bajo consideración. Sea E el conjunto de entradas $e_1, \dots, e_k$. Sea X un operador que une un elemento de E con un elemento de M , tal que para cada e en E y cada m_i en M le corresponde uno y sólo un elemento del conjunto I . In mapea elementos de E para un conjunto M dado, encima del conjunto I en una correspondencia muchos a uno. El conjunto I se llama el conjunto de información de un conjunto de modelos M dado, y el operador In se llama el operador intérprete. Luego para cada elemento de M debe corresponder un subconjunto de información que se representa por $I(M_j)$.

Consideremos un conjunto de tesis T , que son sentencias que nos dicen algo respecto del mundo, tales que los elementos de T no se implican por ningún elemento de E , M , o I . G es una función de dos lugares que transforma a T_i y a M_j dentro de elementos de un sistema de números reales. G representa el grado de confianza en T_i dada por la información en $I(M_j)$. Esta representa la creencia de la tesis, dado que el mundo es aproximadamente descrito por el modelo M_j .

El dubitador selecciona una tesis de A del conjunto T y se compromete a construir el caso para soportar A , en efecto, una defensa de la tesis A . Esto se logra al seleccionar a M_{esimo} como:

$$M_{esimo} = \max_j [G(A, I(M_j))]$$

Esto es, el dubitador afirma que ésta es una manera de mirar la realidad, dada M_{esima} tal que las entradas se puedan interpretar para soportar la tesis A .

La siguiente cosa que hace el dubitador de Hegel es encontrar la tesis B que es la antítesis de A y también encontrar el modelo que soporta a B . B no tiene que porque ser la negación lógica de A .

El siguiente acto del dubitador es mirar los dos modelos del mundo, las M_i 's que soportan la tesis y la antítesis y examina las fuentes del conflicto entre ellas. Se espera resolver o entender y caer en la verdad acerca del problema. El gran modelo de este problema en el contexto en el cual el conflicto se puede entender se llama la síntesis del problema.

La elección del tipo de dubitador depende de nuestro conocimiento y de nuestra vista de donde debe radicar la verdad. Si pensamos que la verdad está en el modelo, y un buen modelo está a la mano del dubitador, los usados en las ciencias físicas son los más apropiados. A este se le llama la aproximación del modelo directo. Por otro lado, si nuestro punto de vista tiene parcialmente la verdad en un modelo y parcialmente en los datos, podemos emplear el modelo kantiano, que es el caso en muchos modelos sociales. Los modelos hegelianos o modelos dialécticos de cuestionar, son más apropiados para aquellos dominios donde en nuestro punto de vista la verdad puede emerger de la oposición entre tesis y antítesis.

El punto principal que se puede concluir es que existen muchas aproximaciones para dudar y dependiendo del dominio del problema, una aproximación puede ser más útil que otra. Cuando llevamos a cabo un proceso de razonamiento, en una situación especial, este se debe realizar en el contexto del modelo de duda útil para el problema entre manos. Así podemos plantearnos lo siguiente: Dados los aspectos observados de un problema, en términos de los eventos observados, un conocimiento del dominio que nos da los hechos básicos para construir los modelos, para varias situaciones posibles en el dominio y determinar algunos modelos, que de acuerdo con los eventos observados, satisfagan algunos criterios específicos de interés, se debe de poder dar una inferencia acerca de algunos aspectos no observados del problema en el contexto de los modelos anteriores.

Un problema, que se enfrenta en muchas situaciones de razonamiento, es que toda la información que debe ser relevante no está disponible y la que está disponible es confusa y no necesariamente relevante.

1.5. RAZONAMIENTO POR APROXIMACIÓN

En muchas situaciones de razonamiento el razonador se enfrenta al problema de tener la información relevante incompleta. Cuando se está en una situación donde se tiene un conocimiento observado incompleto, se pueden generar un número de hipótesis o explicaciones, tales que cada una de ellas puede explicar algunos de los hechos observados. El razonador se enfrenta al problema de manejar la multiplicidad de hipótesis que explican la situación.

Otro de los problemas con que se encuentran los razonadores es la incertidumbre del conocimiento. En muchas situaciones, a menudo no estamos en la posición de insistir que las declaraciones acerca de una cierta entrada o una creencia sean absolutamente ciertas, pero estamos en la posibilidad de admitir que tienen un cierto grado de incertidumbre. Se deben de hacer un número de intentos para modelar esta incertidumbre.

Otro de los problemas que se han encontrado es la imprecisión del conocimiento. El problema se presenta cuando las entradas y las creencias son imprecisas por naturaleza. En muchas situaciones de examen de nuestra vida diaria los modelos se construyen con predicados no muy precisos. En lugar de descartar algunas entradas, tratamos de encontrar versiones menos precisas del modelo en el cual podemos incluir las entradas. El único método que está sistemáticamente dirigido al razonamiento por aproximación es la teoría de los conjuntos borrosos.

Si un razonador trabaja de acuerdo con el modelo de Leibniz, está interesado en determinar el valor de verdad de alguna proposición tratando de formar la red de hechos que la contenga, luego un razonador aproximado está interesado en determinar la incertidumbre asociada con la proposición de interés, dada la incertidumbre asociada con las entradas y las creencias, o la versión imprecisa de la proposición y el modelo, de forma tal que la proposición pueda ser probada como cierta dadas las imprecisiones asociadas con las entradas y las creencias.

Un razonador aproximado que trabaje con un modelo kantiano, está interesado en determinar una lista ordenada de pocos modelos que puedan explicar las entradas. La imprecisión asociada con algunas proposiciones incluidas en el modelo puede ser un criterio para preferir uno sobre los demás modelos.

Un razonador aproximado que trabaja de acuerdo con un modelo hegeliano estará interesado en determinar parejas de modelos que soporten la tesis y la antítesis. La precisión de las proposiciones en los modelos puede ser parte del criterio para preferir un modelo sobre otro.

Veremos ahora la naturaleza de los modelos de situaciones que se pueden emplear por varios tipos de razonadores.

1.6. MODELOS CAUSALES

El conocimiento de relaciones de causalidad nos da las bases para relacionar varios efectos y luego nos da la estructura causal para la situación. La información causal parece ser la base más convincente para relacionar varios eventos. Tversky y Kahneman lo plantean de la siguiente forma:

Es lo más común que las personas se esfuercen para lograr una interpretación coherente de los eventos que los rodean, y la organización de los eventos por esquemas de relación causa-efecto no sirve para lograr esta meta.

Nos muestran que en el contexto de la conducta humana los datos que nos dan una interpretación causal afectan los juicios, ya que los datos que no caen en este esquema son dominados por datos causales relevantes. Los sistemas de razonamiento buscan tratar con las estructuras de situaciones que usan información causal para interconectar varios eventos. Filosóficamente existen muchos problemas con la idea de causalidad. Muchos argumentan que somos nosotros los que siempre buscamos establecer relaciones de causalidad entre varios eventos. Nosotros consideraremos la causalidad como una experiencia humana, y así puede formar parte de las creencias a través de las cuales podemos construir nuestros modelos de la realidad. Usamos estas creencias para

percibir causalidad sin cuestionar la validez de estas relaciones causales percibidas. ¿No cuestionamos causalmente a nuestro perro?

Examinaremos varios sistemas de razonamiento y sus capacidades en varios tipos de situaciones de razonamiento. Los sistemas los clasificaremos en dos grandes rubros, los basados en lógica formal, y los basados en cálculo numérico.

1.7. LÓGICA FORMAL

Muchos de los marcos de razonamiento están basados en lógica formal y por lo mismo se desprecian los problemas del razonamiento aproximado. La teoría de la probabilidad ha avanzado como un marco mejor situado para manejar incertidumbre, pero se han presentado algunos sistemas no numéricos para manejar incertidumbre dentro del marco de lógica formal, como el presentado por Halpern y Rabin en 1987 y Poole en 1988.

Un teorema formal que prueba sistemas es un dubitador de Leibniz. Este mecanismo de prueba de teoremas es el mismo con el que la construcción de una red de posibilidades deriva la cláusula nula, que es una actividad equivalente a la construcción de las redes de sucesos. Las entradas y los axiomas constituyen el único modelo de la situación y el intento de probar un teorema es determinar si el teorema puede deducirse del conocimiento contenido en el modelo.

Los modelos de Kant y de Hegel requieren más de un modelo de la situación. Es posible que los axiomas correspondientes a dos diferentes modelos de una situación sean inconsistentes, por lo tanto el marco de teoremas no corresponde a los modelos kantianos y hegelianos.

La lógica por omisión nos presenta un marco que puede ser empleado para implantar dubitadores. En el marco de esta lógica, los razonadores poseen dos tipos de conocimientos: hechos y omisiones. Los hechos son estas piezas de información que son conocidos como ciertos en el mundo. Las omisiones son aquellas piezas de información que pueden o no ser ciertas en una situación particular. Por ejemplo la frase, los pájaros vuelan, no es siempre cierta, y por lo mismo no es un hecho. Si sólo algunos pájaros vuelan y otros no, esta información se puede asumir como cierta en ciertas descripciones del mundo y en otras no. Luego es posible una pieza de conocimiento por omisión. Un ejemplo de lo anterior es el dado por Poole en 1988:

Omisión 1. Mamífero $\rightarrow \neg$ vuela (X)

Omisión 2. Vampiro (X) $\rightarrow$ Vuela (X)

Omisión 3. Inservible(X) $\rightarrow \neg$ vuela (X)

Hecho 1 Vampiro(X) $\rightarrow$ mamífero(X)

Hecho 2 Vampiro (Drácula)

Hecho 3 inservible (Drácula)

Podemos ahora decir que la entrada disponible para nosotros es: Drácula no vuela.

Podemos construir una teoría que incluya los tres hechos y la omisión 1 y de acuerdo con esta teoría podemos explicar la entrada. Si lo que tenemos disponible como entrada fuera Drácula vuela, podemos construir una teoría con los tres hechos y la omisión 2

para explicarla. El punto es que podemos usar la inferencia lógica dentro de una teoría, pero cuando un número de teorías nos permite la selección de una teoría apropiada, esta teoría es la de máxima importancia. En el ejercicio del razonamiento importa tanto la tarea de inferir como la selección de la teoría conveniente. Poole establece que en lugar de esperar razonamiento para deducciones precisas de nuestro conocimiento, examinemos las consecuencias de ver el razonamiento como un simple caso de la formación de una teoría.

Las componentes básicas usadas para construir una teoría son los hechos y las omisiones. En cualquier teoría, todos los hechos y algunas de las omisiones se incluyen de tal manera que la teoría permanece consistente. En muchas situaciones puede ser posible construir varias teorías para explicar un particular conjunto de entradas. Si la teoría es para servir como modelo de una situación y el razonador está planeando tomar alguna acción basado en el modelo asumido, se puede requerir que sólo sea seleccionado uno de los modelos. Una alternativa para manejar tal situación es determinar todas las teorías que explican la evidencia y luego determina la utilidad de cada posible acción en el contexto de cada teoría. Para seleccionar la acción óptima puede usarse el marco de trabajo de la teoría de decisiones. En esta aproximación necesitamos conocer la oportunidad relativa que cada teoría tiene de convertirse en el modelo de la situación. La otra alternativa es ordenar varias posibles teorías de acuerdo con algunos criterios de interés, y elegir la teoría más interesante. Luego es posible una aproximación durante el ciclo de hipótesis y prueba, donde podemos presentar las hipótesis más interesantes y después probarlas.

La incapacidad del razonador de obtener todas las entradas necesarias para desenredar la teoría actual, puede darse de cualquiera de las siguientes dos formas: la primera opción es cerrarse a un razonamiento kantiano y la segunda a un razonamiento hegeliano. En la primera opción se usan las mediciones numéricas para indicar la incertidumbre, y en la segunda las incertidumbres se representan aceptando más de un modelo como posible.

En muchas publicaciones se ha presentado el problema de representar la incertidumbre asociadas con las proposiciones sin tener que emplear medidas numéricas. Muchos de estos trabajos se han realizado en el ámbito de las lógicas modales, como el de Segerberg en 1971 y Gärdenfors en 1975, han presentado lógicas modales usando fórmulas de la forma $p \geq q$ que se lee como p es más creíble que q . En 1987 Halpern y Rabin presentaron la lógica LL para razonar alrededor de la probabilidad. En esta lógica el operador modal L captura la noción de ser probable. Esto es, la fórmula L_p representa la noción “es probable que p sea cierta” Un modelo M para esta lógica se puede ver como un árbol de estados e . Cada estado e consiste de un conjunto de hipótesis que son tomadas como verdad por ahora. Un estado e_2 es un hijo del estado e_1 si asumimos que e_1 sea cierto, un experto con conocimiento del dominio puede llamar a e_2 un estado creíble. Cada estado representa un conjunto de hipótesis consistentes. También un estado no representa un modelo causal de la situación que se describe, aunque a nuestra vista pueda ser una característica deseable para cualquier razonador.

El problema de manejar el nivel de imprecisión no se ha direccionado dentro del marco de trabajo de las lógicas formales. Para formar una red de hechos con declaraciones menos precisas no se puede hacer en forma automática, si el intento con declaraciones más precisas falla.

1.8. TEORÍAS DE PROBABILIDAD

La relación de probabilidad más empleada desde el punto de vista del razonamiento es la probabilidad condicionada. Consideremos dos variables A y b con conjuntos asociados de posibles valores discretos $(a_1, \dots, a_k)$ y $(b_1, b_2, \dots, b_n)$. Pensemos que el razonador conoce los valores de la probabilidad condicional $P(a_i | b_j)$. Cuando se observa un evento $B = b_i$, el razonador puede calcular los valores de probabilidad para los eventos no observados $A = a_n$.

En una situación de razonamiento donde son conocidas, solamente en términos de una probabilidad condicionada, las relaciones entre varios aspectos de una situación, el conocimiento del evento observado se puede, en general, usar para inferir las probabilidades de otro evento no observado. El trabajo realizado en el área de las redes bayesianas por Pearl en 1986, ha resultado ser un marco de trabajo para estos razonamientos. Las redes bayesianas son gráficas acíclicas en las cuales los nodos corresponden a las variables proposicionales, los cortes representan las relaciones de probabilidad condicional entre las proposiciones ligadas, y también significan relaciones causales. En la red las relaciones causales se cuantifican por probabilidades condicionadas de cada variable dado el estado de sus nodos padres.

En cualquier modelo de una situación es suficiente incluir solamente una sola causa para cada efecto, por lo cual es necesario conocer la causa. Cuando se pueden presentar dos causas individuales, se pueden ver los efectos causados por la interacción de ambas causas individuales, y luego esta interacción es la causa exclusiva responsable del efecto. De esta manera podemos decir que las redes bayesianas que incluyen todas las causas conocidas para cada efecto es una combinación de todos los modelos individuales que pueden ser posibles para la situación. Una inferencia que usa una red se puede decir que está basada en lo siguiente: se deben de conocer los modelos conocidos para la situación, y considerando que cada uno de ellos sea posible, la inferencia refleja los efectos de cada modelo posible.

Las redes bayesianas representan el conocimiento de una relación causal por medio de probabilidades condicionales asociadas con los cortes de la red. Hunter demostró en 1980 que las probabilidades condicionales fallan al capturar la información importante acerca de la dirección de la influencia causal y su información puede ser crucial en algunas instancias de razonamiento. Este tipo de razonadores está más cerca de los dubitadores leibnitzianos. Estos dubitadores, dado un modelo de la situación, determinan sus consecuencias lógicas. Las redes bayesianas son luego un modelo de la situación incorporando todas las relaciones probabilísticas entre los aspectos de interés. Las redes bayesianas no se pueden usar para construir dubitadores del tipo kantiano o hegeliano, en los cuales el razonador necesita identificar un modelo factible de entre todos los modelos.

El empleo de la probabilidad para razonar también ha sido explorado por Spiegelhalter en 1988. Sus técnicas son desarrolladas para un contexto en el cual una estructura gráfica representa la relación de dependencia entre proposiciones. Sus métodos intentan actualizar información probabilística a todos los otros nodos cuando se obtiene evidencia de uno de los nodos. Se realizan actualizaciones sucesivas cuando es disponible evidencia en más de un nodo.

1.9. TEORÍA DE LA EVIDENCIA.

Se han desarrollado cálculos numéricos diferentes a la teoría de la probabilidad para usarlos en muchos sistemas de razonamiento. Un cálculo popular es la teoría de la evidencia desarrollado por Shafer's en 1976.

Si S es el conjunto universal de eventos, luego en la teoría de la probabilidad de incertidumbre, la información consiste de una densidad de probabilidad p en los elementos de S , tal que: $p: S \rightarrow [0,1]$ y:

$$\sum_{s \in S} p(s) = 1$$

En la teoría de la evidencia definimos una función de agrupamiento que es inducida en nosotros por la evidencia disponible, y asigna partes de una cantidad finita de creencia al subconjunto S . Cada asignación de agrupamiento a un subconjunto s de S representa la parte de nuestra creencia que soporta a s sin ser posible localizar esta creencia entre subconjuntos estrictos de s .

Los grados de creencia en un subconjunto A de S $Cre(A)$ se definen como la suma de todas las agrupaciones que soportan A o cualquiera de sus subconjuntos estrictos. Esto es:

$$Cre(A) = \sum_{X \in A, X \neq \Phi} m(X)$$

El grado de plausibilidad de un subconjunto A de S , $Pl(A)$ se define como la suma de todas las agrupaciones que puede posiblemente soportar A . La posibilidad de que un subconjunto s que soporte a A , significa que la agrupación asignada a s puede posiblemente gravitar en ese subconjunto estricto de s , que también es un subconjunto de A . Esto es:

$$Pl(A) = \sum_{X \cap A, X \neq \Phi} m(X)$$

Cuando hay disponible más de una función de agrupamiento para un conjunto S , se combinan usando la regla de Dempster, la cual lleva a cabo una suma ortogonal de las dos funciones de agrupamiento. El proceso de actualizar funciones de creencia se ve desde dos perspectivas: Una ve la creencia como una función de plausibilidad, cuyo supremo e ínfimo envuelve una familia de distribución de probabilidades, otra es el intento de relacionar la teoría de evidencia con la teoría de probabilidad, como una generalización y no en la dirección de las funciones de creencias que se deben de ver a la luz de otro marco de trabajo.

1.10. TEORÍA DE LOS CONJUNTOS BORROSOS

La teoría de los conjuntos borrosos desarrollada por Zadeh, captura y representa nociones imprecisas en sí mismas como alto, joven, etc. Si consideramos los predicados como bajo, en lugar de evaluar la verdad o falsedad para parámetros reales valuados, los

predicados actúan como una constricción elástica en los valores aceptados para los parámetros. Para cada x perteneciente al conjunto de los números reales, uno puede definir $\mu_{\text{bajo}}(x) \in [0,1]$ como la extensión de membresía de los números x en los valores aceptables para bajo. El grado de membresía $\mu_{\text{bajo}}(x)$ es interpretado por Zadeh como la posibilidad $\mu_{\text{bajo}}(x)$ de que la proposición sea cierta con el valor x . Es simple calcular las medidas de necesidad y posibilidad definidas para cada proposición.

Esta teoría es muy útil cuando uno quiere hacer transiciones entre proposiciones muy precisas pero inciertas y proposiciones borrosas y menos precisas, pero estrictas.

1.11. CONJUNTOS BORROSOS DINÁMICOS

A nuestro parecer las cosas no son tan simples como lo aquí presentado. En todos los dubitadores, incluyendo los basados en los conjuntos borrosos de Zadeh, las creencias se enmarcan como fundamentalmente dadas y estáticas por el observador. Parece que aceptamos que el creer es una característica definitoria e inamovible del observador. Dado el conjunto de creencias, la pertenencia o no, sean totales o parciales, de las entradas nos permite una sola clasificación posible. ¿Podemos aceptar que el esquema de creencia acerca de la estatura humana sea el mismo entre los pigmeos y los jugadores de un equipo profesional de balón cesto? Es evidente que no, y lo mismo pasa con la creencia acerca de la riqueza entre los pueblos de tercer mundo y los del primer mundo.

Podemos de lo anterior decir que no existe un conjunto borroso que las defina sino que este conjunto depende del contexto de creencia del observador, por lo que la función de pertenencia debería de ser enmarcada como dependiente del contexto, sea éste el de los pigmeos, el de los jugadores de balón cesto, el de los habitantes del primer mundo o del tercer mundo, sólo así podemos establecer clases y conjuntos.

También tenemos que aceptar dada su claridad que los contextos de creencia son cambiante y no estáticos, por lo que la creencia de la altura humana en un contexto situacional dado no es estática. Por ejemplo, no tiene el mismo significado la palabra hombre o mujer alta en el México del siglo XXI, que en el México del siglo XIX. La alimentación, el acceso a la información de la salud, la medicina etc. Han cambiado la creencia social acerca de ser alto.

Lo que en nuestro días es bueno y lo vemos como absolutamente normal, en la época de nuestros abuelos no lo era, nunca hubiese sido aceptado por ellos. La creencia acerca de la bondad o malicia de nuestros actos, cambian con el contexto social y éste cambia a su vez con el tiempo, por lo que si las entradas son las mismas, pero el conjunto de creencias cambia, la pertenencia o no a una clase también se modifica, cambia el estado de pertenencia de la entrada a la clase.

Este tipo de realidad se presenta en todos los sistemas bajo estudio. Si las características de un medio físico son cambiantes y son cambiantes el conjunto de creencias del observador, las formas de clasificación necesariamente lo son, por lo que debemos aceptar que algo que en un tiempo pertenecía a una clase en forma absoluta, la definía, la modificación de la clase por el cambio del conjunto de creencias lo puede excluir. Es decir puede ir dejando el grado de valor absoluto de pertenencia a no pertenecer al

conjunto, a través de transcurso del tiempo. Ejemplo de lo anterior pueden ser el juicio moral, o la situación económica.

De lo anterior se desprende que no basta la definición de conjunto borroso de Zadeh para poder inferir acerca del mundo real, ya que tanto las entradas como las creencias son cambiantes en el tiempo, por lo que proponemos una extensión de los conjuntos borrosos de Zadeh a una teoría de conjuntos borrosos dependientes del tiempo, a los cuales llamaremos conjuntos borrosos dinámicos, y su lógica asociada a la que llamaremos lógica borrosa dinámica, que es el motivo de este trabajo.

Podemos definir un conjunto borroso dinámico de la siguiente forma:

C_{bd} es el conjunto de todas las x existentes en U tales que satisfagan la característica $p(x(t))$ en un t dado.

En situaciones donde el evento no está bien definido y es cambiante en el tiempo, la salida puede ser una cantidad diferente a cierto o falso, que podemos cuantificar como el grado de creencia en un tiempo t . Estos eventos se modelan por conjuntos borrosos dinámicos ya que su función de membresía puede tomar valores entre uno y cero dependiendo del tiempo.

CAPÍTULO 2

2.1. LOGICA BORROSA

2.2. BORROSIDAD

2.3. LÓGICA BORROSA DE PRIMER ORDEN

2.4. FORMULISMO

2.5. TEORÍAS DE LÓGICA BORROSA DE PRIMER ORDEN

2.1. LOGICA BORROSA

Los logistas tradicionales tenían la convicción de que las sentencias de los lenguajes naturales podían ser sólo ciertas o falsas, y cuando mucho, incorporaron un tercer valor que se ha interpretado como indeterminado. Los logistas modernos han pensado muy seriamente acerca de la semántica formal, y otros estudiosos de la forma del lenguaje, como los filósofos y psicólogos, han llegado a determinar el hecho de que los conceptos empleados en el lenguaje natural tienen sus fronteras vagas y su significado es borroso. Dado lo anterior, las sentencias de un lenguaje natural a menudo no son siempre totalmente ciertas, o totalmente falsas, y además tienen un cierto sentido, ya que en algunos conceptos son ciertas, y en otras falsas.

Para los logistas clásicos el valor de verdad o falsedad de una sentencia en términos de la teoría clásica de conjuntos se define con base a la pertenencia o no, del sujeto de la sentencia, a un conjunto perfectamente definido (predicado). Si decimos Juan es alto, esta sentencia es verdadera si Juan pertenece o no al conjunto de hombres altos. Verdadera si pertenece, falsa si no pertenece. Pero en los lenguajes naturales tenemos el problema de que el concepto de altura, es un concepto relativo y no definido, ya que no es lo mismo un hombre alto como concepto entre los jugadores de un equipo de balón cesto y los pigmeos.

Consideremos además algunas otras situaciones como el hecho de preguntarnos acerca del valor de verdad de algunas sentencias del tipo:

X es un Y

Esta podría ser verdadera sólo en cierto grado en lugar de ser claramente cierta o falsa, y en los siguientes ejemplos se ve:

El canario es un pájaro	cierto
El águila es un pájaro	menos cierto que a
El pingüino es un pájaro	menos cierto que b
El murciélago es un pájaro	muy lejos de la verdad
El perro es un pájaro	falso

Estas mismas sentencias las podríamos escribir, dando un juicio relativo a la categoría de membresía asignando el miembro (elemento) en cuestión a la categoría en la cual tiene el mayor grado de membresía.

El canario es más pájaro que otra cosa.	verdadero
El águila es más pájaro que otra cosa	verdadero
El pingüino es más pájaro que otra cosa	Más verdadero que falso
El murciélago es un pájaro	falso
El perro un pájaro	falso

Como se ve del ejemplo anterior para la lógica tradicional no es siempre fácil asignar los valores de verdad a las sentencias del lenguaje natural.

Si partimos de la hipótesis fundamental de que todo el conocimiento humano acerca del mundo real está basado en teorías, tomadas aquí en el sentido más general, ya que incluye teorías formales y pensamientos intuitivos, y el hecho de que estas teorías son

descripciones, y que la descripción para darse requiere de un lenguaje y el lenguaje de una lógica, debemos de desarrollar una nueva lógica que contemple las características de borrosidad del lenguaje natural y nos permita con él, describir y teorizar.

Se define un lenguaje como un conjunto de sentencias compuestas de símbolos acordes con ciertas reglas.

Una interpretación, es una función que asigna un significado a las palabras de un lenguaje relativo a cierto sistema, y especifica las verdaderas condiciones para las sentencias.

Si empleamos un lenguaje para hacer descripciones, necesitamos alguna lógica, reglas para hacer inferencias, para derivar sentencias conductivas de un conjunto de premisas. Si las reglas tienen las propiedades de preservar la verdad, dando sólo consecuencias ciertas de premisas ciertas, tendremos una clase de sistema deductivo, en este caso borroso, que se puede identificar con una teoría.

De acuerdo con esto, desarrollar una lógica, acorde al lenguaje natural que empleamos para definir los fenómenos actuales, nos permitirá hacer de él una teoría. Tenemos conocimiento intuitivo o empírico de algunas sentencias que son ciertas, y emplearemos de acuerdo a esta lógica algunas reglas de inferencia para obtener nuevas sentencias ciertas.

La interpretación, a menudo, está en estos casos íntimamente relacionada con el lenguaje, cada palabra tiene su propio significado, en este caso, más o menos definido (borroso).

En la aproximación algebraica a la lógica, los predicados se interpretan como funciones desde el dominio de una interpretación a una apropiada red completa. Nótese que un conjunto borroso sobre un dominio U es una función de U dentro de un intervalo $[0,1]$ de la recta real.

Históricamente, hasta el final del siglo XIX, las teorías matemáticas se construyeron en una forma intuitiva o axiomática. En otras palabras éstas fueron siempre construidas sobre ideas intuitivas relacionadas con nociones primitivas o en las bases de propiedades de nociones primitivas establecidas en un conjunto de axiomas. El descubrimiento de paradojas en la teoría intuitiva de conjuntos de Cantor, como la de Russell de 1903, fuerzan la necesidad de construir teorías matemáticas axiomáticas. Las paradojas de una teoría intuitiva de conjuntos se podían eliminar en un sistema axiomático de una teoría de conjuntos. La primera teoría axiomática de conjuntos fue propuesta por Zermelo en 1908, sin embargo el método axiomático no preveía la aparición de otra clase de paradojas, las paradojas semánticas.

Esto inspiró la introducción de una gran precisión en la construcción de las teorías matemáticas, llevándolas al concepto de teoría formalizada, en las cuales no solamente se formulaban por un conjunto de axiomas los conceptos de nociones primitivas, sino que se describían precisamente tanto el lenguaje de la teoría, como las herramientas lógicas usadas en cualquier proceso de deducción. Las formas matemáticas de razonamiento estaban basadas en el principio de que cualquier proposición matemática

es siempre cierta o falsa, esto es, es una lógica bivaluada, y usualmente formalizada dentro del marco del cálculo de predicados clásico de primer orden.

Durante mucho tiempo se ha conocido la fuerte relación entre la lógica clásica y la teoría de las álgebras Booleanas. Las investigaciones en lógica de Boole llevaron al concepto de las hoy llamadas álgebras Booleanas. Los resultados de Stone en 1936 en la representación de las álgebras Booleanas por campos de conjuntos, extienden esta relación al campo de los conjuntos. La interpretación de fórmulas del cálculo proposicional clásico como una función sobre cualquier álgebra Booleana, es una generalización del bien conocido método de las tablas de verdad. Una extensión de este método al cálculo clásico de predicados, a través del trabajo de Rasiowa y Sikorski en 1950, permitió una interpretación de las fórmulas como funciones del universo de una interpretación, hacia álgebras Booleanas completas, y hacia campos de todos los subconjuntos de un espacio.

En 1908 Brouwer formuló los principios, desde punto de vista filosófico, de la fundamentación de las matemáticas que se conoce como intuicionismo. Lo anterior fue anticipado por Kant, Kronecker y Poincaré, pero fue desarrollado por Brouwer y su escuela como una forma radical de constructivismo.

El propósito del intuicionismo es prevenir paradojas en las matemáticas. Las diferencias básicas entre el punto de vista de muchos matemáticos y el de los intuicionistas en el significado de la lógica fundamental y los conceptos de la teoría de conjuntos caen en un diferente entendimiento de la noción de conjuntos infinitos, en la palabra “existe”. Por ejemplo, “debe de existir un entero positivo x que satisface la condición $A(x)$ ” es considerado por los intuicionistas como verdad si existe un método de construir un entero positivo x que satisfaga la condición $A(x)$, esto es, si tal entero es presente. En general, los matemáticos entienden lo anterior como verdadero si existe una prueba desde los axiomas de la aritmética, con la base de la lógica clásica, en particular probando que no para todo x , no $A(x)$. Luego si aplicamos la ley de la lógica clásica que “si no para todo x , no $A(x)$, debe de existir un x tal que $A(x)$, y de la regla del modus ponens, se infiere que debe de existir un x que satisfaga $A(x)$. Esta prueba por reducción a una contradicción no nos da ningún método para construir algún x que satisfaga $A(x)$. Lo anterior es rechazado por los intuicionistas. La eliminación de las pruebas no constructivas del razonamiento matemático, causa la retracción de todas las leyes de la lógica clásica que pueden caer de tales pruebas. En particular, la ley del medio excluido que establece la bivalencia de la lógica clásica, es rechazada por el intuicionismo.

Es sorprendente que la investigación algebraica de la lógica intuicionista, que fue ideada para formalizar el razonamiento constructivo, nos lleve a interpretaciones, por ejemplo, en el álgebra de todos los subconjuntos abiertos de la línea real en el intervalo $[0,1]$, que sugiere una interpretación en el álgebra de los valores borrosos de verdad, si subconjuntos abiertos de $[0,1]$ son tratados como valores borrosos de verdad.

El criticismo filosófico relacionado con la implicación material clásica, resultó en la formulación de la lógica modal de Lewis, en la cual, los conectivos modales proposicionales “es necesario que” y “es posible que” ocurren al lado de los conectivos proposicionales de la lógica clásica. Lewis introduce las lógicas modales como un sistema deductivo formalizado del cálculo proposicional. Otros autores construyeron otras lógicas modales, jugando un importante papel en las ciencias computacionales

teóricas, y en la inteligencia artificial, por ejemplo, en la lógica de programación, en las lógicas no monotónicas y en la lógica del conocimiento de Chellas y Turner.

Las investigaciones de McKinsey y Tarski establecieron la conexión entre el cálculo modal proposicional y las álgebras booleanas topológicas, así como con el campo topológico de los conjuntos. Las fórmulas de este cálculo se pueden interpretar como funciones definidas en cualquiera de estas álgebras. Es de especial interés el resultado concerniente a la interpretación, en el álgebra, de todos los subconjuntos de la recta real o del intervalo $[0,1]$, tratados como un campo topológico de conjuntos. Rasiowa y Sikorski iniciaron un tratamiento algebraico del cálculo modal de predicados. Esta aproximación permite una interpretación de las fórmulas del cálculo modal de predicados como una función del universo de interpretación dentro de álgebras booleanas topológicas, y dentro de los campos topológicos de todos los subconjuntos de los espacios topológicos.

El razonamiento del sentido común toma en consideración las incertidumbres, particularmente las relacionadas con declaraciones acerca de eventos en el futuro. La idea de que estas declaraciones no siempre son ciertas o falsas, lleva a Lukasiewicz a su cálculo proposicional de tres valores, y luego a la de m valores, y después a la generalización de un número incontable de valores. En forma independiente Post introdujo su cálculo proposicional de m valores. Todos estos cálculos no fueron construidos como sistemas deductivos axiomáticos formalizados, sino por el método de tablas de verdad.

En la lógica polivaluada de Lukasiewicz los conectivos proposicionales $\neg, \cup, \cap$ y $\Rightarrow$ están basados en reglas aritméticas. Se puede desarrollar un sistema de tres valores de los tres valores lógicos: 0, $1/2$, y 1, donde 0 y 1 corresponden a lo falso y cierto respectivamente, usando las siguientes ecuaciones para cada $x, y, z, \in \{0, 1/2, 1\}$:

$$\neg x = 1 - x, x \cup y = \max(x,y), x \cap y = \min(x,y)$$

$$x \Rightarrow y = \begin{cases} 1 & \text{si } x \leq y \\ 1 - x + y & \text{si } x > y \end{cases}$$

Se usan las mismas ecuaciones para desarrollar sistemas de m valores en los cuales los valores lógicos son:

$$1, \frac{m-2}{m-1}, \dots, \frac{1}{m-1}, 0 \quad (1)$$

En el cálculo proposicional de infinito número de valores de Lukasiewicz son números en el intervalo $[0,1]$, y los conectivos proposicionales están definidos por la misma ecuación (1). Esta lógica es considerada por los especialistas en lógica borrosa como la base del razonamiento aproximado.

El primer sistema axiomático para álgebras correspondientes a la lógica de m valores de Post fue dado por Rosembloom. El artículo de Epstein inició un desarrollo intensivo de la teoría de estas álgebras. El primer sistema ecuacional axiomático fue propuesto por Traczyk, y otro importante sistema ecuacional fue desarrollado por Rousseau.

Rasiowa formalizó el cálculo de predicados de múltiples valores de Post. Las fórmulas de este cálculo de predicados se puede interpretar como funciones del universo de una interpretación dentro de una cadena $e_0 \leq e_1 \leq \dots \leq e_{m-1}$ donde e_0 y e_{m-1} juegan el papel de falso y cierto respectivamente.

La tesis fundamental de la teoría de conjuntos borrosos es que el conjunto de función de membresía puede ser parcial o gradual, esto es, no tiene porque ser siempre como en la teoría ordinaria de conjuntos. Un conjunto borroso F no tiene porque tener una frontera definida, nítida, y el grado de membresía de un elemento u en él, está dado por una función de membresía $\mu_F: U \rightarrow L$, donde U es el universo del discurso y L es un conjunto parcialmente ordenado. En la práctica L se elige como el intervalo real $[0,1]$ con $\mu_F(\cdot) = 1$ que corresponde a la total pertenencia y $\mu_F(\cdot) = 0$ que es total no pertenencia, y $\mu_F(\cdot) = .5$ que es total borrosidad ya que no se sabe si pertenece o no, ya que tanto pertenece como no pertenece.

En este contexto es fácil ver que si uno decide representar la extensión del predicado de primer orden F con un conjunto borroso, luego, cuando se interprete, la verdad de F estará en el intervalo $[0,1]$. En este caso podemos decir que F es un predicado borroso, vago, y la lógica que subyace detrás, será una lógica multivaluada con la ley del medio excluido no válido. Lo anterior es debido al valor de membresía $.5$ que determina un absoluto desconocimiento respecto a si un objeto tiene o no la característica F , por lo que el valor 0.5 lo podemos interpretar como:

1. No existe evidencia concluyente de que la proposición “ u tiene la propiedad F ” sea verdadera.
2. No existe evidencia concluyente acerca de que la proposición “ u tiene la propiedad F ” no sea verdadera.
3. No existe evidencia que sugiera la posibilidad de que u tenga la propiedad F ”.

Esto es similar al tipo de conocimiento que se da en sistemas de razonamiento monotónico que aluden algún hecho de que “ u tenga la propiedad F ”, para ser asegurada como verdad si existe la posibilidad de que sea cierta y nunca se contradiga esto. Luego la existencia natural de este valor especial de verdad de 0.5 en la lógica borrosa, cae en la profundidad del llamado razonamiento frustrante: queremos la opción de quitar “ u tiene la propiedad F ” si decrecemos el valor de la función de membresía de u en F , es decir llevarla a no ser, o incrementar el valor de su función de membresía para llevarlo a totalmente ser.

Es bastante sorprendente que en todas las lógicas borrosas existentes, la presencia de este valor especial de verdad siempre se silencie por hacerlo miembro de un conjunto de valores designados de verdad, o por modificar la noción de satisfacibilidad de una fórmula, de manera que sea válida en lógica borrosa, si y sólo si ésta es válida en la lógica clásica de primer orden. El resultado subsecuente es que la lógica de primer orden es un caso especial de lógica borrosa en el conjunto $\{0,1\}$. De esta manera la lógica borrosa es reducida a una lógica clásica de primer orden, y por lo mismo es una lógica monotónica. Sin embargo podemos decir que el valor de verdad de 0.5 es parte del conocimiento que se tiene de algo, por lo que se debe de tratar a $F(u)$ como una suposición en lugar de un hecho.

Tradicionalmente los sistemas de lógica formal se han caracterizado por contener los siguientes elementos:

1. Un lenguaje formal L define el conjunto de fórmulas bien formadas
2. Un conjunto de fórmulas bien formadas llamadas axiomas.
3. Un conjunto de reglas de inferencia para derivar teoremas de axiomas.

Las diferentes lógicas se pueden distinguir por un mínimo conjunto de axiomas y teoremas cerrados bajo las reglas de inferencia. Además extendiendo el conjunto de axiomas nunca se podrá impedir la derivación de teoremas ya derivados del conjunto original de axiomas. Como es su intención esta propuesta ha probado ser adecuada para modelar el dominio de las matemáticas, pero desafortunadamente ha probado no ser adecuada para modelar el dominio del razonamiento del sentido común por dos razones: primero porque los agentes de razonamiento son forzados a obtener conclusiones basados en una especificación incompleta de la información relevante. Se hacen suposiciones acerca de la información perdida y se derivan conjeturas basadas en las suposiciones en lugar de teoremas. Las nuevas evidencias pueden probar las suposiciones inválidas, y no se derivan más conjeturas, o que nos muestra la naturaleza no monotónica del formalismo del sentido común. Segundo, la caracterización de la lógica como un mínimo conjunto de teoremas, generados por axiomas y reglas de inferencia es muy limitante para la caracterización del tipo de lógica necesaria para modelar el razonamiento humano. Una propuesta mejor es concentrarse en que sigue a que, o de manera más formal, la noción semántica de vínculo. En este sentido los axiomas y las reglas de necesidad se representan como las fórmulas bien formadas en un lenguaje lógico. Este conjunto se puede llamar el conjunto de premisas, y se puede considerar como el conjunto inicial de hipótesis que limita el conjunto de modelos considerados cuando se emplea la relación de vínculo.

Las suposiciones acerca de la información perdida y las conjeturas resultantes, se clasifican como necesidades. Algunas de las características más importantes de razonamiento por necesidad son:

1. Las conjeturas se realizan en un contexto de información incompleta.
2. Se realiza alguna forma de verificación de posibilidad antes de brincar a alguna conclusión, asegurando que la conjetura es coherente al contexto del razonamiento actual.
3. Las conjeturas son tan débiles como sus derivadas en la forma normal.

2.2. BORROSIDAD

Sea $P(x)$ un predicado de primer orden donde P es un símbolo de predicado y x es una variable individual. Sea I una interpretación bajo un dominio no vacío U que asigna a P una relación unaria borrosa en forma de una función de pertenencia $\mu_P: U \rightarrow [0,1]$. Si p es un predicado eneario, luego I puede asignarle una relación borrosa enearia caracterizada por una función de pertenencia de n lugares.

Además dada I y una valuación V que asigna a x un elemento de U , el valor de verdad de la proposición $P(u)$, es el grado de membresía de u en μ_P . Luego, la verdad de un

predicado borroso es gradual en contraste con los predicados precisos cuyo valor de verdad está en el conjunto $\{0,1\}$.

Se puede hacer una fuerte distinción entre la noción de verdad graduada, debido a la existencia de predicados borrosos en la presencia de información completa y grados de incertidumbre debidos a la información incompleta o incierta como exactamente se debe de tomar valores de verdad borrosos o un predicado preciso. Se pueden identificar los siguientes tres casos:

1. El caso de proposiciones inciertas que son siempre ciertas o falsas, pero debido a la falta de información completa sólo se puede estimar que tan posible o necesario sea que la proposición sea cierta.
2. El caso de las proposiciones borrosas que caen en grados intermedios de verdad que se puede especificar con absoluta certeza.
3. El caso general de proposiciones borrosas inciertas, las cuales debido a la presencia de predicados borrosos, tienen valores de verdad intermedios, pero la falta de información completa no permite que los grados de verdad sean especificados con absoluta certeza.

2.3. LÓGICA BORROSA DE PRIMER ORDEN

Se puede asumir la estructura de los valores de verdad de la forma:

$$L = (L, \vee, \wedge, \otimes, \rightarrow, 0, 1)$$

donde 0, 1 son los elementos menor y mayor respectivamente $\vee, \wedge$, son las operaciones de supremo e ínfimo respectivamente, $\otimes$ es la operación de producto, la cual es isotónica en ambas variables y $(L, \otimes, 1)$ es un monoide conmutativo, y $\rightarrow$ es la operación de residuo, la cual es antitóna en la primera variable e isótóna en la segunda. La propiedad de adjunción

$$A \otimes B < c \text{ si } a < b \rightarrow c$$

vale para toda $a, b, c, \in L$.

Como un caso particular que se usa generalmente en aplicaciones, los grados de verdad se asumen basados en el intervalo unitario de números reales $[0,1]$. Luego si $\vee = \max, \wedge = \min$,

$$A \otimes B = 0 \vee a + b - 1) \text{ y } a \rightarrow b = 1 \wedge (1 - a + b) \forall a, b \in [0,1].$$

El uso de este intervalo es muy natural pero algunas veces causa confusión entre valores de verdad y probabilidades. Insistiremos que es sólo una similitud formal, como el mismo conjunto de medidas se han usado para los dos que en esencia son fenómenos diferentes.

Para modelar la semántica del lenguaje natural es necesario introducir nuevas operaciones adicionales

$$a \leftrightarrow b = (a \rightarrow b) \wedge (b \rightarrow a)$$

que es la biresiducción y

$$a^p = a \otimes a \dots \otimes a$$

que es la potencia $\forall a, b \in [0,1]$.

El lenguaje básico de la lógica borrosa de primer orden consiste de:

1. Variables $x, y, \dots$
2. constantes $c, d, r, \dots$
3. símbolos funcionales $f, g, \dots$
4. un conjunto de símbolos para tablas de verdad $\{a: a \in L\}$
5. símbolos predicados $p, q, \dots$
6. conectivos binarios $\Rightarrow$ y un conjunto de conectivos adicionales $\{o_i: i \in J\}$
7. un símbolo para el cuantificador general “para cada”

2.4. FORMULISMO

1. Un símbolo a para un valor de verdad $a \in L$ es una fórmula atómica.
2. Si $q_1, \dots, q_n$ son términos y p es un símbolo predicado, $p(q_1, \dots, q_n)$ es una fórmula atómica.
3. Si $A, B, A_1, \dots, A_m$ son fórmulas y “o” “y” conectivos, luego $A \Rightarrow B$, o $(A_1, \dots, A_m)$, y para cada x , A son fórmulas.

La fórmula $\neg A$ es una abreviación de $A \Rightarrow 0$. Similarmente podemos definir $A \vee B$ disyunción, $A \wedge B$ conjunción, $A \& B$ conjunción fuerte, $A \Leftrightarrow B$ equivalencia, (debe existir x) A (cuantificador existencial) y $A^k = A \& A \& \dots \& A$ potencia.

Como en lógica clásica, introducimos los conceptos de variables libres, variables acotadas y el de término sustituible. Si q es un término y A es una fórmula, luego $A_x[q]$ es una fórmula resultante de A cuando se sustituye el término q en lugar de cada ocurrencia libre de x en A . El trabajar con muchos valores de verdad es una característica de la lógica borrosa, por lo que es permisible introducir un conjunto borroso de axiomas. De cualquier modo esto nos lleva a una sintaxis evaluada, donde cada paso al derivar nuevas fórmulas, es evaluado por valores de verdad. Luego podemos extender la noción de una regla de inferencia de la forma siguiente: Una regla n -aria de inferencia es un doblete de la forma:

$$r = (r^{\sin}, r^{\text{sem}})$$

donde $r^{\sin}$ es la parte sintáctica y r^{sem} es la parte semántica. La parte sintáctica es una operación n -aria parcial en F_j y la parte semántica es una operación n -aria en L que preserva uniones arbitrarias no vacías en cada argumento. La parte sintáctica tiene el mismo papel que en la lógica clásica, y la parte semántica evalúa cada paso en el proceso de derivación de nuevas fórmulas.

Un conjunto borroso X de fórmulas es cerrado con respecto a r si:

$$X(r^{\sin}(A_1, \dots, A_n)) \supset r^{\text{sem}}(X(A_1), \dots, X(A_n))$$

vale para toda $A_1, \dots, A_n \in F_j$ para la cual se define $r^{\sin}$. Una regla de inferencia es sonora si la ecuación anterior es válida para cualquier función $G: f_f \rightarrow L$ que es un homomorfismo con respecto a los conectivos.

Las reglas de inferencia se escriben usualmente de la forma:



donde $a_i \in L$ son valuaciones verdaderas de las fórmulas A_i .

Las siguientes también son reglas de inferencia:

1. $r^{\text{M.P.}}: \frac{A, A \Rightarrow B}{B} \left(\frac{a, b}{a \otimes b} \right)$ que es el modus ponens.
2. $r^{\text{R.a.}}: \frac{B}{a \Rightarrow B} \left(\frac{b}{a \rightarrow b} \right)$ que es la regla de levantamiento a.
3. $r_G: \frac{A}{(\text{para cada } x)A} \left(\frac{a}{a} \right)$ que es la regla de generalización

Sea X el conjunto de fórmulas, luego:

$$(C^{\sin}X)A = \Lambda \{U(A): U \subseteq F_j \text{ U}$$

es cerrada con respecto a todas las reglas de inferencia y

$$A, \{X \subseteq U\}$$

define un conjunto borroso de consecuencias sintácticas del conjunto borroso X .

Una prueba de la fórmula A de un conjunto borroso X , es una secuencia:

$$w: A_1[a_1;P_1], A_2[a_2;P_2], \dots, A_n[a_n;P_n]$$

tal que A_n es A y P_i , $i < n$ es AL o AE si A es un axioma lógico o especial respectivamente, o P_i es si A es una fórmula

$$r^{\sin}(A_{i_1}, \dots, A_{i_n}), i_1, \dots, i_n < n$$

y r es una regla de inferencia n -aria sonora. Por lo tanto:

$$a_i = \text{Val}_x(w(i))$$

es un valor de la prueba:

$$w(i): A_1[a_1;P_1], A_2[a_2;P_2], \dots, A_n[a_n;P_n]$$

definida de la forma:



Teorema.

$$(C^{\sin}X)A = V\{Val_x(w): w \text{ es una prueba de } a \text{ desde } X \subseteq F_j\}$$

Luego encontrando una prueba, sea w , de una fórmula A asegura solamente que el grado en el cual A es un teorema es mayor o igual a $Val_x(w)$. Si $Val_x(w) < 1$, luego es difícil asegurar que no podemos encontrar una prueba con mayor valor.

2.5. TEORÍAS DE LÓGICA BORROSA DE PRIMER ORDEN

Una teoría T en el lenguaje J de la lógica borrosa de primer orden (teoría borrosa) es un triplete:

$$T = (A_L, A_s, \mathbf{R}),$$

donde $A_L \subseteq F_J$ es el conjunto borroso de los axiomas lógicos, $A_s \subseteq F_J$ es un conjunto borroso de axiomas especiales y R es el conjunto de reglas de inferencia que deben al menos contener las reglas r_{MP} , $\{r_{Ra}; a \in L\}$ y r_G . Por $J(\Gamma)$ denotamos el lenguaje en la teoría borrosa L . El cálculo de predicados borrosos es la teoría borrosa con $A_s = \emptyset$.

Sea $\mathbf{D}$ una estructura para $J(\Gamma)$. Luego $\mathbf{D}$ es un modelo de la teoría Γ , $\mathbf{D} \models \Gamma$, si:

$$A_s(A) \leq \mathbf{D}(A),$$

es válida para toda $A \in F_J$. Tenemos $A_L(A) \leq \mathbf{D}(A)$ en cualquier modelo $\mathbf{D} \models \Gamma$ para toda fórmula $A \in F_J$. Si $(C^{\text{sem}}X)A = a$, luego la fórmula A es cierta en grado a en la teoría Γ y podemos escribir:

$$\Gamma \models aA.$$

Si $(C^{\text{sem}}X)A = a$, entonces A es un teorema en grado a de la teoría Γ y se escribe:

$$\Gamma \vdash aA$$

Escribimos $\Gamma \vdash A$, $\Gamma \models A$, en lugar de $\Gamma \vdash_1 A$, $\Gamma \models_1 A$ respectivamente y decimos que A es un teorema cierto de la teoría Γ .

Si w es una prueba en la teoría Γ , luego escribimos $Val_T(w)$ para este valor.

Una teoría Γ es contradictoria si existe una fórmula A y pruebas w y w' de A y $\neg A$ respectivamente, tal que:

$$\text{Val}_T(w) \otimes \text{Val}_T(w') > 0$$

y es consistente en el caso opuesto.

El teorema de cerradura nos dice que: Sea $A \in F_J$ y A' es su cerradura, luego:

$$\Gamma \mid \text{---} aA \text{ si } T \mid \text{---} A'$$

Un lenguaje J' es una extensión de J si $J \subseteq J'$. Luego es obvio que

$$F_{J(T)} \subseteq F_{J'(T)}.$$

Sean:

$$\Gamma = (A_L, A_s, \mathbf{D}), \Gamma' = (A'_L, A'_s, \mathbf{R}),$$

teorías en sus respectivos lenguajes. Pongamos

$$\boxed{} \text{ si } A \in F_J \text{ y } \overline{A}'_s(A) = 0$$

en cualquier otro caso. Si

$$\boxed{},$$

luego Γ' es una extensión de T . Para simplificar la notación rescribiremos A_s en lugar de $\overline{A}_s$, y entenderemos que $A_s(A) = 0$ para toda $A \in F_{J'} - F_J$.

La extensión Γ' es una extensión conservativa de Γ si $\Gamma' \mid \text{---} bA$ y $\Gamma \mid \text{---} aA$ implica $a = b$ para cualquier fórmula $A \in F_J$.

El teorema de deducción nos dice que si A es una fórmula cerrada y $\Gamma' = \Gamma \cup \{1/A\}$, entonces:

1. Si $\Gamma \mid \text{---} A^n \Rightarrow B$ y $\Gamma' \mid \text{---} bB$ para alguna n , luego $a \leq b$.
2. Para cada prueba w' de B en Γ' hay n y una prueba w de $A^n \Rightarrow B$ en Γ tal que:

$$\text{Val}_{T'}(w') = \text{Val}_T(w)$$

Sea Γ una teoría consistente, $\Gamma \mid \text{---} (\exists x)(A(x))^n$ para cada n , y $t \in J(\Gamma)$ sea una nueva constante, luego la teoría

$$\Gamma' = \Gamma \cup \{1/A_x(t)\}$$

en el lenguaje $J(\Gamma) \cup \{t\}$ es una extensión conservativa de la teoría Γ .

Sea A una fórmula de Γ , p un nuevo símbolo predicado n -ario y sea:

$$\Gamma' = \Gamma \cup \{1/p(x_1, \dots, x_n) \Leftrightarrow A\}$$

luego Γ' es una extensión conservativa de Γ , y para cada fórmula B de Γ' hay una fórmula B^* de Γ tal que:

$$\Gamma' \mid \text{---} B \Leftrightarrow B^*$$

Sea A_0 la fórmula cerrada elegida y $\Gamma \mid \text{---} aA_0$. Propongamos:



Lo que sigue lo podemos ver como una generalización de los teoremas de completos de Gödel

Una teoría borrosa Γ es consistente si tiene un modelo. Si Γ es consistente, luego cada $A \in F_j$ y b definida en la forma anterior es un modelo $\mathbf{D}$ tal que $\mathbf{D}(A) \leq b$.

Sea Γ' una teoría consistente, luego:

$$\Gamma \mid \text{---} aA \text{ si } T \mid = aA,$$

vale para cada fórmula $A \in F$.

CAPÍTULO 3

3.1. RAZONAMIENTO APROXIMADO

3.2. REGLAS DE TRASLACIÓN DEL RAZONAMIENTO APROXIMADO

3.3. REGLAS DE INFERENCIA EN RAZONAMIENTO APROXIMADO

3.1. RAZONAMIENTO APROXIMADO

El razonamiento aproximado es una herramienta matemática cuya dirección es modelar la forma humana de razonamiento cuando están presentes nociones bajas. La teoría propuesta por Zadeh está basada en las reglas derivadas y en relación con las operaciones en las relaciones borrosas. Las aplicaciones exitosas de la teoría nos han llevado a suponerla como una teoría que funciona bien. Por lo anterior es necesario desarrollar una teoría concisa y formal que nos ayude a comprender el problema. En este sentido la lógica borrosa de primer orden nos sirve como la base matemática de la teoría del razonamiento aproximado.

La formulación de un concepto nos lleva a la formación de agrupamiento de elementos que tienen una cierta propiedad φ . En otras palabras existe una cierta propiedad de objetos φ , tal que es construido por medio de agrupar

$$X = \{x: \varphi(x)\},$$

de todos los elementos x que tienen la propiedad φ .

Consideremos el elemento x_0 y preguntemos si tiene la propiedad φ . La respuesta puede tomar la forma de una proposición como $\varphi(x)$ es cierta en grado α . Además muchas palabras del lenguaje natural pueden ser interpretadas como nombres de la anterior propiedad φ . Luego en principio los conjuntos borrosos se pueden usar para modelar la semántica del lenguaje natural.

Una característica especial del pensamiento humano es el uso efectivo del lenguaje natural, aún en el proceso del razonamiento lógico, por lo que podríamos afirmar que el modelo matemático del razonamiento humano se puede basar en los conjuntos borrosos y en la lógica borrosa.

En lo que continúa consideraremos un conjunto S de sintagmas¹ del lenguaje natural, y un lenguaje de la lógica borrosa de primer orden S , que se usará para modelar el significado de los sintagmas de S . Denotaremos los elementos de S por $A, B, \dots$. Sea M_j y F_j el conjunto de términos y fórmulas de un lenguaje S respectivamente. Luego la teoría del razonamiento aproximado consiste especialmente en la formación de dos clases de reglas:

1. Reglas de traslación. Por ejemplo las reglas para obtener conjuntos borrosos de los sintagmas tomados de S .
2. Reglas de inferencia. Por ejemplo las reglas para obtener conclusiones de algunas premisas establecidas en la forma de los elementos de S .

3.2. REGLAS DE TRASLACIÓN DEL RAZONAMIENTO APROXIMADO

Una regla de traslación es una tripleta:

¹ Un sintagma es una parte de una sentencia que se construye de acuerdo con las reglas gramaticales.

$$s = \langle s^{\text{form}}, s^{\text{fuzz}}, s^{\text{ext}} \rangle$$

donde

$$s^{\text{form}}: \mathbf{S} \rightarrow \mathbf{S} \cup F_j \cup M_j$$

es una función parcial de asignación de fórmulas $A(x) \in F_j$, en general abiertas, a algún sintagma $\mathbf{A} \in \mathbf{S}$.

s^{fuz} es una función de asignación que asigna un conjunto borroso

$$s^{\text{fuzz}}(A(x)) = \{a_t/A_x[t]: t \in M_j\} \subseteq F_j,$$

de una fórmula cerrada a $A(x)$.

s^{ext} es una función que asigna una extensión

$$s^{\text{ext}}(s^{\text{fuzz}}(A(x))) = \{a_t/t; t \in M_j\} \subseteq M_j$$

al conjunto borroso $s^{\text{fuzz}}(A(x))$. Escribiremos

$$A(x) \text{ en lugar de } s^{\text{form}}(\mathbf{A}), Z(A(x)), \text{ en lugar de } s^{\text{fuz}}(A(x)),$$

y

$$E(A(x)) \text{ en lugar de } s^{\text{ext}}(s^{\text{fuzz}}(A(x))).$$

Para ser independientes del modelo concreto, consideremos el conjunto M_j de todos los términos. Puede denotar los elementos de cualquier modelo. Más allá de nuestras deliberaciones procederemos dentro de la sintaxis en lugar que dentro de la semántica. Nótese que la misma $Z(A(x))$ puede tener diferentes interpretaciones en diferentes interpretaciones en varios modelos $\mathbf{D}$.

Luego un sintagma $\mathbf{A}$ es interpretado como una fórmula $A(x)$, que representa la propiedad expresada por $\mathbf{A}$. Como la propiedad es generalmente vaga, la representamos con un conjunto borroso de las fórmulas cerradas $Z(A(x))$. La fórmula $A(x)$ junto con $Z(A(x))$ también se puede entender como una expresión formal de una intensión de la propiedad bajo consideración. Esta extensión es luego un conjunto borroso de elementos $E(A(x))$.

Por ejemplo, sea A : = la temperatura es alta. Luego $A(x)$ es una fórmula que representa la propiedad de ser alta. Es la propiedad de elementos (grados) a través de los cuales la variable x (temperatura) varía. Luego en una traslación de $A(x)$ a un conjunto borroso de fórmula cerrada

$$\{a_t/A_x[\gamma]; t \in M_j\}.$$

Cada $A_x[\gamma]$ es una expresión formal de la proposición la temperatura es alta y se le asigna un valor de verdad a_γ . El último paso $E(A(x))$, nos lleva a un conjunto borroso abstracto de temperaturas $\gamma \in M_j$ sin significado especial.

La función s^{form} se puede definir de la forma siguiente: Sea N un nombre. Luego:

$$S^{\text{form}}(N) = x.$$

donde x es alguna variable². Sea A un adjetivo, luego:

$$s^{\text{form}}(AN) = s^{\text{form}}(N \text{ es } A) = A(x)$$

donde (Ax) es una fórmula representando la propiedad llamada por el adjetivo A . Sea m un modificador lingüístico, por lo que:

$$s^{\text{form}}(m) = o_m$$

donde o_m es un conectivo unario y:

$$s^{\text{form}}(mAN) = s^{\text{form}}(N \text{ es } mA) = o_m(A(x))$$

Sea B_∞, B_ϵ son sintagmas tales que

$$s^{\text{form}}(B_1) = B_1(y) \text{ y } s^{\text{form}}(B_2) = B_2(z). \text{ Si } 'B_\infty \text{ y } B'_\epsilon \in \mathbf{S}, B'_\infty \text{ y } B'_\epsilon \in \mathbf{S} \text{ y si } B'_\infty$$

luego $B'_\epsilon \in \mathbf{S}$, luego:

$$s^{\text{form}}(B_1 \text{ y } B_2) = B_1 \nabla B_2$$

$$s^{\text{form}}(B_1 \text{ o } B_2) = B_1 \square B_2$$

$$s^{\text{form}}(\text{si } B_1 \text{ luego } B_2) = B_1 \Rightarrow B_2$$

donde ∇ denota un conectivo del tipo conjuntivo, $\square$ denota un conectivo del tipo disyuntivo, y $\Rightarrow$ denota un conectivo del tipo implicativo.

Un punto crucial es la determinación de la función s^{form} . Desde el punto de vista de la lógica borrosa podemos considerar las siguientes clases de $Z(A(x))$.

$$Z_s(A(x)) = \{a_t/A_x[\gamma]: \gamma \in M_j \text{ y } a_t = A_s(A_x[\gamma])\}$$

cada $a_t, \gamma \in M_j$, es un grado de verdad que $A_x[\gamma]$ es un axioma especial.

$$Z_{|-}(A(x)) = \{a_t/A_x[t]: \gamma \in M_j \text{ y } \Gamma \vdash a_t A_s(A_x[t])\}$$

cada $a_t, \gamma \in M_j$, es un grado de probabilidad de $A_x[t]$ en una teoría dada Γ .

$$Z_D(A(x)) = \{a_t/A_x[\gamma]: t \in M_j \text{ y } a_t = \mathbf{D}(A_x[\gamma])\}$$

cada $a_t, \gamma \in M_j$, es un grado de verdad de $A_x[\gamma]$ en el modelo dado $\mathbf{D}$.

$$Z_{|=}(A(x)) = \{a_t/A_x[\gamma]: \gamma \in M_j \text{ y } \Gamma \models a_t(A_x[\gamma])\}$$

cada $a_t, \gamma \in M_j$, es un grado de verdad de $A_x[\gamma]$ en una teoría dada Γ .

Por el teorema de completos

$$Z_{|-}(A(x)) = Z_{|=}(A(x))$$

² De aquí en adelante entenderemos que varios N se asignan a varias variables. Por simplicidad no lo expresaremos en la notación

es válido para cualquier fórmula en A .

Los conjuntos borrosos Z_s y Z_D son suplementarios, ya que en el razonamiento aproximado deseamos obtener intensiones de deliberación, no sintagmas originales de conocimiento. Como una relación muy compleja entre fórmulas que debe de existir, $Z|_ =$ es la que mejor podemos evaluar, y nos da inmediatamente $Z|_ =$.

Sea $C: \rightarrow L$ un homomorfismo y $E \subseteq F_j$. Pongamos $A_s(B) = C(B)$ para toda $B \in E$ y $A_s = 0$ en otro caso. Luego:

$$(C^{\sin} A_s)B = A_s(B), \text{ para toda } B \in E.$$

En el razonamiento aproximado, debemos definir un conjunto de fórmulas cerradas $s^{\text{borr}}(s^{\text{for}}(A))$ para algunas palabras primarias. Estos conjuntos borrosos en efecto son subconjuntos borrosos de conjuntos borrosos de axiomas especiales. Así podemos entender el razonamiento aproximado como una derivación de nuevas fórmulas en la teoría borrosa definida por el conjunto borroso de nuevos axiomas.

3.3. REGLAS DE INFERENCIA EN RAZONAMIENTO APROXIMADO

La situación del razonamiento aproximado es la siguiente: Dados un conjunto I de sintagmas y un lenguaje J , nuestro objetivo es conocer la intensión de todos los sintagmas $A \in I$, como todas las fórmulas:

$$A(x) \in F_j$$

y un conjunto de fórmulas:

$$Z|_ (A(x)) \subseteq F_j$$

El conjunto borroso anterior se obtuvo en el marco de la teoría borrosa determinada por un conjunto borroso de axiomas, $A_s \subseteq F_j$, que consiste de:

1. Un subconjunto borroso $Z_s(A(x))$ para algunas palabras primarias $A \in I$.
2. Un subconjunto borroso $Z_s(B(y))$ para algunos sintagmas adicionales (declaraciones) $B \in I$, que usamos para describir el proceso.

Suponemos $Z_s = Z|_$ en 1 y 2. En muchos casos, el sintagma (2) representa declaraciones condicionales, pero en general no están restringidos a ellos. Luego la tarea del razonamiento aproximado es encontrar $A(x)$ y $Z|_ (A(x))$ para alguna A , o viceversa. Esta tarea es difícil ya que se tiene que encontrar las pruebas con los objetos $A_x [\gamma]$ para todas $\gamma \in M_j$. La solución puede ser más o menos aproximada usando las reglas de inferencia que definiremos.

Dada una regla de inferencia r , de lógica borrosa de primer orden, una regla de inferencia $\mathbf{R}$ en razonamiento aproximado basado en una regla r es un cuádruple de funciones:

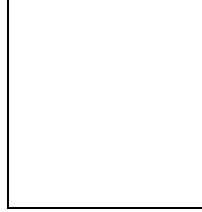
$$\mathbf{R} = (\mathbf{R}^{\text{ling}}, r^{\text{sin}}, \mathbf{R}^{\text{borr}}, \mathbf{R}^{\text{ext}})$$

que corresponden a los cuatro niveles de traslación y que se definen como sigue:

En el nivel de lenguaje $\mathbf{R}^{\text{lin}}: \mathbf{I}^n \rightarrow \mathbf{I}$ es una función parcial que se puede escribir en una forma más transparente como:



donde $\mathbf{A}_1, \dots, \mathbf{A}_n, \mathbf{B} \in \mathbf{I}$. El sintagma $\mathbf{B}$ es tal que el diagrama:



conmuta donde r^{sin} es una parte sintáctica de r .

Sea
 $s^{\text{form}}(\mathbf{A}_i) = \mathbf{A}_i(x_i, y), y = 1, 2, \dots, n$ y $s^{\text{form}}(\mathbf{B}) = \mathbf{B}(y)$.

luego la función $\mathbf{R}^{\text{lin}}$ es una representación lingüística de la prueba:

$$w: \mathbf{A}_1(x_1, y), \dots, \mathbf{A}_n(x_n, y), r^{\text{sin}}(\mathbf{A}_1(x_1, y), \dots, \mathbf{A}_n(x_n, y)): = \mathbf{B}(y)$$

Los siguientes pasos conducen a los conjuntos borrosos a una fórmula cerrada

$$Z(\mathbf{A}_i(x_i, y)) = Z(\mathbf{B}(y)).$$

Como podemos considerar cualquier etapa de inferencia,

$$Z|_{\text{—}}(\mathbf{A}_i(x_i, y)) = Z(\mathbf{A}_i(x_i, y))$$

no es necesaria. Luego, si:

$$Z(\mathbf{A}_i(x_i, y)) = \{\mathbf{A}_{\text{tis}}/\mathbf{A}_{\text{xiy}}[\gamma_i, s]: \gamma_i s \in M_j\}$$

luego $a_{\text{tis}} \geq A_s(A_{\text{xis}}[\gamma_i, s])$, pero no es necesario que sea un máximo. Dados los términos $\gamma_1, \dots, \gamma_n, s \in M_j$, podemos considerar la prueba:

$$w_{t_1, \dots, t_n, s} = \mathbf{A}_{1, x_1 y}[\gamma_1 s, s][a_{\gamma_1 s}], \dots, \mathbf{A}_{n, x_n y}[\gamma_n s, s][a_{\gamma_n s}], \mathbf{B}_y[s][r^{\text{sem}}(a_{\gamma_1 s}, \dots, a_{\gamma_n s}): r]$$

Por lo que para cada $s \in M_j$ existe un conjunto de pruebas:

$$W_s = \{w_{\gamma_1, \dots, \gamma_n, s}: \gamma_1, \dots, \gamma_n \in M_j\}$$

siendo un subconjunto de todas las pruebas de $\mathbf{B}_y[s]$.

Podemos dar la siguiente definición:

$$\mathbf{R}^{\text{borr}}: (L^{\text{fj}})^n \rightarrow L^{\text{fj}}$$

es una función parcial n-aria que asigna un conjunto borroso:

$$Z(B(y)) = \{b_s / \mathbf{B}_y [s] : s \in M_j\}$$

al conjunto borroso:

$$Z(A_i(x, y)) = \{a_{i, \gamma_{is}} / \mathbf{A}_{i, xiy} [\gamma_i, s] : \gamma_i, s \in M_j\}$$

tal que:

$$b_s = \{r^{\text{sem}}(a_{\gamma_{1s}}, \dots, a_{\gamma_{ns}}) : \gamma_1, \dots, \gamma_n \in M_j\}$$

para todo $b_s, s \in M_j$.

La ecuación de b_s es exactamente la fórmula para calcular los grados de membresía de los consecuentes en la regla de razonamiento aproximado. Se justifican con base en la lógica de primer orden, por lo que los grados b_s son solamente una pobre estimación de los grados de probabilidad de $\mathbf{B}_y[s]$, $s \in M_j$, ya que puede haber otras pruebas de esta fórmula con mayores valores, se tiene:

$$Z(B(y)) \subseteq Z|_{\text{—}} (B(y))$$

por lo que debe haber condiciones especiales bajo las cuales la inclusión se puede cambiar en igualdad, así: Sea $s \in M_j$. si hay un modelo $\mathbf{D}$ tal que:

$$\mathbf{D}(\mathbf{B}_y[s]) = c = \mathbf{V}\{r^{\text{sem}}(a_{\gamma_{1s}}, \dots, a_{\gamma_{ns}}) : \gamma_1, \dots, \gamma_n \in M_j\}$$

para alguna $a_{\gamma_{1s}}, \dots, a_{\gamma_{ns}} \in M_j$, luego $\Gamma |_{\text{—}} b_s \mathbf{B}_x [s]$, $s \in M_j$.

Esta simple proposición se hace especial cuando $\mathbf{A}_1, \dots, \mathbf{A}_n$ son las únicas declaraciones lingüísticas conducidas al conjunto borroso de axiomas especiales. Demostraremos que es la regla del modus ponens. Sea:



$$s^{\text{form}}(A) = A(x), s^{\text{form}}(\text{si } \mathbf{A} \text{ luego } \mathbf{B}) = A(x) \Rightarrow \mathbf{B}(y)$$

y

$$A_s(A_x[t]) = a_\gamma, \gamma \in M_j$$

$$A_s(A_x[\gamma]) \Rightarrow \mathbf{B}_y[s] = b_{\gamma s}, \gamma, s \in M_j$$

$$A_s(C) = 0 \text{ para el otro, } C \in F_j$$

Sea $\mathbf{D}_0$ una estructura canónica donde:

$$\mathbf{D}_0(A_x[\gamma]) = a_\gamma, \gamma \in M_j$$

$$\mathbf{D}_0(\mathbf{B}_y[s]) = \mathbf{V}\{a_\gamma \otimes b_{\gamma s} : \gamma \in M_j\} = c_s, s \in M_j$$

donde:

$$\mathbf{D}_0(A_x[\gamma] \Rightarrow \mathbf{B}_y[s]) = a_t \rightarrow c_s \geq b_{\gamma s}$$

si

$$a_t \otimes b_{\gamma s} \leq c_s$$

y por lo tanto $\mathbf{D}_0$ es un modelo. Luego las reglas del modus ponens generalizado usadas en el razonamiento aproximado, hasta aquí se pueden considerar como las que nos dan

el mejor resultado cuando se asume que se está trabajando en una teoría restringida con **A** y si **A** luego **B** como el único axioma especial.

El último paso concierne a la función:

$$\mathbf{R}^{\text{ext.}}: (L^{\text{mj}})^n \rightarrow L^{\text{mj}}.$$

Definimos su analogía como $\mathbf{R}^{\text{borr}}$ cuando reemplazamos todas las ocurrencias de $A_{x_i y}[t_i, s]$ o $B_y[s]$ simplemente con (γ_i, s) o s respectivamente. Las extensiones $E(A_i(x_i, y))$, $E(B(y))$ son conjuntos borrosos abstractos y b_s es la fórmula para la composición de relaciones borrosas.

El nivel de extensiones es muy importante cuando no siempre conocemos los sintagmas resultantes **B** y siempre la fórmula $B(y) = s^{\text{lin}}(\mathbf{B})$. Luego podemos empezar de un conjunto borroso $C \subset M_j$ obtenido de $E(A_1(x_1, y)), \dots, E(A_n(x_n, y))$, usando b_s y encontrando $\mathbf{B} \in \mathbf{I}$ y $B(y) \in F_j$, tal que $E(B(y)) \approx C$.

CAPÍTULO 4

- 4.1. CONJUNTOS BORROSOS DINAMICOS
- 4.2. NOTACIÓN, TERMINOLOGÍA Y OPERACIONES BÁSICAS.
- 4.3. OPERACIONES ALGEBRAICAS.
- 4.4. RELACIÓN BORROSA DINÁMICA.
- 4.5. PROYECCIÓN Y CONJUNTOS BORROSOS DINÁMICOS CILÍNDRICOS
- 4.6. EL PRINCIPIO DE EXTENSIÓN
- 4.8. CONJUNTOS DINÁMICOS BORROSOS DEL TIPO II

4.1. CONJUNTOS BORROSOS DINAMICOS

La teoría de los conjuntos borrosos dinámicos se puede considerar como un cuerpo de conceptos y técnicas para trabajar en forma sistemática con el tipo de imprecisiones que se presentan cuando las fronteras de una clase de objetos no están claramente definidas y cambian en el tiempo. Un buen ejemplo de estas clases es la edad de un hombre. El hombre en el tiempo puede pertenecer a diferentes conjuntos borrosos. En la primera edad el hombre pertenece a la clase borrosa de los niños. Al pasar el tiempo se es miembro de la clase borrosa de los adolescentes, después pertenece a la clase borrosa de los jóvenes, pasando el tiempo se pertenece a la clase borrosa de los de edad madura y al final se pertenece a la clase borrosa de los viejos. Luego, un conjunto borroso dinámico se puede pensar como una clase en la cual existe una progresión gradual dependiente del tiempo; de no pertenecer a una clase, a pertenecer parcialmente a la clase, a pertenecer totalmente a la clase, a pertenecer parcialmente a la clase y a no pertenecer a la clase. En forma más precisa un objeto puede cambiar su función de pertenencia entre cero (no miembro) a uno (totalmente miembro) en el tiempo. Desde este punto de vista un conjunto borroso en el sentido matemático convencional se puede ver como un caso degenerado de los conjuntos borrosos dinámicos. Esto es, un conjunto borroso sólo admite grados de pertenencia fijos.

La teoría de los conjuntos borrosos dinámicos tiene dos ramas distintas, en el mismo sentido que los conjuntos borrosos ordinarios. En una, un conjunto borroso dinámico es tratado como una construcción matemática relacionada con que se puedan realizar aseveraciones que cambian en el tiempo. En la otra, se pueden ver como una teoría borrosa de los conjuntos borrosos dinámicos, en los cuales se introducen las borrosidades variantes en el tiempo en la lógica, la cual, contiene las reglas de manipulación de los conjuntos borrosos dinámicos, con el sentido de que las aseveraciones acerca de ellos pueden variar en el tiempo. Esta relacionada con el cambio del significado lingüístico que permita el desarrollo de una lógica borrosa dinámica.

En esta lógica, el valor de verdad, así como las reglas de inferencia son imprecisas y cambian en el tiempo, por lo que la lógica borrosa dependiente del tiempo es el fundamento de un razonamiento borroso dinámico.

4.2. NOTACIÓN, TERMINOLOGÍA Y OPERACIONES BÁSICAS.

Un conjunto borroso dinámico generalmente se asume inmerso en un universo del discurso no borroso, el cual puede ser cualquier colección de objetos, conceptos, o construcciones matemáticas. El universo del discurso puede ser un conjunto de números reales, el conjunto de velocidades en un proceso dinámico, etc. Los universos del discurso los denotaremos por los símbolos $U, V, W, \dots$. Un conjunto borroso dinámico en U , o equivalentemente, un subconjunto borroso de U , se denota por los símbolos $A, B, C, D, \dots$.

El uso de los conjuntos borrosos dinámicos surge por la necesidad de asignar grados de pertenencia variantes en el tiempo, bajo ciertos criterios dependientes del tiempo a los elementos de un conjunto dinámico. La frase “muy grande” no está definida, pero es posible establecer nuevas correspondencias en relaciones inductivas dependientes del tiempo. Esta frase puede tener diferentes significados para diferentes individuos en diferentes tiempos.

La asignación que se puede hacer acerca de que tan cerca se encuentran dos puntos se puede representar como una función $\mu_{A(t)}(t)$ del conjunto de los números reales, al intervalo unitario $[0,1]$. Si $\mu_{A(t)}(t) = 1$ luego se es miembro de $A(t)$ y los valores intermedios nos dan menores grados de pertenencia en el tiempo, si $\mu_{A(t)}(t) = 0$ no se es miembro del conjunto.

Formalmente, si $X(t)$ es un conjunto borroso dinámico y $A(t) \subset X(t)$, luego la función:

$$f_{A(t)}: x(t) \rightarrow [0,1]$$

se denominará como la función de membresía de un conjunto borroso dinámico $A(t)$ en el universo del discurso $U(t)$.

Sea $U(t)$ el universo del discurso de los objetos que tienen por elemento genérico a $x(t)$. Luego $X = [x(t)]$, donde $x(t)$ es una función del tiempo y $t \in [0, \infty]$.

Sea $A(t)$ en $U(t)$. $A(t)$ es un subconjunto dinámico en $U(t)$, si podemos caracterizar $U(t)$ por una función $f_{A(t)}[x(t)]$ tal que a cada punto en $U(t)$ asocie un número real en el intervalo $[0,1]$, donde el valor de $f_{A(t)}[x(t)]$ representa el grado de pertenencia temporal de $x(t)$ en $A(t)$.

El soporte de $A(t)$ es el conjunto de puntos en $U(t)$ para el cual $f_{A(t)}[x(t)]$ es positiva en un tiempo t definido. La altura dinámica de $A(t)$ es el supremo de $f_{A(t)}[x(t)]$ en un tiempo t definido sobre $A(t)$. El punto de cruce dinámico de $A(t)$ es un punto en $U(t)$ cuyo grado de membresía en $A(t)$ sea 0.5 en un tiempo t definido. $A(t)$ es normal si su altura es unitaria y subnormal si no es el caso.

Debemos de hacer notar que en muchas aplicaciones el grado de pertenencia de $f_A[x(t)]$, se puede interpretar como el grado de compatibilidad de $U(t)$, con el concepto representado por $A(t)$, en un tiempo definido. En otros casos $f_{A(t)}[x(t)]$ se debe de interpretar como el grado de posibilidad de $U(t)$ dado $A(t)$ en un tiempo definido t .

Un conjunto borroso dinámico $A(t)$ de $U(t)$ se expresa en forma lineal como:

$$A(t) = f_1[x(t)] U_1(t) + f_2[x(t)] U_2(t) + \dots + f_n[x(t)] U_n(t),$$

$$A(t) = \sum_{i=1}^n f_i[x(t)] U_i(t)$$

que es consistente con la representación de un conjunto borroso finito en forma lineal en $U_i(t)$. Un conjunto borroso dinámico arbitrario de U se puede expresar en la forma:



Definición

Sea $x(t)$ un conjunto en el tiempo t y $A(t)$ un subconjunto de $U(t)$. La función de membresía de $A(t)$ es la función:

$$f_A(t), t: x(t) \rightarrow [0,1]$$

El propósito de introducir el tiempo es el de poder analizar las situaciones en las cuales los valores de las funciones de membresía de los objetos en el conjunto borroso cambian con el tiempo.

Una restricción en el tiempo t , es un subconjunto borroso $C(t)$ de $U(t)$. En muchos casos se puede obtener un ordenamiento parcial de $U(t)$, digamos $\beta \leq A(t)$ en un t dado. Decimos $\alpha \leq A(t)$ si cuando α es una alternativa deseable [$\alpha \in A(t)$], luego β es deseable [$\beta \in A(t)$] y si $\beta \notin A(t)$ luego $\alpha \notin A(t)$.

Cada función de membresía $[\mu_A(t), t]$ representa las posibilidades de todas las alternativas al tiempo t , respecto al criterio al tiempo t que nos da $A(t)$. Si $t' > t$ está dado en el último tiempo con respecto al criterio anterior.

Los valores de $f_{A(t)}[x(t)]$ cercanos a la unidad nos dan los valores máximos de la función de membresía para $x(t)$ en $A(t)$ y el valor de $f_{A(t)}[x(t)]$ cercano a cero es el mínimo valor del grado de membresía de $x(t)$ en $A(t)$. Estas funciones en el caso que $f_{A(t)}[x(t)] = f_{A(t)}(x) \forall t$ se reducen a las funciones de membresía de los conjuntos borrosos clásicos, y si la función de membresía sólo puede tomar los valores 1 y 0 se tienen las funciones de la teoría clásica de conjuntos.

Si $A(t)$ es un subconjunto borroso dinámico de $U(t)$, luego el conjunto dinámico de nivel α de $A(t)$ es un conjunto clásico denotado por $A_\alpha(t)$ que contiene todos los elementos de $U(t)$, y cuyo grado de membresía en $A(t)$ es mayor o igual a α . Esto es:

$$A_\alpha(t) = \{U(t) | f_A[x(t)] \geq \alpha\} \forall t.$$

Definición:

Un conjunto borroso dinámico se define como temporalmente vacío si:

$$f_A[x(t)] = 0 \text{ para algún } t.$$

luego:

$$A = \Phi^T$$

En el caso en el cual la función de membresía temporal

$$f_{A(t)}[x(t)] = 0 \forall t,$$

luego

$$A = \Phi$$

Y se dice que tenemos un conjunto borroso dinámico absolutamente vacío.

Sean $A(t)$ y $B(t)$ dos conjuntos borrosos dinámicos en X , Luego $A(t)$ es temporalmente igual a $B(t)$ si para algún t :

$$f_A[x(t)] = f_B[x(t)] \text{ para algún } x(t) \in X, \text{ y en un } t \text{ dado.}$$

Y se denota por:

$$A \overset{T}{=} B$$

En este caso cuando

$$f_A[x(t)] = f_B[x(t)] \forall x(t) \in X \text{ y } \forall t,$$

luego $A(t)$ es absolutamente igual a $B(t)$, y se denota por:

$$A = B.$$

Definimos el complemento de un subconjunto borroso dinámico $A(t)$ como:

$$1 - f_A[x(t)] \forall x(t) \in X, \text{ y algún } t.$$

Y lo denotaremos por A'^T .

Tendremos el complemento absoluto de $A(t)$ que denotaremos por A' , si:

$$1 - f_A[x(t)] \forall x(t) \in X, \text{ y } \forall t.$$

Sean $A(t)$ y $B(t)$ dos subconjuntos borrosos dinámicos en $X(t)$. Decimos que $A(t)$ está temporalmente incluido en $B(t)$ si a un t dado:

$$f_A[x(t)] \leq f_B[x(t)] \in X, \text{ y algún } t.$$

Y lo denotamos por:

$$A \overset{T}{\subset} B$$

si

$$f_A[x(t)] \leq f_B[x(t)] \in X, \text{ y } \forall t.$$

Decimos que A está incluido en B y lo denotamos por:

$$A \subset B.$$

Unión:

La unión temporal de dos conjuntos borrosos dinámicos A y $B \in X$ es otro conjunto borroso dinámico $C \in X$ tal que:

$$f_C[x(t)] = \max\{f_A[x(t)], f_B[x(t)]\}$$

en algún t con $x \in X$, y se denota por:

$$C \overset{\Delta}{=} A \overset{T}{\cup} B$$

La unión de dos conjuntos borrosos dinámicos $A(t)$ y $B(t)$ está definida como el menor de los conjuntos borrosos dinámicos que incluyen temporalmente a $A(t)$ y $B(t)$.

Si:

$$f_C[x(t)] = \max\{f_A[x(t)], f_B[x(t)]\} \quad \forall t \text{ con } x \in X$$

se tiene la unión absoluta de dos conjuntos borrosos dinámicos y se denotará por:

$$C = A \cup B$$

La unión temporal de los conjuntos borrosos dinámicos es asociativa, así sean $A(t)$, $B(t)$ y $C(t)$ conjuntos borrosos dinámicos existentes en X , luego:

$$A \overset{T}{\cup} (B \overset{T}{\cup} C) = (A \overset{T}{\cup} B) \overset{T}{\cup} C$$

Demostración:

$$D = B \overset{T}{\cup} C,$$

donde:

$$f_D[x(t)] = \max\{f_B[x(t)], f_C[x(t)]\} \text{ para cualquier } t \text{ con } x \in X$$

$$A \overset{T}{\cup} D = E$$

donde

$$f_E[x(t)] = \max\{f_A[x(t)], f_D[x(t)]\} \text{ para cualquier } t \text{ con } x \in X.$$

De lo anterior tenemos:

$$f_E[x(t)] = \max\{f_A[x(t)], \max\{f_B[x(t)], f_C[x(t)]\}\} \text{ para cualquier } t \text{ con } x \in X.$$

$$f_E[x(t)] = \max\{f_A[x(t)], f_B[x(t)], f_C[x(t)]\} \text{ para cualquier } t \text{ con } x \in X.$$

$$A \overset{T}{\cup} B = F$$

donde:

$$f_F[x(t)] = \max\{f_A[x(t)], f_B[x(t)]\} \text{ para cualquier } t \text{ con } x \in X.$$

$$F \overset{T}{\cup} C = M$$

donde

$$f_M[x(t)] = \max\{f_F[x(t)], f_C[x(t)]\} \text{ para cualquier } t \text{ con } x \in X.$$

$$f_M[x(t)] = \max\{f_F[x(t)], \max\{f_A[x(t)], f_B[x(t)]\}\} \text{ para cualquier } t \text{ con } x \in X.$$

$$f_M[x(t)] = \max\{f_A[x(t)], f_B[x(t)], f_C[x(t)]\} \text{ para cualquier } t \text{ con } x \in X.$$

Por la definición de identidad temporal de dos conjuntos borrosos dinámicos:

$$M^T = E$$

y por lo tanto:

$$A^T \cup (B^T \cup C) = (A^T \cup B)^T \cup C$$

Se puede también asumir la notación:

$$f_C^T = f_A^T \vee f_B^T$$

para la unión temporal, y:

$$f_C = f_A \vee f_B,$$

para la unión absoluta.

Intersección:

La intersección de dos conjuntos borrosos dinámicos A y B es otro conjunto borroso dinámico I definido por la función de membresía:

$$f_C[x(t)] = \min\{f_A[x(t)], f_B[x(t)]\}$$

para alguna t y se denota por:

$$A^T \cap B^T = I$$

o en forma simplificada:

$$f_I = f_A \wedge f_B$$

La intersección de dos conjuntos borrosos dinámicos A(t) y B(t) se entiende como el mayor de los conjuntos borrosos dinámicos que contiene temporalmente a A(t) y B(t).

La intersección absoluta de dos conjuntos borrosos dinámicos A(t) y B(t) está definida por la función de membresía:

$$f_I[x(t)] = \min\{f_A[x(t)], f_B[x(t)]\} \forall t$$

y se denota por:

$$A \cap B = I$$

O en forma simplificada:

$$f_I = f_A \wedge f_B$$

Se dice que dos conjuntos borrosos dinámicos son temporalmente disjuntos si:

$$A^T \cap B^T = \Phi$$

Con las anteriores definiciones se pueden construir las siguientes propiedades:

1. Leyes de Morgan:

$$(A \cup B)^T = A'^T \cap B'^T$$

Prueba:

$$(A \cup B)^T = 1 - (A \cap B)^T = 1 - \max\{f_A[x(t)], f_B[x(t)]\}^T = \min\{1 - f_A[x(t)], 1 - f_B[x(t)]\}^T = A'^T \cap B'^T$$

$$(A \cap B)^T = A'^T \cup B'^T$$

Prueba:

$$(A \cap B)^T = 1 - (A \cup B)^T = 1 - \min\{f_A[x(t)], f_B[x(t)]\}^T = \max\{1 - f_A[x(t)], 1 - f_B[x(t)]\}^T = A'^T \cup B'^T$$

2. Leyes Distributivas:

$$\begin{aligned} C \cap (A \cup B)^T &= (C \cap A)^T \cup (C \cap B)^T \text{ para alguna } t \\ C \cup (A \cap B)^T &= (C \cup A)^T \cap (C \cup B)^T \text{ para alguna } t \end{aligned}$$

Como un ejemplo se probará el caso:

$$C \cap (A \cup B)^T = (C \cap A)^T \cup (C \cap B)^T \text{ para alguna } t$$

prueba:

$$\begin{aligned} C \cap (A \cup B)^T &= \min\{f_C[x(t)], \max\{f_A[x(t)], f_B[x(t)]\}\}^T = \\ &= \max\{\min f_C[x(t)], f_A[x(t)], \{f_C[x(t)], f_B[x(t)]\}\}^T = (C \cap A)^T \cup (C \cap B)^T \end{aligned}$$

4.3. OPERACIONES ALGEBRAICAS.

a.- Suma temporal algebraica:

La suma temporal algebraica de A y B se denota por $A + B^T$ y está definida por:

$$f_{A+B}^T = f_A^T + f_B^T$$

con la condición:

$$f_A[x(t)] + f_B[x(t)] \leq 1$$

b.- Diferencia absoluta temporal:

La diferencia absoluta de dos conjuntos borrosos dinámicos A y B se denota por:

$$\left| A - B^T \right|$$

y se puede definir de la forma:

$$f_{|A-B|}^T = \left| f_A^T[x(t)] - f_B^T[x(t)] \right| \text{ para alguna } t$$

Nótese que en el caso de conjuntos clásicos, ésta se reduce al complemento relativo de:

$$A \cap B \text{ en } A \cup B.$$

c. Producto.

El producto de dos conjuntos borrosos dinámicos A y B está denotado por:

$$A \cdot B$$

y lo podemos definir de la forma:

$$f_{A \cdot B}^T = f_A^T[x(t)] \cdot f_B^T[x(t)] \text{ para alguna } t$$

d.- Suma acotada:

La suma acotada de dos conjuntos borrosos dinámicos A y B está denotada por:

$$A \oplus B$$

y se puede definir de la forma:

$$f_{A \oplus B}^T = 1 \cap f_A^T[x(t)] + f_B^T[x(t)] \text{ para alguna } t$$

e.- Diferencia acotada.

La diferencia acotada de dos conjuntos borrosos dinámicos A y B se denota por:

$$A \otimes B$$

y se puede definir de la forma:

$$f_{A \otimes B}^T = 1 \cup f_A^T[x(t)] - f_B^T[x(t)] \text{ para alguna } t$$

f.- Combinación Convexa.

Sean A, B y Λ tres conjuntos borrosos dinámicos arbitrarios. La combinación convexa de A, B y Λ se denota por:

$$(A, B; \Lambda = \Lambda A + \Lambda' B,$$

donde Λ' es el complemento temporal de Λ . Expresado en términos de su función de pertenencia:

$$f_{[A,B;A]}^T[x(t)] = f_A[x(t)]f_A[x(t)] + \{1 - f_A[x(t)]f_B[x(t)]\}$$

para alguna t.

f.- Producto Cartesiano.

Si $A_1, \dots, A_n$ son subconjuntos borrosos dinámicos de $U_1, \dots, U_n$ respectivamente, el producto cartesiano de $A_1, \dots, A_n$ se denota por $A_1 \times \dots \times A_n$ y se define como un subconjunto borroso dinámico $U_1 \times \dots \times U_n$ cuya función de membresía se expresa por:

$$f_{A_1} \times \dots \times f_{A_n}(U_1, \dots, U_n) = f_{A_1}(U_1) \cap \dots \cap f_{A_n}(U_n)$$

4.4. RELACIÓN BORROSA DINÁMICA.

Una Relación Borrosa Dinámica es un subconjunto borroso dinámico de $X(t) \times X(t)$ con una función de membresía dada por:

$$f_R[x(t)]: x(t) \times x(t) \rightarrow [0,1]$$

Si se tienen dos relaciones dinámicas $R_A[x(t)]$ y $R_B[x(t)]$, se puede definir la composición dinámica de la forma:

$$R_A[x(t)] \circ R_B[x(t)]$$

Por medio de la función de membresía:

$$\max_z^t, \min^t \{f_{R_A}[x(t), z(t)] f_{R_B}[z(t), x(t)]\}$$

para alguna t y $z(t)$ y $x(t) \in X$.

4.5. PROYECCIÓN Y CONJUNTOS BORROSOS DINÁMICOS CILÍNDRICOS

Si $R(t)$ es una relación de n elementos en $U_1 \times \dots \times U_n$, luego su proyección en $U_{i1} \times \dots \times U_{ik}$ es una relación dinámica borrosa de k elementos $R_q(t)$ que está definida por:

$$R_q^T = \bigcup_{U(q')}^T f_R(U_1(t), \dots, U_n(t))$$

donde q es el índice de la secuencia, q' es el complemento de q, y $\bigcup_{U(q')}$ es el supremo de $f_R(U_1(t), \dots, U_n(t))$ sobre las u's que están en $U(q')$ en el tiempo t.

Pueden tener proyecciones idénticas en $U_{i1} \times \dots \times U_{ik}$ distintas relaciones dinámicas en $U_{i1}(t) \times \dots \times U_{ik}(t)$. Por lo que dado una relación borrosa R_q en $U_{i1}(t) \times \dots \times U_{ik}(t)$, debe de existir una sola relación dinámica suprema en un tiempo t, $R'_q(t)$ en $U_{i1}(t) \times \dots \times U_{ik}(t)$, cuya proyección dinámica en $U_{i1}(t) \times \dots \times U_{ik}(t)$, sea R_q en el tiempo t. La función de membresía $f_{R'_q}$, está dada por:

$$f_{R'q}(U_1(t) \times \cdots \times U_n(t)) = f_{Rq}(U_{i1}(t) \times \cdots \times U_{ik}(t))$$

en el entendimiento de que la ecuación es válida en un tiempo t para toda $U_{i1}(t) \times \cdots \times U_{ik}(t)$ tal que los argumentos $i_1, \dots, i_k$ en $f_{R'q}$ sean iguales en un tiempo t a los argumentos correspondientes en f_{Rq} . Por esta razón, R'_q es una extensión cilíndrica de R_q , con R_q constituyendo la base dinámica de R_q .

4.6. EL PRINCIPIO DE EXTENSIÓN

El principio de extensión para conjuntos borrosos dinámicos es en esencia una identidad dinámica básica que permite que el dominio de definición de un mapeo o una relación sean extendidos de puntos en $U(t)$ en el tiempo t a subconjuntos borrosos dinámicos de $U(t)$. Más específicamente, supóngase que f es un mapeo de $U(t)$ a $V(t)$ y $A(t)$ es un subconjunto borroso dinámico de $U(t)$ expresado como:

$$A(t) = f_1^T[x(t)]U_1(t) + \cdots + f_n^T[x(t)]U_n(t) \text{ para alguna } t$$

luego el principio de extensión nos asegura que:

$$F[A(t)]^T = f_1^T[x(t)]U_1(t) + \cdots + f_n^T[x(t)]U_n(t) = f_1^T F[U_1(t)] + \cdots + f_n^T F[U_n(t)]$$

Luego la imagen dinámica de A bajo f se puede deducir del conocimiento temporal de las imágenes dinámicas de $U_1(t) \times \cdots \times U_n(t)$ bajo f .

Si el soporte de $A(t)$ es un continuo, luego se tiene:

$$F[A(t)]^T = \int_{U(t)} f_A^T[x(t)]U(t) / U(t)$$

y el principio de extensión toma la forma:

$$F[A(t)]^T = \int_{U(t)} f_A^T[x(t)]U(t) / U(t) = \int_{V(t)} f_A^T[x(t)]U(t) / F[U(t)]$$

En este caso se puede entender que $F[U(t)]$ es un punto en $V(t)$ y $f_{A(t)}[U(t)]$ es su grado de membresía temporal en $F[A(t)]$, que es un subconjunto dinámico de $V(t)$.

4.7. CONJUNTOS DINÁMICOS DE NIVEL DE UN CONJUNTO BORROSO DINÁMICO

Si $A(t)$ es un conjunto borroso dinámico de $U(t)$, luego el conjunto dinámico de nivel α es el que comprende a todos los elementos de $U(t)$ cuyo grado de membresía en $A(t)$, en un tiempo t , es mayor o igual a α . Este conjunto se denota por $A_\alpha(t)$ y se representa por:

$$A_{\alpha} = \bigcup_{t \in T} \{U(t) | f_A(x(t)) \geq \alpha\} \text{ en un tiempo } t$$

Un conjunto borroso dinámico $A(t)$ se puede descomponer en sus conjuntos de nivel, en un tiempo t , por medio de la resolución dinámica de identidad que es:

$$A(t) = \sum_{\alpha} \alpha A_{\alpha}(t) \text{ para un } t$$

o

$$A(t) = \int_0^1 \alpha A_{\alpha}(t) \text{ para un } t.$$

en el caso continuo. $\alpha A_{\alpha}(t)$ es el producto de un escalar con un conjunto $A_{\alpha}(t)$ y la sumatoria y la integral son la unión de $A_{\alpha}(t)$ con α corriendo de 0 a 1.

Podemos usar una forma modificada del principio de extensión descomponiendo $A(t)$ en sus conjuntos dinámicos de nivel en lugar de los siguletes borrosos dinámicos. De esta forma se puede escribir:

$$F[A(t)] = F\left[\sum_{\alpha} \alpha A_{\alpha}(t)\right] = \sum_{\alpha} \alpha F[A_{\alpha}(t)]$$

Si tenemos una función F de n -arreglos, la cual es un mapeo del producto cartesiano $U_1(t) \times \dots \times U_n(t)$ al espacio $V(t)$, y un conjunto borroso dinámico $A(t)$ en $U_1(t) \times \dots \times U_n(t)$ con U_i $i = 1, \dots, n$ denotando un punto genérico en $U_i(t)$. Una aplicación directa del principio de extensión dinámico nos lleva a:

$$F[A(t)] = F_{U_1 \times \dots \times U_n} \left[\sum_{\alpha} \alpha A_{\alpha}(t) \right]$$

Por lo que en muchos casos lo que conocemos no es $A(t)$ sino sus proyecciones $A_1(t), \dots, A_n(t)$ en $U_1(t), \dots, U_n(t)$, respectivamente. En este caso la función de membresía de $A(t)$ es:

$$f_A[x(t)](U_1, \dots, U_n) = f_{A_1}[x(t)](U_1(t)) \cap \dots \cap f_{A_n}[x(t)](U_n(t))$$

donde $f_{A_i}[x(t)]$ es la función de membresía de $A_i(t)$. Esto es equivalente a asumir que $A(t)$ es el producto cartesiano de sus proyecciones dinámicas.

Consideremos que $X(t)$, $Y(t)$, y $Z(t)$ son tres conjuntos y sea $f[x(t)]$ un mapeo dinámico $f[x(t)]: X(t) \times Y(t) \rightarrow Z(t)$, el principio dinámico de extensión nos lleva a extender este mapeo de ser un punto que mapea en el espacio $X(t) \times Y(t)$ a ser un mapeo en el espacio $I(t)^{X(t)} \times I(t)^{Y(t)}$. El principio de extensión nos lleva a extender el mapeo $f[x(t)]$ a ser un subconjunto dinámico borroso con argumentos dinámicos.

4.8. CONJUNTOS DINÁMICOS BORROSOS DEL TIPO II

Un conjunto dinámico borroso del tipo II se define como un subconjunto borroso dinámico cuya función de membresía es un subconjunto borroso dinámico. Estos son de particular utilidad en situaciones en las cuales debe de existir alguna imprecisión en la descripción de la función de membresía, y además cambia en el tiempo. Consideremos que $[A(t)(x)]$ es un subconjunto dinámico borroso en el intervalo unitario. Si asumimos

$$D(t) = A(t) \overset{T}{\cap} B(t)$$

usando el principio dinámico de extensión se está en la posibilidad de obtener $D(t)$ como un subconjunto borroso dinámico del tipo II de $X(t)$. En particular:

$$D(t) = \overset{T}{\cup}_{x(t)} \left\{ \frac{A(t) \overset{T}{\cap} B(t)}{x(t)} \right\}$$

En lo anterior para cada x en un tiempo t ,

$$D[X(t)] = \overset{T}{\min}[A(t), B(t)].$$

Asumamos a $A(t)$ como un subconjunto borroso dinámico del tipo II. Definiremos la negación $A'(t)$ a ser tal que

$$A'(t) = \overset{T}{1} - A(t)$$

Si $A(t)$ en sí mismo es un subconjunto borroso dinámico, luego el empleo del principio dinámico de extensión nos da:

$$A'(t) = \overset{T}{\left\{ \frac{A_{x(t)}[Z(t)]}{1 - Z(t)} \right\}}$$

CAPÍTULO 5

5.1. LÓGICA BORROSA DINÁMICA

5.2. EL EFECTO DE LA BORROSIDAD EN HECHOS Y REGLAS.

5.1. LÓGICA BORROSA DINÁMICA

A partir de los conjuntos borrosos dinámicos, podemos de alguna manera extender la lógica borrosa, a la lógica borrosa dinámica. Tomemos la sintaxis de la lógica borrosa; "P" y "Q",.... tendrán valores en el intervalo cerrado $\{0,1\}$, y sean "P" "Q",..., significados por el valor de las variables proposicionales. La valuación de las conectivas se definen como:

$$\begin{aligned} |\neg P| &= 1 - |P| \\ |P \wedge Q| &= \min(|P|, |Q|) \\ |P \vee Q| &= \max(|P|, |Q|) \\ |P \rightarrow Q| &= 1 \text{ si } |P| \leq |Q| \end{aligned}$$

En lógica borrosa dinámica las definiremos de la siguiente forma:

$$\begin{aligned} |\neg P(t)| &= 1 - |P(t)| \\ |P(t) \wedge Q(t)| &= \min(|P(t)|, |Q(t)|) \\ |P(t) \vee Q(t)| &= \max(|P(t)|, |Q(t)|) \\ |P(t) \rightarrow Q(t)| &= 1 \text{ si } |P(t)| \leq |Q(t)| \end{aligned}$$

Se puede en lógica borrosa definir la vinculación semántica de la forma:

$$P \rightarrow Q \text{ si } P \leq Q$$

En lógica borrosa dinámica podemos definir la vinculación borrosa dinámica de la forma:

$$P(t) \xrightarrow{T} Q(t) \text{ si } P(t) \leq Q(t)$$

que significa:

Sea $P(t)$ = Juan es muy alto en el tiempo t y $Q(t)$ = Juan es alto en el tiempo t . Claramente $P(t)$ vincula semánticamente a $Q(t)$, ya que en el tiempo t el valor de muy alto es menor o igual al de alto, pero en otro tiempo esto no tiene por que ser así..

" $\rightarrow$ " es el correlativo lógico de la implicación material que será válida en lógica borrosa en todos los casos donde se tenga una implicación lógica real entre "P", "Q", esto es que:

$$P \leq Q$$

sea necesario.

$P \rightarrow P$, y $\neg P \vee P$, son tautologías en lógica borrosa. En lógica borrosa dinámica decimos que:

$$P(t) \xrightarrow{T} P(t) \forall t \text{ y } \neg P(t) \wedge P(t)$$

también son tautologías en todo tiempo.

Supongamos que:

$P = \text{esta pared es roja en el tiempo } t$.

Supongamos que la pared es roja en 0.6 de verdad en el tiempo t , luego:

"La pared es roja o no en el tiempo t "

es cierta en un 60% de verdad en ese tiempo t , de acuerdo a la semántica dada.

Similarmente:

$$P(t) \wedge \neg P(t)$$

no es una contradicción, ya que su valor de verdad es de un 40% en el tiempo t .

Desde el punto de vista de la lógica borrosa el modus ponens, no sólo es el principio de inferencia, sino que se puede generalizar de tal forma que preserve el grado de verdad:

$$\begin{array}{l} \vdash_a P \\ P \rightarrow Q \\ \hline \therefore \vdash_a Q \end{array}$$

Si consideramos a P cierto en un grado a , y sabemos que $P \rightarrow Q$, luego se puede asegurar que Q es cierto en grado a . En la lógica borrosa dinámica:

$$\begin{array}{l} \vdash_a^T P(t) \\ P(t) \xrightarrow{T} Q(t) \\ \hline \therefore \vdash_a^T Q(t) \end{array}$$

Podemos entonces ver, cuales de las tautologías clásicas son válidas en la lógica borrosa dinámica proposicional, y son las siguientes:

$$\begin{array}{l} P(t) \xrightarrow{T} P(t) \\ \left(P(t) \xrightarrow{T} \left(Q(t) \xrightarrow{T} R(t) \right) \right) \xrightarrow{T} \left(P(t) \xrightarrow{T} Q(t) \right) \xrightarrow{T} \left(P(t) \xrightarrow{T} R(t) \right) \\ \left(\neg P(t) \xrightarrow{T} \neg Q(t) \right) \xrightarrow{T} \left(Q(t) \xrightarrow{T} P(t) \right) \\ \neg \neg P(t) \xrightarrow{T} P(t) \\ \left(P(t) \xrightarrow{T} Q(t) \right) \wedge \neg Q(t) \xrightarrow{T} \neg P(t) \end{array}$$

$$\left(P(t) \overset{T}{\rightarrow} Q(t) \right) \overset{T}{\rightarrow} \left(\left(Q(t) \overset{T}{\rightarrow} R(t) \right) \overset{T}{\rightarrow} \left(P(t) \overset{T}{\rightarrow} R(t) \right) \right)$$

Leyes de Morgan.

Leyes asociativas.

Leyes distributivas.

Leyes conmutativas.

Las tautologías no válidas en la lógica borrosa dinámica proposicional son:

$$\begin{aligned} & P \overset{T}{\vee} \neg P \text{ para un } t \\ & P(t) \overset{T}{\rightarrow} \left(Q(t) \overset{T}{\rightarrow} P(t) \right) \text{ para un } t \\ & \neg P(t) \overset{T}{\rightarrow} \left(Q(t) \overset{T}{\rightarrow} P(t) \right) \text{ para un } t \\ & \left(\left(P(t) \overset{T}{\wedge} Q(t) \right) \overset{T}{\rightarrow} R(t) \right) \overset{T}{\leftrightarrow} \left(P(t) \overset{T}{\rightarrow} \left(Q(t) \overset{T}{\rightarrow} R(t) \right) \right) \text{ para un } t \\ & \left(P(t) \overset{T}{\rightarrow} \left(Q(t) \overset{T}{\wedge} \neg Q(t) \right) \right) \overset{T}{\rightarrow} \neg P(t) \text{ para un } t \\ & \left(P(t) \overset{T}{\wedge} \neg P(t) \right) \overset{T}{\rightarrow} Q(t) \text{ para un } t \\ & Q(t) \overset{T}{\rightarrow} \left(P(t) \overset{T}{\wedge} \neg P(t) \right) \text{ para un } t \end{aligned}$$

La lógica borrosa dinámica proposicional se reduce a la lógica borrosa proposicional cuando las variables proposicionales son válidas para todo tiempo.

Podemos inferir por extensión la lógica borrosa dinámica de predicados de la lógica borrosa dinámica proposicional, definiendo valuaciones para predicados dinámicos y para cuantificadores dinámicos. Sea $F(t)$ un predicado de n valores. Definimos el valor de $Fx(t)_1, \dots, Fx(t)_n$ como el grado de pertenencia de la eneada ordenada $(x_1(t), \dots, x(t)_i)$ en el conjunto borroso $F(t)_i$, donde $x_i(t)$ es la denotación de $x(t)$ en una valuación dada. En otras palabras:

$$\left| F(t)_{x_1}, \dots, F(t)_{x_n} \right| = \psi(t)(\bar{x}_1(t), \dots, \bar{x}(t)_n)$$

Los cuantificadores se valúan de la forma:

$$\forall x(t) F(t)_x$$

se define como el mínimo de los valores de $F(t)x(t)$ para todas las asignaciones de $x(t)$ a los elementos del universo del discurso.

El valor de

$$\exists x(t) F(t)_x$$

es el máximo correspondiente, o sea:

$|\forall x(t)F(t)_x|^T = \min\{|F(t)_x|\}$ para todas las asignaciones de $x(t)$ en el universo del discurso.

$|\exists x(t)F(t)_x|^T = \max\{|F(t)_x|\}$ para todas las asignaciones de $x(t)$ en el universo del discurso.

Podemos tener una lógica borrosa diámica modal, sumando los operadores " $\overset{T}{\uparrow}$ " y " $\overset{T}{\downarrow}$ ", y dándoles los valores:

$$\left| \overset{T}{\uparrow} P(t) \right| / W(t) = \min\{|P(t)| / W(t)\} \text{ para toda } W(t)' \text{ tal que } R(t)WW.$$

$$\left| \overset{T}{\downarrow} P(t) \right| / W(t) = \max\{|P(t)| / W(t)\} \text{ para toda } W(t) \text{ tal que } R(t)WW.$$

donde "/" es la relación alternativa. El signo ($\overset{T}{\uparrow}$) es el functor de necesidad y ($\overset{T}{\downarrow}$) es el functor de posibilidad.

Nótese que el valor " $\overset{T}{\uparrow} P(t)$ ", es igual a algún α , $0 < \alpha < 1$. Si se interpreta como que $P(t)$ es necesariamente cierto. En la lógica borrosa dinámica modal se deben de tener valores de grados de esta necesidad de verdad

$$\text{si } \overset{T}{\uparrow} P(t) \text{ cae entre } 0 \text{ y } 1.$$

Si deseamos hacer esta extensión guardando la evaluación:

$$(|\neg P(t)|) = 1 - |P(t)|$$

y consideramos la presuposición como:

$$P(t) \text{ presupone } Q(t), \text{ si } P(t) \parallel \neg Q(t) \text{ y } \neg P(t) \parallel \neg Q(t)$$

tomando la definición de vínculo tenemos: $P(t)$ presupone $Q(t)$, si:

$$P(t) \leq Q(t) \text{ y } 1 - P(t) \leq Q(t)$$

en todos los modelos. Pero esto significa que $P(t)$ presupone $Q(t)$ sólo si $Q(t) < 50\%$ en el tiempo t , que es un resultado extraño.

Una forma de abolir esto, es dando una valuación típica de alguna forma. En vez de asignar un valor verdadero a la proposición, se da una valuación asignando una pareja ordenada, $(V(t), F(t))$ que tiene un valor cierto y otro falso en el tiempo, cuya suma es:

$$V(t) + F(t) = 1$$

Para la lógica temporal presuposicional se puede extender esto a una tercera ordenada ($V(t)$, $F(t)$, $N(t)$) de forma que:

$$V(t) + F(t) + N(t) = 1$$

y donde $N(t)$ es que no tiene sentido. Así:

$$\neg P(t) = (\alpha, \beta, \gamma) \text{ si } P(t) = (\alpha, \beta, \gamma) \text{ donde } \alpha + \beta + \gamma = 1$$

en otras palabras, el verdadero valor de $\neg P$ será el valor falsificado de P y viceversa, y P y $\neg\neg P$, tendrán el mismo valor de N .

Se ha definido " $\rightarrow$ " a tener valores 0 y 1. Si suponemos que toma valores intermedios, se tendrá que generalizar la noción de vínculo temporal a la de vínculo temporal en un grado α . Se escribe :

$$\left\| \begin{array}{c} t \\ - \\ \alpha \end{array} \right\|$$

y se tiene:

$$P(t) \left\| \begin{array}{c} t \\ - \\ \alpha \end{array} \right\| Q(t) \text{ si } \left\| \begin{array}{c} t \\ - \\ \alpha \end{array} \right\| P(t) \rightarrow Q(t)$$

Veremos otras ideas sobre lógica borrosa dinámica como las proposiciones condicionales borrosas dinámicas, y las reglas composicionales dinámicas de inferencia.

En cálculo proposicional clásico consideremos la expresión si A luego B :

$$A \rightarrow B$$

se tiene, como ya se vio que:

$$A \rightarrow B = \neg A \vee B$$

Un condicional borroso dinámico $A(t) \xrightarrow{T} B(t)$, con $A(t)$ como antecedente y $B(t)$ como consecuente, son conjuntos borrosos en lugar de variedades proposicionales.

Ejemplos de esto son:

Si grande luego pequeño en el tiempo t .

Si resbaloso luego peligroso en tiempo de lluvia.

que son abreviaciones de:

Si $x(t)$ es grande, luego $y(t)$ es pequeño en el tiempo t .

Si la carretera está resbalosa en tiempo de lluvia, luego es peligroso manejar.

En esencia estas proposiciones describen relaciones entre dos variables borrosas dinámicas. Esto da a entender que las proposiciones condicionales borrosas dinámicas están definidas como relaciones borrosas, y no como conectivos lógicos. Para esto se define el producto cartesiano de dos conjuntos borrosos dinámicos como:

"Sea $A(t)$ un subconjunto borroso del universo $U(t)$, y sea $B(t)$ un subconjunto borroso de otro diferente universo del discurso, luego el producto cartesiano de $A(t)$ y $B(t)$ se define por:

$$A(t) \times^T B(t) = \int_{U \times V} \psi_{A(t)}(U(t)) \wedge^T \psi_{B(t)}(V(t)) / (U(t), V(t))$$

El significado de una condicional borrosa de la forma $A(t)$ luego $B(t)$, se aclara considerando que es un caso particular de si $A(t)$ luego $B(t)$, de otra forma $C(t)$ en un tiempo t . Donde $A(t)$, $B(t)$, $C(t)$, son subconjuntos borrosos de los universos del discurso $U(t)$, $V(t)$. Esto en términos de lo anterior se define como:

Si $A(t)$ luego de otra forma:

$$C(t) = A(t) \times^T B(t) + \left(\neg^T A(t) \times^T B(t) \right)$$

En el caso de que

$$C(t) = U(t), A(t) \rightarrow^T B(t) = A(t) \times^T B(t) + \left(\neg^T A(t) \times^T V(t) \right)$$

Si

$$a(t) = U(t), A(t) \rightarrow^T B(t) = U(t) \times^T B(t) = \left(\neg^T A(t) \times^T V(t) \right)$$

La consideración de $C(t) = V(t)$ implica que en ausencia de una indicación de lo contrario, el consecuente de:

$$\neg^T A(t) \rightarrow^T C(t)$$

puede ser cualquier subconjunto borroso dinámico del universo del discurso.

La regla de inferencia borrosa dinámica la definiremos así:

Si $R(t)$ es una relación borrosa de $U(t)$ a $U(t)$, y $X(t)$ es un subconjunto borroso de $U(t)$, luego el subconjunto borroso $Y(t)$ de $U(t)$, que es inducido por $X(t)$ está dado por la composición:

$$Y(t) = X(t) \circ^T R(t)$$

en la cual:

$$X \circ^T R = \int_{U(t) \times V(t)} \psi_{Y(t)} \left[\left(\psi(t) \times^T (U(t), V(t)) \wedge^T \psi(t) R(t) (U(t), V(t)) \right) \right] / (U(t), V(t))$$

$y \times^T$ juega el papel de una relación unitaria. Esto lo podemos ver como modus ponens generalizado.

Definiremos una variable lingüística dinámica como:

Una variable lingüística dinámica se entiende como una variable cuyos valores son palabras o proposiciones en un lenguaje natural o artificial cuyo significado cambia en el tiempo. Por ejemplo EDAD, es una variable lingüística si sus valores son lingüísticos y no numéricos, por ejemplo: Joven, muy joven, etc., que va cambiando con el tiempo

En términos más formales:

Una variable lingüística dinámica se caracteriza por un quintuplete $(X(t), T_v(X), U(t), G(t), M(t))$ en el cual:

- $X(t)$ es el nombre de la variable.
- $T(t) (X)$ Es el conjunto de términos de X , o sea la colección de sus valores lingüísticos.
- $U(t)$ Es el universo del discurso.
- $G(t)$ Es la regla sintáctica que genera los términos en $T(X)$.
- $M(t)$ Es la regla semántica que asocia a cada valor lingüístico x , su significado $M(X)$, donde M es un subconjunto borroso de U .

Clásicamente se define el universo del discurso como:

Como todo nombre connotativo debe denotar, pero hay nombres descriptivos que no se supone sean aplicables a algo en el mundo real o verdadero, pero que tienen aplicación en alguna esfera o universo que se debe de distinguir del universo real.

De aquí surge la idea del universo del discurso introducida por De Morgan que lo explica así:

"Si recordamos que en muchas proposiciones la extensión que abarca el pensamiento es más reducida que el universo entero, se empieza a descubrir que toda la extensión que abarca un tema de discusión es para el propósito de ésta, la cual llamaremos un universo, o sea una extensión de ideas que es expresado o comprendido como contentiva de todo el asunto que se discute".

O sea que De Morgan restringe su aplicación a una extensión limitada por la significación de los nombres positivos correspondientes.

Definiremos el universo dinámico del discurso como un conjunto de objetos, acciones, relaciones o conceptos, que forman el sujeto del discurso que cambia en el tiempo.

El significado de un valor lingüístico $x(t)$ se caracteriza por una función de compatibilidad $C(t): \rightarrow [0,1]$, que asocia con cada $u(t)$ en U su compatibilidad con $x(t)$. La función de la regla semántica es para relacionar la compatibilidad de los términos primarios en un valor lingüístico compuesto que cambia en el tiempo, a la compatibilidad del valor compuesto. Para esto las limitantes lógicas dinámicas, como las conectivas (y, o) las trataremos como operadores no lineales que modifican en el tiempo el significado de sus operandos en una forma definida. Consideraremos los valores lingüísticos como expresiones de la forma: $h(t)U$, donde $h(t)$ es una limitante lingüística dinámica, y U es un término primario. En unión con los conectivos lógicos, (y, o) y las limitantes lingüísticas dinámicas, los términos primarios sirven como los generadores de los valores de una variable lingüística dinámica, constituyendo el

sentido de cada valor en el tiempo, un subconjunto borroso dinámico de U. Por ejemplo: el valor de verdad lingüístico dinámico cierto, se puede representar como el subconjunto borroso dinámico:

$$\text{Cierto} = \int_0^T \psi(t) \text{cierto}(V(t)) / V(t)$$

donde la integral indica la unión de elementos únicos :

$$y(t) \text{ cierto } (v(t))/v(t)$$

oscilando v en el intervalo [0,1].

Si cierto se define de esa forma, no cierto se definirá como:

$$\text{no cierto} = \int_0^T (1 - \psi(t) \text{cierto}(V(t))) / V(t)$$

y falso:

$$\text{Falso} = \int_0^T \psi(t) \text{cierto}(1 - V(t)) / V(t)$$

Con estos elementos se crea lo que llamaremos razonamiento aproximado dinámico. Discutiremos más a fondo la idea de la presentación de la lógica borrosa dinámica. *"La lógica borrosa dinámica es diferente de la lógica bivaluada, y de la lógica borrosa. Esta es la base del razonamiento aproximado dinámico, que es un modo de razonar nuevo, en el cual los valores de verdad y las reglas de inferencia son borrosos, y cambian en el tiempo"*.

Cambiaremos el nombre de proposición al de declaración dinámica, y son de la forma:

U(t) es de A(t) en un tiempo t.

donde U(t) es el nombre de un objeto, y A(t) es el nombre de un posible subconjunto borroso dinámico del universo del discurso U. Si interpretamos a A(t) como un predicado borroso, luego la declaración: U(t) es A(t), en el tiempo t, se puede parafrasear como: U tiene la propiedad A en el tiempo t.

Equivalentemente "U(t)" es "A(t)", se puede interpretar como una ecuación de asignación en la cual, un conjunto borroso dinámico A(t) se asigna como un valor a una variable lingüística dinámica que denota un atributo de U(t) en el tiempo t. Esta proposición está asociada con dos subconjuntos borrosos dinámicos:

1. El significado de A(t), M(t)(A), que es un subconjunto borroso de U(t).
2. El valor de verdad de "U(t)" es "A(t)", denotado por V(A(t)), y se define como un posible subconjunto de un universo de valores de verdad.

Decimos que verdad es el nombre de una variable lingüística dinámica en la cual el primer término es cierto, con falso como la negación de cierto, así:

$$\psi(t) \text{ falso}(V(t)) = \psi(t) \text{ cierto}(1 - V(t)) \text{ con } 0 \leq V(t) \leq 1$$

Para construir la base de una lógica borrosa dinámica, se extiende el significado de las operaciones lógicas $\neg, \vee, \wedge, \rightarrow$, a operadores con valores lingüísticos dinámicos en vez de numéricos. Si $A(t)$ es un subconjunto borroso del universo del discurso $U(t)$, $A(t)$, $yU(t)$, luego se tiene que:

1.) El grado de pertenencia de $U(t)$ en el conjunto borroso dinámico $A(t)$, es $y(U)$.

2.) El valor de verdad del predicando borroso $A(t)$, es $yA(t)(U)$.

proposiciones que son equivalentes.

Si $V_T(A(t))$ es un punto en $V_T = [0,1]$, representando el valor de verdad de la proposición " $U(t)$ es $A(t)$ ", donde $u(t)$ es un elemento de $U(t)$, el valor de verdad de $A(t)$ no se da por:

$$V_T(\neg A(t)) = 1 - V_T(A)$$

para los conectivos borrosos. Sea $V_T(A(t))$ y $V_T(B(t))$ los valores de verdad lingüísticos de las proposiciones $A_T(t)$ y $B_T(t)$. Entonces:

$$V_T(A(t)) \overset{T}{\wedge} V_T(B(t)) \text{ para } V_T(A(t) \text{ y } B(t))$$

$$V_T(A(t)) \overset{T}{\vee} V_T(B(t)) \text{ para } V_T(A \text{ o } B)$$

$$V_T(A(t)) \overset{T}{\rightarrow} V_T(B(t)) \text{ para } V_T(A(t) \rightarrow B(t))$$

y

$$\neg V_T(A(t)) \text{ para } V_T(\neg A(t))$$

donde $\overset{T}{\wedge}$ es el min en un tiempo t (conjunción) y $\overset{T}{\vee}$ es el max en un tiempo t (disyunción).

El valor de verdad de $A \overset{T}{\rightarrow} B$, depende de la forma en la cual el conectivo " $\overset{T}{\rightarrow}$ " se define para valores numéricos de verdad. Si definimos:

$$V_T(A(t) \overset{T}{\rightarrow} B(t)) = \neg V_T(A(t)) \overset{T}{\vee} V_T(A(t)) \overset{T}{\wedge} V_T(B(t))$$

para el caso donde $V_T(A(t))$ y $V_T(B(t))$ son puntos en $[0,1]$, la aplicación del principio dinámico de extensión da:

$$V_T(A(t) \rightarrow B(t)) = \sum_{i,j} \left(x_i(t) \wedge b_j(t) \right) / \left((1 - v_i(t)) \vee (v_i(t) \wedge w_j(t)) \right)$$

Se define el valor indeterminado de verdad, o $q(t) = ?$, como:

$$\theta(t) = \int_0^1 \theta(t) / V(t)$$

y:

$$? = V(t) = \text{universo de valores de verdad} = [0,1] = \int_0^1 1 / W(t)$$

que se interpretan como el grado de pertenencia de desconocimiento.

Para cada conjunto el grado de pertenencia de un punto u en U puede tener una de estas tres formas posibles:

- 1.) Un número en $[0,1]$
- 2.) 0, indefinido.
- 3.) ?, desconocido.

Si se sabe como calcular los valores de verdad de $A(t)$ y $B(t)$, $A(t) \vee B(t)$, y no $B(t)$, dado el valor lingüístico de verdad de $A(t)$ y $B(t)$, es simple calcular

$$V_T(A(t) \wedge B(t)) \text{ y } V_T(A(t) \vee B(t)) \text{ y } V_T(\neg B),$$

cuando $V_T(B(t)) = ?$. Supóngase que:

$$V_T(A(t)) = \int_0^1 \psi(t)(u) / u$$

y:

$$V_T(B(t)) = ? \int_0^1 1 / W$$

aplicando el principio de extensión:

$$V_T(A(t)) \wedge ? = \int_0^1 \psi(t)(u) / U(t) \wedge \int_0^1 1 / W = \int_0^1 \int_0^1 \psi(t)(U(t)) / (U(t) \wedge W(t))$$

donde:

$$\int_0^1 \int_0^1 = \int_{[0,1] \times [0,1]}$$

y esto se reduce a:

$$V_T(A) \wedge^T ? = \int_0^1 (V(t)_{[w,1]} \psi(t)(V(t)) / W)$$

o sea que el valor de verdad de A(t) y B(t) cuando $V_T(B(t)) = ?$, es un subconjunto borroso de [0,1], en el cual el grado de pertenencia de un punto w(t), está dado por el supremo de (V_T) , sobre $[w(t),1]$.

Si el universo de valores de verdad en cuatro valores q(t), T(t), F(t) y $T(t)+F(t) = ?$, se encuentra que:

$$\begin{aligned} T(t) \wedge^T 0 &= 0 \\ T(t) \wedge^T (T(t) + F(t)) &= T(t) \wedge^T T(t) + F(t) \wedge^T F(t) = T(t) + F(t) \\ F(t) \wedge^T (T(t) + F(t)) &= F(t) \wedge^T T(t) + F(t) \wedge^T F(t) = F(t) + F(t) = F(t) \\ (T(t) + F(t)) \wedge^T (T(t) + F(t)) &= T(t) \wedge^T T(t) + T(t) \wedge^T F(t) + F(t) \wedge^T T(t) + F(t) \wedge^T F(t) \\ &= T(t) + F(t) + F(t) + F(t) = T(t) + F(t) \end{aligned}$$

define el concepto de una variable de verdad compuesta como:

$$F(t) = (F(t)_1, \dots, F(t)_n)$$

denotando una eneada de valores lingüísticos dinámicos de verdad compuestos, en la cual cada $F_i(t)$ es un valor único de verdad lingüístico asociado con un conjunto de términos $T_i(t)$, un universo del discurso V_i , y una base de variables V_i .

De acuerdo a todo lo que se ha visto, la regla básica de inferencia en la lógica tradicional es el modus ponens ya explicado aquí. De acuerdo a esto, determinaremos que por el Modus ponens se puede inferir la verdad a una proposición borrosa dinámica B(t) de la verdad de otra A(t), por la implicación:

$$A(t) \rightarrow^T B(t).$$

En el razonamiento humano el modus ponens se emplea más en forma aproximada en vez de la forma exacta. Típicamente sabemos que si A(t) es verdadera, y que $A(t)^* \rightarrow B(t)$, donde $A^*(t)$ es en algún sentido una aproximación a B(t), luego de A(t) y $A(t)^* \rightarrow B(t)$, se puede inferir que B(t) es aproximadamente cierto.

En base a esto se puede definir una regla composicional de inferencia que es sólo una generalización de lo siguiente: Supóngase que se tiene la curva $Y = f(x)$, y esta dado que $x = a$, luego de $y = f(x)$ y $x = a$, se puede inferir $y = b = f(a)$.

Si generalizamos este proceso considerando que a está en un intervalo, y que f(x) es una función valuada en el intervalo. Entonces para encontrar el intervalo $y = b$, que corresponda al intervalo a, primero se construye un conjunto cilíndrico a con base a, y

se encuentra su intersección I con la curva valuada en el intervalo. Luego se proyecta la intersección en el eje x, encontrando la y deseada en el intervalo b.

Más específicamente sean $y_A(t)$, $y_B(t)$, $y_F(t)$ las funciones de pertenencia de $A(t)$, $B(t)$, $F(t)$, luego por definición de $A(t)$:

$$\begin{aligned}\psi(t)_A(U_1, \dots, U_n) \\ \psi(t)_A(U_{i_1}, \dots, U_{i_k}) \\ \psi(t)_A(x(t), y(t)) = \psi(t)_A(x(t))\end{aligned}$$

y en consecuencia:

$$\begin{aligned}\psi_{A+F}(x, y) &= \psi_A(x, y) \wedge \psi_F(x, y) \\ \psi_A(x) &\wedge F\psi(x, y)\end{aligned}$$

Proyectando $A(t)$, $F(t)$ en el eje OY se tiene:

$$\psi_B(t)(y(t)) = V_{T_x} \wedge \psi_{T_A}(x(t)) \wedge \psi(t)(x(t), y(t))$$

con esto se ve que $B(t)$ se puede representar como: $B(t) = A(t) \circ B(t)$, que como ya se dijo antes es la composición. Proponemos la siguiente regla:

Sean U y V dos universos del discurso con bases variables $u(t)$ y $v(t)$ respectivamente. Sea $R(u(t))$, $R(u(t).v(t))$ y $R(v(t))$ restricciones en $U(u(t), v(t))$ y $v(t)$, en el entendido de que $R(u(t))$, $R(u(t), v(t))$ y $R(v(t))$ son borrosas dinámicas como relaciones en U , $U \times V$, y V . Sean $A(t)$ y $F(t)$ subconjuntos borrosos particulares de U y $U \times V$, luego la regla composicional de inferencia dice que la solución de la ecuación relacional de asignación:

$$R(t)(U) = A(t)$$

y

$$R(t)(U, V) = F(t)$$

están dadas por:

$$R(t)(V) = A(t) \circ F(t)$$

En este sentido se puede inferir:

$$R(t)(v) = A(t) \circ F(t) \text{ de } R(t)(u) = A(t) \text{ y } R(t)(u, v) = F(t)$$

De esta manera se ve el Modus Ponens como un caso especial de la regla composicional de inferencia.

En lógica tradicional, la implicación material se define como un conectivo lógico para variables proposicionales.

Si $A(t)$ y $B(t)$ son dos variables proposicionales, luego $A(t) \xrightarrow{T} B(t)$, o si $A(t)$ luego $B(t)$. En muchos discursos humanos se tiene que $A(t)$ y $B(t)$ son conjuntos borrosos dinámicos o predicados borrosos dinámicos, en vez de variables proposicionales. Para extender la noción de implicación material a conjuntos borrosos dinámicos, sean U y V dos posibles universos del discurso, y sea $A(t), B(t)$ y $C(t)$ los conjuntos borrosos de U, V y V . Primero se definirá el significado de la expresión si $A(t)$ luego $B(t)$ de lo contrario $C(t)$, y luego se definirá si $A(t)$ entonces $B(t)$ como un caso particular del anterior.

La expresión si $A(t)$ entonces $B(t)$ de lo contrario $C(t)$, es una relación dinámica binaria borrosa definida en $U \times V$ por :

Si $A(t)$ luego $B(t)$ de lo contrario

$$C = A(t) \times B(t) + \neg A(t) \times C(t)$$

Para el caso si $A(t)$ luego $B(t)$ se considera que $C(t) = V$ y:

$$\text{Si } A(t) \text{ luego } B(t) = A(t) \times B(t) + \neg A(t) \times B(t)$$

En base a lo anterior se puede definir el modus ponens de la forma:

"Sean $A_1(t)$, $A_2(t)$ y $B(t)$ los subconjuntos borrosos de U , $U \times V$ respectivamente. Si $A_1(t)$ se asigna a la restricción $R(t)(u)$ y la relación

$$A_2(t) \xrightarrow{T} B(t)$$

se asigna a la restricción $R(t)(u(t), v(t))$. Luego:

$$\begin{aligned} R(t)(u(t)) &= A_1(t) \\ R(t)(u(t), v(t)) &= A_2(t) \xrightarrow{T} B(t) \end{aligned}$$

Como se ve esta ecuación relacional de asignación se puede resolver para la restricción en $v(t)$, obteniéndose:

$$R(t)(v(t)) = A_1(t) \text{ o } (A_2(t) \xrightarrow{T} B(t))$$

Una expresión para esta conclusión en la forma:

$A_1(t)$	Premisa
$A_2(t) \xrightarrow{T} B(t)$	Implicación.
$A_1(t) \text{ o } (A_2(t) \xrightarrow{T} B(t))$	Conclusión.

constituye el establecimiento de un Modus Ponens generalizado. Esto difiere del modus ponens tradicional en dos aspectos:

1. A1, A2, y B, son subconjuntos borrosos.
2. A1 no debe de ser idéntico con A2.

Con estos desarrollos de la lógica borrosa se define al razonamiento aproximado dinámico, como una forma de razonamiento que implica los elementos fundamentales siguientes:

1. El empleo de valores de verdad lingüísticos dependientes del tiempo para expresar el grado de verdad de aseveraciones borrosas y sus combinaciones.
2. La representación de aseveraciones borrosas como ecuaciones relacionales dinámicas.
3. La solución de ecuaciones de asignación relacionales dinámicas, para las restricciones borrosas sobre variables específicas dependientes del tiempo.
4. Aproximación a la solución de ecuaciones de asignación relacionales por valores lingüísticos dinámicos.
5. El empleo del Modus Ponens generalizado dinámico para deducir conclusiones aproximadas de aseveraciones borrosas dinámicas:

$$(U_{1(t)}, \dots, U_{n(t)}) \text{ es } A(t)$$

y enunciados convencionales de la forma:

Si B(t) entonces C(t) de lo contrario D(t).

Una definición más formal del universo del discurso será de la forma siguiente: Primero se define lo que llama un espacio esencial dinámico, y lo definiremos como:

Un espacio esencial $K = \{W\}^T$, con elementos genéricos denotados por $w(t)$, puede ser cualquier conjunto de objetos y construcciones, así que a partir de esto se define el Universo del discurso de la forma:

"Sea R una colección dada de objetos que forman el espacio esencial dinámico. Sea E un conjunto que contiene K y que está generado de K para la aplicación finita de las operaciones, unión, producto directo, y F colección de conjuntos borrosos dinámicos, luego el universo del discurso $U(K)$, o simplemente la U, es un subconjunto de E".

El empleo de la lógica borrosa dinámica, como un marco de trabajo para el manejo de conceptos dinámicos inciertos, nos permite considerar una serie de elementos que no se podrían operar con efectividad o correctamente por las técnicas convencionales. Los puntos más importantes de estos elementos son:

La borrosidad de los antecedentes y los consecuentes en reglas de la forma:

1. Si $x(t)$ es A, luego $y(t)$ es B
2. Si $x(t)$ es A, luego y es B, con factor de certeza $= a^T$

donde el antecedente de $x(t)$ es A , y el consecuente $x(t)$ es B , son proposiciones borrosas dinámicas, y a es un valor numérico del factor de certeza, por ejemplo:

Si $x(t)$ es pequeño, luego $y(t)$ es largo, con $FC = 0.8$

en el cual el antecedente $x(t)$ es pequeño, y el consecuente $y(t)$ es largo, son proposiciones borrosas dinámicas, ya que la denotación de los predicados, pequeño y largo, son subconjuntos borrosos *pequeño* y *largo* dinámicos.

Relación parcial entre el antecedente de una regla y el hecho aportado por el propositor.

El número de reglas que se operan es pequeño, y habrá muchos casos en los que un hecho tal como $x(t)$ es $A(t)$, no case exactamente con el antecedente de ninguna regla de la forma:

Si $x(t)$ es $A(t)$, luego $y(t)$ es $B(t)$ con $FC(t) = a$.

La regla convencional no puede encajar en forma ad hoc, ya que la relación parcial no nos lleva por sí misma a un análisis dentro de los confines de una lógica bivaluada, por el contrario los grados de verdad y pertenencia en lógica borrosa dinámica nos dan una forma natural de trabajar con relaciones parciales en el tiempo a través del empleo de reglas dinámicas de composición de inferencia e interpolación.

La presencia de cuantificadores borrosos dinámicos en el antecedente y/o consecuente de una regla, tales como: varios, pocos, muchos más, usualmente, etc. Como una ilustración considérese una disposición dinámica como la siguiente: (disposición es una proposición que implica cuantificadores borrosos dinámicos).

d D(t) Los estudiantes en primaria son niños.

que se puede interpretar como la proposición:

p D(t) Muchos estudiantes en primaria son niños.

que se puede expresar como una regla, o equivalentemente como la proposición condicional:

r D(t) Si $x(t)$ es un estudiante, es posible que sea joven.

en la cual, la posibilidad de posible, tiene la misma denotación expresada como un subconjunto borroso dinámico del intervalo unitario que el cuantificador muchos.

5.2. EL EFECTO DE LA BORROSIDAD EN HECHOS Y REGLAS.

Típicamente la base de conocimientos consiste en una colección de proposiciones que representan los hechos, y una colección de proposiciones condicionales que constituyen las reglas, por ejemplo los hechos pueden ser:

Carolina en este momento es una estudiante de Posgrado.
 La población del D.F.hoy es de alrededor de 20,000,000
 El D.F. es una ciudad contaminada en 1999 .
 Juan tiene ahora una úlcera duodenal. $FC = 0.3$

Si poder tener una úlcera duodenal es un predicado borroso dinámico, luego que Juan la pueda tener en un grado, hace que el significado de factor de certeza sea indefinido, más específicamente, que el

$$FC = 0.3,$$

significa que Juan hoy tiene una úlcera duodenal en un grado del 0.3, o sea que la posibilidad del evento borroso: "Juan hoy tiene una úlcera duodenal" es de 0.3. Como se ve, para interpretar significativamente es necesario poder definir la posibilidad de un evento borroso dinámico. Esto se puede hacer en lógica borrosa dinámica.

Como otro ejemplo considérese una regla de la forma general:

Si $x(t)$ es $A(t)$, luego $y(t)$ es $B(t)$, con posibilidad $p(t)$.

donde $x(t)$ y $y(t)$ son variables dinámicas, $A(t)$ y $B(t)$ son predicados dinámicos, y $p(t)$ una posibilidad borrosa dinámica expresada como un número borroso y alrededor de 0.8, o como una posibilidad lingüística: muy posible. Por ejemplo:

Si Juan tiene hoy un Ferrari rojo nuevo, luego es probable que su esposa sea joven.

Expresada como una posibilidad condicional, la regla se puede escribir como:

$$Po\{y(t) \text{ es } B(x(t) \text{ es } A(t))\} \text{ es } p(t)$$

Si la regla se trata como una probabilidad condicional, se sigue que:

$$Pr(t)\{y \text{ no es } B(t)/x(t) \text{ es } A(t)\} \text{ es } 1 - p(t)$$

pero esta condición es incorrecta si $A(t)$ es un conjunto borroso. La conclusión correcta será:

$$P(t)r\{y \text{ no es } B(t)/x(t) \text{ es } A(t)\} + Pr\{y(t) \text{ es } B(t)/x(t) \text{ es } A(t)\} \geq 1$$

donde las probabilidades en cuestión son números borrosos. En concreto:

$$P(H / E) \equiv 1 - P(\neg H / E)$$

no es válida cuando E es una proposición borrosa dinámica.

CAPÍTULO 6

- 6.1. INFERENCIA EN LOGICA BORROSA DINAMICA.
- 6.2. TIPOS DE PROPOSICIONES.
- 6.3. DEDUCCION.
- 6.4. INFERENCIA DE PROPOSICIONES CUANTIFICADAS.
- 6.5. CARDINALIDAD DE UN CONJUNTO BORROSO.
- 6.6. EL SILOGISMO INTERSECCION/PRODUCTO
- 6.8. CADENAS DE PROPOSICIONES.
- 6.9. SEMANTICA Y PRAGMATICA DE LOS SISTEMAS BORROSOS.
- 6.10. RESTRICCIONES BORROSAS DINÁMICAS
- 6.11. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO I
- 6.12. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO II
- 6.13. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO III
- 6.14. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO IV
- 6.15. RAZONAMIENTO BORROSO DINÁMICO
- 6.16. EL CONCEPTO DE DISTRIBUCIÓN DINÁMICA DE POSIBILIDADES
- 6.17. ECUACIONES DE ASIGNACIÓN DE POSIBILIDAD DINÁMICA
- 6.18. PROYECCIÓN Y PARTICULARIZACIÓN
- 6.19. EL CONCEPTO DE UNA VARIABLE LINGÜÍSTICA DINÁMICA

6.1. INFERENCIA EN LOGICA BORROSA DINAMICA.

El mayor poder expresivo de la lógica borrosa dinámica, deriva de que esta contiene como casos especiales a la lógica borrosa tradicional, y la lógica clásica de dos valores, así como las lógicas multivaluadas. Las principales características de la lógica borrosa dinámica son las siguientes:

1. En la lógica bivaluada una proposición P , es siempre falsa o cierta. En los sistemas lógicos multivaluados, una proposición puede ser cierta o tener un valor intermedio de verdad que puede ser un elemento de un conjunto finito de valores de verdad T . En lógica borrosa dinámica, los valores de verdad se consideran dentro de un intervalo sobre el subconjunto borroso dinámico de T . Por ejemplo, si T es el intervalo unitario, luego el valor de verdad en lógica borrosa dinámica, por ejemplo muy cierto, se puede interpretar como un subconjunto borroso dinámico del intervalo unitario que define la posibilidad dinámica de distribución asociada con el valor de verdad en cuestión. En este sentido un valor de verdad borroso dinámico puede ser visto como una caracterización imprecisa de un valor de verdad intermedio que cambia con el tiempo.
2. Los predicados en lógica bivaluada están restringidos a ser decisivos en el sentido de que la denotación de un predicado debe de ser un conjunto no borroso del universo del discurso. En lógica borrosa dinámica, los predicados pueden ser decididos, por ejemplo: Mortal, siempre, padre de, o más, generalmente, etc. o borrosos dinámicos como por ejemplo: enfermo, cansado, largo, alto, etc.
3. La lógica bivaluada, así como la multivaluada, consideran sólo dos cuantificadores, todo y alguno, en contraste la lógica borrosa dinámica considera el empleo de cuantificadores borrosos dinámicos como: mucho, varios, pocos, mucho de, frecuentemente, etc. Estos cuantificadores se pueden interpretar como números borrosos que nos dan una caracterización imprecisa de la cardinalidad de uno o más conjuntos borrosos dinámicos. Desde esta perspectiva, un cuantificador borroso dinámico se puede ver como un predicado borroso dinámico de segundo orden. Basados en este punto de vista, los cuantificadores borrosos dinámicos se pueden usar para representar el significado de proposiciones que contengan probabilidades borrosas dinámicas, y permiten el manejo de probabilidades dentro de la lógica borrosa dinámica.
4. La lógica borrosa dinámica nos da un método para representar el significado de los modificadores tanto borrosos dinámicos como borrosos y no borrosos, como no, mucho, más o menos, extremadamente, etc. Esto nos permite el empleo de variables lingüísticas dinámicas, o sea, variables cuyos valores son palabras o sentencias en el lenguaje natural, cuyo significado cambia en el tiempo. Por ejemplo: Edad es una variable lingüística, cuyos valores pueden ser: Joven, viejo, muy joven, no muy viejo, etc., que cambia con el tiempo.
5. En lógica bivaluada, una proposición p , se puede cualificar principalmente por la asociación con p , de un valor de verdad cierto, o falso. Un operador modal tal como, posible o necesario, y un operador intencional tal como, creíble. En lógica borrosa dinámica, los tres modos principales de cualificación son:

1. Cualificación de verdad, expresada como $p(t)$ es $u(t)$, en la cual $u(t)$ es un valor de verdad borroso dinámico.
2. Cualificación de probabilidad, expresada como $p(t)$ es l , en la cual l es una probabilidad borrosa dinámica.
3. Cualificación de posibilidad, expresada como $p(t)$ es r , donde r es una posibilidad borrosa.
4. por lo que conocer y creer, se consideran como predicados borrosos dinámicos.

6.2. TIPOS DE PROPOSICIONES.

En la lógica bivaluada, la principal herramienta para inferencia es el Modus Ponens. La validez de este proceso de inferencia se puede poner en duda cuando la información en los conocimientos es borrosa dinámica. Las ideas básicas para el proceso de deducción en lógica borrosa dinámica son:

Si consideramos que la base de conocimientos consiste en una colección de proposiciones $\{p_1(t), \dots, p_n(t)\}$, de las cuales algunas o todas son de naturaleza borrosa dinámica, las proposiciones se pueden dividir en cuatro categorías principales:

1. La proposición incondicional y no cualificada. Ej: Juan tiene un hermano joven, toma la forma canónica:

$$\text{es } Fx(t)$$

donde $x(t)$ es una variable, y F es un predicado borroso dinámico.

2. La proposición incondicional cualificada, por ejemplo: Es muy probable que Juan tenga una hermana joven, lo que muestra en forma más directa la proposición. María es muy vivaz la mayoría del tiempo. En este caso el cualificador mayoría del tiempo, juega el papel de un cualificador borroso dinámico.

Formas canónicas:

1. $x(t)$ es $F(t)$ es l , donde $x(t)$ es una variable, $F(t)$ es un predicado borroso, y l es una probabilidad borrosa.
2. $Q's(t)$ son $F's(t)$, donde $q(t)$ es un cuantificador borroso y $F(t)$ es un subconjunto borroso del universo del discurso U .
3. La proposición condicional cualificada, por ejemplo: Si un carro es viejo, es muy probable que no sea rentable. Sus formas canónicas son dos:
 - a) Si $x(t)$ es $F(t)$, luego $y(t)$ es $G(t)$ es l , donde $x(t)$, $y(t)$ son variables, y $F(t)$ y $G(t)$ son predicados borrosos dinámicos.
 - b) $QF's(t)$ son $G's(t)$, donde $Q(t)$ es un cuantificador borroso, y $F(t)$ y $G(t)$, son predicados borrosos dinámicos.

4.) La proposición condicional no cualificada, por ejemplo: Si $x(t)$ es un hombre, $x(t)$ es mortal, cuya forma canónica es: Si $x(t)$ es $F(t)$, luego $y(t)$ es $G(t)$, donde $x(t)$, $y(t)$, son variables, y $F(t)$ y $G(t)$ son predicados borrosos dinámicos.

Muchas de las proposiciones del tipo (3), son disposiciones, es decir proposiciones que implícitamente contienen cuantificadores borrosos dinámicos. Las disposiciones juegan un papel de especial importancia en la representación del conocimiento común.

Una proposición p en un lenguaje natural se puede ver como una colección de variables focales $x_1(t), \dots, x_n(t)$, tomando valores en $U_1, \dots, U_n$ respectivamente, las cuales están restringidas por un sistema de constricciones borrosas $F_1(t), \dots, F_m(t)$. en lo general, las variables $x_1(t), \dots, x_n(t)$ y sus constricciones asociadas $F_1(t), \dots, F_m(t)$, son implícitas en vez de explícitas en $p(t)$. Siendo más concretos, cualquier proposición incondicional $p(t)$, se puede representar en la forma canónica $cf(p(t))$ como:

$$p(t) \xrightarrow{T} cf(p(t)) \text{ D } x(t) \text{ es } F(t)$$

en la cual:

$$x(t) \text{ D } (x_1(t), \dots, x_n(t))$$

es una enxada de variables focales cuyos componentes son: $x_1(t), \dots, x_n(t)$; y $F(t)$ es una relación borrosa en $U_1, \dots, U_n$, la cual representa una restricción elástica en $x(t)$. Informalmente lo que esto significa, es que la proposición $p(t)$ se puede ver como un sistema de constricciones elásticas, y que la representación del significado de $p(t)$, es el proceso por el cual las constricciones implícitas en $p(t)$ se hacen explícitas, expresando $p(t)$ en una forma que ponga en evidencia la variable restringida $x(t)$, y la restricción borrosa dinámica $F(t)$ que es inducida por $p(t)$.

Sea $Px(t)$ la posibilidad dinámica de distribución de $x(t)$, esto es, el conjunto borroso dinámico de posibles valores que $x(t)$ puede tomar en :

$$\boxed{}$$

luego la proposición canónica $x(t)$ es F se puede interpretar como la ecuación de posibilidad de asignación:

$$Px(t) = F(t)$$

o más explícitamente como:

$$\Pi_x(U) \Delta \text{Pos} \left\{ x(t) = U \right\}^T = \mu_{F(t)}(U)$$

donde $U = (U_1, \dots, U_n)$ es un punto genérico en $U = U_1 \times \dots \times U_n$, con la función:

$$\Pi_{x(t)}: U \xrightarrow{T} [0,1]$$

como la función de posibilidad de distribución asociada con Px , y:

$$\mu_{x(t)}: U \xrightarrow{T} [0,1]$$

es la función de pertenencia de la relación borrosa $F(t)$, y

$$\text{Pos}^T \left\{ x(t) = U \right\}$$

se debe de leer como la posibilidad dinámica de que $x(t)$ pueda tomar a U como su valor en un tiempo t . En esta forma, $p(t)$ se puede traducir dentro de su forma canónica, o equivalentemente su ecuación de posibilidad dinámica de asignación:

$$p(t) \rightarrow \Pi_{x(t)}^T = F(t)$$

en la cual:

$$\Pi(x_1(t), \dots, x_n(t))^T = F(t)$$

es la distribución de posibilidad dinámica inducida por $p(t)$.

Si $p(t)$ es una proposición condicional, por ejemplo, si Roberto ahora trabaja en Cuernavaca, luego Roberto ahora vive cerca de Cuernavaca, la forma canónica de $p(t)$ se puede expresar como:

$$cf(p(t)) \text{ D Si } x(t) \text{ es } F(t) \text{ luego } y(t) \text{ es } G(t).$$

donde las variables focales dinámicas $x(t)$ y $y(t)$, toman valores en U y V respectivamente; $F(t)$ y $G(t)$ son subconjuntos borrosos de U y V ; y $x(t)$ es consecuente de $p(t)$. Correspondientemente, la ecuación dinámica de posibilidad de asignación asociada con p , llega a ser:

$$p(t) \rightarrow \Pi_{y(t)/x(t)}^T \text{ es } H(t)$$

en la cual $P(y(t)/x(t))$ denota la distribución de posibilidad condicional de $y(t)$ dada $x(t)$, y $H(t)$ está definido en términos de $F(t)$ y $G(t)$ por:

$$\mu_{H(t)}(U, V) = 1 \wedge^T [1 - \mu(t)F(t)(U) + \mu_{G(t)}(V)]$$

donde $u(t)$ y $v(t)$ son valores genéricos de $x(t)$ y $y(t)$ respectivamente, y:

$$\mu_{H(t)}(U, V) = U \times V \rightarrow^T [0,1]$$

es la función de pertenencia de $H(t)$, y:

$$\mu_F(t): U \rightarrow [0,1] \text{ y } \mu_{G(t)}: V \rightarrow [0,1]$$

son las funciones de pertenencia de $F(t)$ y $G(t)$, y $(\wedge)^T$ es el operador min, luego expresando en términos de funciones de distribución dinámica de posibilidades, implica que:

$$\Pi^i(x_1(t), \dots, x_n(t))$$

6.3. DEDUCCION.

Considerando que la base de conocimientos consiste de una colección finita de proposiciones $\{p_1(t), \dots, p_n(t)\}$, sea:

$$\Pi^i(x_1(t), \dots, x_n(t))$$

o simplemente:

$$\Pi^i$$

la distribución de posibilidad dinámica inducida por $p_j(t)$, $j = 1, \dots, n$. Luego bajo la consideración de que las $p_j(t)$ no son interactivas, la posibilidad dinámica de distribución de la distribución de posibilidad global, se puede expresar como:

$$\Pi^i(x_1(t), \dots, x_n(t)) = \bigwedge_{j=1}^n \Pi^j(x_1(t), \dots, x_n(t))$$

donde $\bigwedge^T$ es el min, y :

$$\Pi^i(x_1(t), \dots, x_n(t)) = \bigwedge_{j=1}^n \Pi^j(x_1(t), \dots, x_n(t))$$

donde $u_i(t)$ e U_i , con $i = 1, \dots, n$.

Supongamos que estamos interesados en los posibles valores en el tiempo de un subconjunto $\{x_{ij}(t), \dots, x_{ik}(t)\}$ del conocimiento de la distribución global dinámica de posibilidad $P(x_{ij}(t), \dots, x_{ik}(t))$. La distribución de posibilidad dinámica deseada está dada por la proyección de $P(x_1(t), \dots, x_n(t))$ en $(U_{i1} \times \dots \times U_{ik})$, que por simplicidad se escribe como:

$$P_s(t) = P(x_{ij}(t), \dots, x_{ik}(t))$$

Por conveniencia, sea $x_s(t)$ las subvariables de la variable: $x(t) = (x_1(t), \dots, x_n(t))$ que es el foco de nuestro interés, por ejemplo:

$$x_s(t) = (x_{ij}(t), \dots, x_{jk}(t))$$

donde el índice de la secuencia:

$$S \Delta (i_1, \dots, i_k)$$

es una subsecuencia de la secuencia $(1, \dots, n)$.

Sea $P_{x_s}(t)$ la proyección de la distribución dinámica global de posibilidades $P_x(t)$ en $U_s(t) \subset U_{i1} \times \dots \times U_{ik}$, luego por definición:

$$\Pi_s(t)[U_s(t)] \Delta \text{Sup}_{U_s(t)} \Pi_x(U)$$

donde $s'=(j_1, \dots, j_m)$ es el índice de la subsecuencia que es complementaria a S.

Ahora, el principio de vinculación de la lógica borrosa dinámica asegura que de cualquier proposición borrosa $p(t)$ se puede inferir una proposición borrosa dinámica $q(t)$, si la distribución de posibilidad inducida por $p(t)$, está contenida en la inducida por $q(t)$. Esto se puede representar en forma esquemática como:

$$\begin{aligned} p(t) &\xrightarrow{T} \Pi_{x(t)}^{p(t)} = F(t) \\ q(t) &\xleftarrow{T} \Pi_{x(t)}^{p(t)}(U) \quad \forall u(t) \in U \end{aligned}$$

donde " $\xleftarrow{T}$ " significa que $q(t)$ es una retranslación de la ecuación de asignación dinámica de posibilidades

$$\Pi_{x(t)}^{p(t)} \xrightarrow{T} G(t)$$

si la última es una translación de $q(t)$. Por simplicidad se diría que $\Pi_{x(t)}^{p(t)}$, se puede inferir de $\Pi_{x(t)}^{q(t)}$, si $\Pi_{x(t)}^{p(t)} \supset \Pi_{x(t)}^{q(t)}$, esto es:

$$\Pi_{x(t)}^{q(t)}(U) \xrightarrow{T} \Pi_{x(t)}^{p(t)}(U) \quad \forall u(t) \in U \text{ en un tiempo } t$$

de la definición de $Px(t)(U_s)$, se sigue que:

$$\Pi_s x(t)(U_s) \xrightarrow{T} \Pi x(t)(U) \text{ para } u \in U$$

en un tiempo t

y luego que $Px_s(t)$ se puede inferir de $Px(t)$ como consecuencia, $Px_s(t)$ representa la distribución de probabilidad deseada de la variable de interés, llamémosle:

$$x_s(t) \Delta (x_{i_1}(t), \dots, x_{j_k}(t))$$

Si x_s está dada por la proyección dinámica de $Px(t)$ en U_s , nos referiremos a la regla de inferencia que lleva $Px_s(t)$ como la P-regla dinámica con P como proyección dinámica. En resumen, dada una base de conocimientos, la distribución de posibilidad dinámica $Px_s(t)$ de una subvariable dinámica específica $x_s(t)$, se puede obtener por la proyección de la distribución dinámica global de posibilidades $Px(t)$ en U_s . La distribución dinámica de posibilidades resultante $Px_s(t)$, se puede expresar como la composición con respecto a $x_s(t)$ de aquellas distribuciones dinámicas de posibilidad constituyentes, las cuales contienen variables que están encadenadas directamente, o transitivamente a las variables en $x_s(t)$, por esta razón, la p regla dinámica puede ser descrita más sugestivamente, como la regla dinámica composicional de inferencia, expresada como una secuencia de operaciones. Se puede representar como la cadena:

$$\begin{array}{ccc}
(p_1(t), \dots, p_n(t)) & \xrightarrow[\text{traslacion}]{T} & (\Pi^1, \dots, \Pi^n) \xrightarrow[\text{conjuncion}]{T} \\
(\Pi^1 \wedge, \dots, \wedge \Pi^n) & \xrightarrow[\text{proyeccion}]{T} & (\Pi x_s(t)) \xrightarrow[\text{retraslacion}]{T}
\end{array}$$

donde la proposición dinámica inferida $q(t)$, es la retranslación dinámica de $Px_s(t)$, que es la distribución dinámica de posibilidades marginal de la subvariable de $(x_1(t), \dots, x_n(t))$, que es el blanco del proceso dinámico de inferencia.

La regla de inferencia tradicional que se puede ver como un caso especial de la regla composicional de inferencia, es el *Modus Ponens*. Para establecer este hecho, primero derivaremos de la regla composicional de inferencia una regla de inferencia más general que el *Modus ponens*, y que se llamara en lógica borrosa dinámica, el *Modus ponens generalizado dinámico*.

Consideremos un par de proposiciones dinámicas $\{p_1(t), p_2(t)\}$ de la forma:

$$\begin{aligned}
p_1(t) \text{ D si } x(t) \text{ es } F(t), \text{ luego } y(t) \text{ es } G(t). \\
p_2(t) \text{ D } x(t) \text{ es } F^*(t)
\end{aligned} \tag{6.1}$$

en las cuales $F(t)$, $F^*(t)$ y $G(t)$ son predicados borrosos dinámicos. La traslación se $p_1(t)$ y $p_2(t)$, se puede expresar como:

$$p_1(t) \xrightarrow{T} \Pi_{(y(t)/x(t))} \xrightarrow{T} H \tag{6.2}$$

donde $mH(u,v)$ está dada por (6.1).

$$p_2(t) \xrightarrow{T} \Pi_{x(t)} \xrightarrow{T} F' \tag{6.3}$$

Aplicando la regla de inferencia composicional a (6.2) y a (6.3), la distribución de posibilidades y se encuentra que es:

$$\Pi_{y(t)} \xrightarrow{T} H \circ F' \tag{6.4}$$

donde (6.4) en su lado derecho representa la composición dinámica de H y F^* con respecto a $x(t)$. Más concretamente:

$$\mu_{y(t)}(U) = \sup(\mu_H(t)(U, V) \wedge \mu_F(t)(U)) = \sup(\mu_f(t)(U) \wedge (1 - \mu_f(t)(U) + \mu_G(t)(V))$$

Esta conclusión se puede establecer en la forma de silogismo:

$$\begin{aligned}
&\text{Si } x(t) \text{ es } F(t), \text{ luego } y(t) \text{ es } G(t) \\
&\quad x(t) \text{ es } F^*(t) \\
&\quad y(t) \text{ es } G(t)
\end{aligned}$$

Este silogismo representa el Modus ponens dinámico. El Modus ponens dinámico difiere de su versión clásica, en dos aspectos: Primero, $F^*(t)$ no requiere ser idéntica a

$F(t)$ como en el caso clásico, y segundo los predicados dinámicos $F(t)$, $G(t)$, $F^*(t)$, no requieren ser descritos.

Si $F(t) \stackrel{T}{=} F(t)^*$, y los predicados son descritos $H(t)$ o $F^*(t)$, se reducen a $G(t)$, y:

Si $x(t)$ es $F(t)$, luego $y(t)$ es $G(t)$.
 $x(t)$ es $F(t)$.
 $y(t)$ es $G(t)$.

que es la forma clásica del modus ponens.

6.4. INFERENCIA DE PROPOSICIONES CUANTIFICADAS.

Por simplicidad restringiremos nuestro estudio a proposiciones cuantificadas cuyas formas canónicas sean de la forma:

$$q(t) \text{ D } Q(t) \text{ A's}(t) \text{ como B's}(t) \quad (6.5)$$

por ejemplo: Muchos coches con el tiempo son inseguros. Donde $A(t)$ y $B(t)$ son predicados borrosos y $Q(t)$ es un cuantificador borroso. Por conveniencia se dirá que $A(t)$ y $B(t)$ son el antecedente y el consecuente dinámicos de $q(t)$, esto está motivado por el hecho de que $q(t)$ se puede interpretar como la asignación dinámica de probabilidad condicional:

$$\text{Prob } \{B(t)/A(t)\} \text{ es } Q(t) \quad (6.6)$$

donde $\text{Prob } \{B(t)/A(t)\}$ denota la probabilidad condicional dinámica del evento borroso dinámico $\{x(t) \text{ es } B\}$, dado el evento borroso dinámico $\{x(t) \text{ es } A\}$. En (6.5) y (6.6) Q tiene el papel de un número borroso que es un subconjunto borroso dinámico del intervalo unitario.

6.5. CARDINALIDAD DE UN CONJUNTO BORROSO.

Para que cobre significado un cuantificador borroso dinámico, es necesario definir una forma de contar el número de elementos en un conjunto borroso dinámico, o sea determinar su cardinalidad.

Existen varias maneras de hacer esto. Para nuestro propósito será suficiente emplear el concepto de conteo sigma dinámico, que se define de la forma:

Sea F un subconjunto borroso dinámico de $U = \{u_1(t), \dots, u_n(t)\}$, expresado simbólicamente como:

$$F = \mu_1(t) / U_1 + \dots + \mu_n(t) / U_n \stackrel{T}{=} \sum_i \mu_i(t) / U_i$$

o simplemente como:

$$F \stackrel{T}{=} \mu_1 / U_1 + \dots + \mu_n / U_n$$

en el cual el término μ_i/u_i , con $i = 1, 2, \dots, n$, significa que μ_i es el grado de pertenencia de u_i en F , y el signo (+) representa la unión.

El conteo sigma de F está definido como la suma aritmética de las $\mu_i(t)$:

$$\sum_{\text{cuenta}(F)} \Delta \sum_{i=1}^n \mu_i(t)$$

con el entendido que la suma se redondea si es necesario al entero más próximo. Por lo tanto se puede estipular que los términos cuyos grados de pertenencia caigan abajo de un umbral específico, se excluyen de la sumatoria. El propósito de tal exclusión es evitarse la situación en la cual un gran número de términos con grados de pertenencia bajos, lleguen a ser contados como equivalentes a un pequeño número de términos con un grado de pertenencia alto. La cuenta sigma relativa denotada por $\sum_{\text{cuenta}}(F(t)/G(t))$ se puede interpretar como la porción de elementos de F que están en G en un tiempo t .

$$\sum_{\text{cuenta}} (F(t) / G(t)) = \frac{\sum_{\text{cuenta}} (F(t) \wedge G(t))}{\sum_{\text{cuenta}} G(t)} \text{ en un tiempo } t$$

donde $F \wedge G$ se define como:

$$\mu_{F(t)G(t)}(u_i) = \mu_{F(t)}(u_i) \wedge \mu_{G(t)}(u_i) \text{ con } u_i \in U$$

luego en términos de las funciones de pertenencia de F y G , la cuenta dinámica sigma relativa de F en G está dada por:

$$\sum_{\text{cuenta}} (F(t) / G(t)) = \frac{\sum_{\text{cuenta}} (\mu_F(t)(u_i) \wedge \mu_G(t)(u_i))}{\sum_{\text{cuenta}} \mu_G(t)(u_i)}$$

El concepto de suma sigma dinámica relativa nos da una base para interpretar el significado de proposiciones de la forma $QA's(t)$ son $B's(t)$. Por ejemplo: Los compañeros que hoy son jóvenes son saludables. Más específicamente, si la variable focal en la proposición dinámica se toma como la proposición de $B's(t)$ en $A's(t)$, luego la regla de traslación se puede expresar de la forma:

$$QA's(t) \text{ son } B's(t) \rightarrow \sum_{\text{cuenta}} (B(t)/A(t)) \text{ es } Q(t).$$

o equivalentemente como:

$$QA's(t) \text{ son } B's(t) \rightarrow \sum_{\text{cuenta}} P_x(t)'s = Q(t)$$

donde:

$$x(t) = \frac{\sum_{\text{cuenta}} (\mu_A(t)(u_i) \wedge \mu_B(t)(u_i))}{\sum_{\text{cuenta}} \mu_A(t)(u_i)}$$

6.6. EL SILOGISMO INTERSECCION/PRODUCTO

Esta es una regla básica de inferencia en lógica borrosa dinámica para proposiciones cuantificadas dinámicas y que se expresa en el esquema:

$$\begin{aligned} Q_1 A's(t) \text{ son } B's(t) \\ Q_2 (A(t) \text{ y } B(t)) \text{ son } C's(t) \\ (Q_1 \text{ t } Q_2) A's(t) \text{ son } (B(t) \text{ y } C(t))'s \end{aligned} \quad (6.7)$$

en los cuales $A(t)$, $B(t)$, $C(t)$ son predicados borrosos y $Q_1 \text{ t } Q_2$ es un número borroso que es el producto de los números borrosos Q_1 y Q_2 . Como un ejemplo podemos escribir:

Muchas de las personas que hoy son estudiantes son solteros.

Un poco más de la mitad de las personas que hoy son estudiantes solteros son mujeres.

(Muchos $\otimes$ pocos más de la mitad) de las personas que hoy son estudiantes son solteros y mujeres.

Si la intersección de $B(t)$ y $C(t)$ está contenida en $C(t)$, el siguiente corolario es consecuencia de (6.7).

$$\begin{aligned} Q_1 A's(t) \text{ son } B's(t) \\ Q_2 (A(t) \text{ y } B(t))'s \text{ son } C's(t) \\ \geq (Q_1 \text{ t } Q_2) A's(t) \text{ son } C's(t) \end{aligned} \quad (6.8)$$

donde el número borroso dinámico $\geq (Q_1 \text{ t } Q_2)$ se lee como el menor $(Q_1 \otimes Q_2)$ en el entendimiento de que $\geq (Q_1 \text{ t } Q_2)$ representa la composición de la relación binaria no borrosa $\geq$ con la relación borrosa dinámica unitaria $(Q_1 \otimes Q_2)$. En particular si los cuantificadores borrosos dinámicos Q_1 y Q_2 son monótonamente crecientes, por ejemplo cuando $Q_1 = Q_2$ D mucho se tiene:

$$\geq (Q_1 \overset{T}{\otimes} Q_2) = Q_1 \overset{T}{\otimes} Q_2$$

y (6.8) es:

$$\begin{aligned} Q_1 A's(t) \text{ son } B's(t) \\ Q_2 A's(t) \text{ son } C's(t) \\ (Q_1 \text{ t } Q_2) A's(t) \text{ son } C's(t) \end{aligned}$$

6.7. EL SILOGISMO DE CONJUNCIÓN CONSECUENTE.

Este silogismo se puede expresar con el esquema siguiente:

$$\begin{aligned} Q_1 A's(t) \text{ son } B's(t) \\ Q_2 A's(t) \text{ son } C's(t) \\ QA's(t) \text{ son } (B(t) \text{ y } C(t))'s \end{aligned} \quad (6.9)$$

donde:

$$0 \overset{T}{\oplus} (Q_1 \overset{T}{\otimes} Q_2 - 1) \geq Q \geq Q_1 \overset{T}{\oplus} Q_2$$

en la cual los operadores $\overset{T}{\otimes}, \overset{T}{\oplus}, \bullet$, y $-$ y las desigualdades $\geq$, son extensiones de $\wedge, \vee, +, -, \geq$ para números borrosos.

Como un ejemplo de este silogismo se tiene:

p1 D Muchos hombres mexicanos no son altos.
p2 D Muchos hombres mexicanos no son bajos.

de (6.9) se puede inferir:

Q hombres mexicanos no son altos ni son bajos.

En el ejemplo anterior la variable de interés es la conjunción de los subsecuentes no altos y no bajos. En el caso más general la variable de interés puede ser una función específica de las variables restringida por la base de conocimientos.

6.8. CADENAS DE PROPOSICIONES.

Una pareja ordenada $(p_1(t), p_2(t))$ de proposiciones cuantificadas de la forma:

$$\begin{aligned} p_1(t) &D Q_1 A's(t) \text{ son } B's(t) \\ p_2(t) &D Q_2 A's(t) \text{ son } C's(t) \end{aligned}$$

se dice que forman una cadena. Más generalmente una cadena de n se puede representar como una eneadada ordenada:

$$(Q_1 A_1's(t) \text{ son } B_1's(t), \dots, Q_n A_n's(t) \text{ son } B_n's(t))$$

en la cual:

$$B_1(t) \overset{T}{=} A_2(t), B_2(t) \overset{T}{=} A_3(t), \dots, B_{n-1}(t) \overset{T}{=} A_n(t)$$

Consideremos que $p_1(t)$ y $p_2(t)$ aparecen como premisas en el esquema de inferencia:

$$\begin{aligned} &Q_1 A's(t) \text{ son } B's(t) \text{ (premisa mayor)} \\ &Q_2 B's(t) \text{ son } C's(t) \text{ (premisa menor)} \\ &?Q A's(t) \text{ son } C's(t) \text{ (conclusión)} \end{aligned}$$

en la cual $?Q$ es un cuantificador borroso dinámico a ser determinado.

Cuando $Q_1 \overset{T}{=} Q_2 \overset{T}{=} \text{todo}$, la transitividad del contenido de los conjuntos borrosos dinámicos, o la regla de propiedad hereditaria en todo, implica que $Q \overset{T}{=} \text{todo}$.

Todos A's(t) son B's(t)
 Todos B's(t) son C's(t)
 Todos A's(t) son C's(t)

La propiedad de herencia, es una propiedad frágil en el sentido que siempre y cuando Q_1 y Q_2 son arbitrariamente cercanos a todo, hay nada que se pueda decir acerca de Q , que significa que es un número borroso, $Q = [0,1]$. Luego para restringir Q , es necesario hacer consideraciones restrictivas acerca de Q_1 y Q_2 , A(t) y B(t).

Una regla de encadenamiento importante que se usa en las relaciones naturales, es la del caso en que la premisa mayor en la cadena de inferencia:

Q_1 A's(t) son B's(t)
 Q_2 B's(t) son C's(t)
 Q A's(t) son C's(t)

es reversible en el sentido de que:

$$Q_1A's(t) \text{ son } B's(t) \overset{T}{\leftrightarrow} Q_1B's(t) \text{ son } A's(t)$$

donde $\overset{T}{\leftrightarrow}$ denota una equivalencia semántica temporal aproximada, por ejemplo:

Hoy muchos coches americanos son grandes $\overset{T}{\leftrightarrow}$ hoy muchos coches grandes son americanos.

Bajo la condición de reversibilidad, el siguiente silogismo vale en un sentido aproximado:

$$\begin{aligned} &Q_1A's(t) \text{ son } B's(t) \\ &Q_2B's(t) \text{ son } C's(t) \\ &\geq (0 \bullet (Q_1 \overset{T}{\oplus} Q_2 - 1)) \text{ son } C's. \end{aligned}$$

Este silogismo representa la regla dinámica de reversibilidad de la premisa mayor.

6.9. SEMANTICA Y PRAGMATICA DE LOS SISTEMAS BORROSOS.

Diseñar es establecer relaciones entre valores, y por lo mismo es acotarse en una forma intrincada dentro de una evaluación, así que en sistemas borrosos dinámicos, cualquier decisión de diseño es siempre implícita o explícitamente basada en una evaluación. Se presentan preguntas fundamentales en la hipótesis de que es describible lingüísticamente un sistema multiobjetivo de muchos atributos, y que todos los lenguajes naturales y artificiales de interés, se pueden dar por alguna transformación gramatical de algún tipo. Estas preguntas son:

1. ¿Cuál es el significado de una evaluación?
2. ¿Cuáles son las formas de las reglas operacionales, por las cuales se construye el

significado de evaluación borrosa dinámica de los espacios, elementos, y contextos de un sistema complejo y contradictorio?

La aproximación clásica para el uso de información empírica en la evaluación, ha sido el desarrollo de una teoría axiomática, en la cual se han definido ciertos conceptos y sus interrelaciones entre ellos. Estos conceptos se refieren a demandas, costos, valores, etc. La precisión es de fundamental importancia en la aproximación clásica, y es bien conocido que es difícil de lograr la descripción precisa, y por lo tanto la evaluación precisa de los micro y macro sistemas. Un sistema complejo en el cual el humano juega un papel dominante en comunicar, informar, decidir, e implantar políticas, claramente representa un macro sistema cambiante en el tiempo, cuya conducta debe de estar sujeta a principios de incompatibilidad dinámica. Este principio asegura la relación inversa entre complejidad y precisión.

La conducta de estos sistemas sólo puede ser descrita y evaluada realísticamente en una forma borrosa dinámica. Los conjuntos borrosos dinámicos nos ofrecen un marco de trabajo para descripciones y evaluaciones de sistemas borrosos dinámicos, y sería de interés la manipulación en un sentido posible de los elementos y contextos del sistema en un espacio prescrito. Sería bueno acordar técnicas para emplear una teoría lingüística.

No es nuevo el propósito de clasificar el estudio de los signos en sintaxis, semántica y pragmática. Anteriormente los problemas pragmáticos han sido tratados de una manera informal por los filósofos en el lenguaje ordinario, o por algunos lingüistas, pero los lógicos y filósofos de un cuadro formalista tienden a ignorarlos, o simplemente los tratan como problemas semánticos y sintácticos. Por supuesto, los pragmáticos bordan en lo semántico, así que es entendible la tentación de tratar un problema pragmático como si fuera semántico.

Tanto la semántica como la pragmática, tienen que ver con el concepto de significado. La semántica considera el significado en sentido formal (Epistemológico) y la pragmática en sentido práctico (Ontológico). Por supuesto, es mejor apreciado el aspecto total del significado, vía su aspecto formal.

Un aspecto importante de los lenguajes naturales, es la expresión de proposiciones dinámicas. Las reglas operacionales específicas para igualar sentencias con las proposiciones dinámicas que ellos expresan, es un problema semántico. En la generalidad de los casos, estas reglas no igualan directamente sentencias con proposiciones, pero pueden hacer esto solamente en lo relativo a aspectos de varios contextos, en los cuales se usa la sentencia. La sintaxis tiene que ver con las sentencias, la semántica trata con proposiciones, y la pragmática examina actos y el contexto en el cual ellas se representan. Las reglas de sintaxis y semántica para un lenguaje, determinarán e interpretarán la sentencia o cláusula, determinarán una proposición de verdad que puede cambiar en el tiempo. Una sentencia interpretada, corresponde a una función del contexto en un tiempo dado, y palabras posibles son determinantes parciales del valor de verdad de lo expresado por la sentencia dada en ese tiempo.

Podemos considerar una proposición como una función del contexto posible de las palabras en su tiempo a los valores de verdad, de esta forma se puede considerar una teoría semántica práctica, como un estudio de la manera en la cual, no las proposiciones,

pero si los valores de verdad, dependen del tiempo y del contexto, y parcialmente dependerán del contexto las posibles palabras con las cuales la sentencia se profiere. Para poner el análisis pragmático anterior en sistemas borrosos dinámicos tenemos:

(Evaluación borrosa dinámica) $\overset{T}{=}$ (interpretación dinámica) y (Proposición dinámica)

así que la evaluación de un sistema complejo, es una función del tiempo y de los contextos situacionales, a los valores de pertenencia de un subconjunto borroso dinámico del espacio pragmático. Basados en la riqueza y ambigüedad de la experiencia moderna, los sistemas revelan profunda complejidad y contradicción. En estos sistemas son inevitables numerosos niveles de significados y combinaciones de enfoques, así que su espacio, sus elementos intrínsecos y contextos de donde derivan, se pueden trabajar e interpretar en primera instancia, en una multiplicidad de formas. El problema de evaluación en estos sistemas, puede ser descrito como un ejercicio matemático en la manipulación de propiedades descritas lingüísticamente, de un sistema sintáctico dentro de objetos traducidos imprecisamente de un espacio semántico, así que la solución llega a ser una interpretación subjetiva de objetos en un espacio pragmático. El concepto dinámico de significado, el cual corresponde a un efecto evaluativo de un contexto situacional, es crucial en el diseño y evaluación de un sistema borroso dinámico en un espacio pragmático.

Según Popper, la estructura intencional de la física clásica consiste en:

1. La construcción de objetos fenoménicos idealizados y objetivizables, y
2. Una explicación de la realidad en términos de partículas clásicas, modelos espacialmente contruidos, y campos clásicos.

por lo tanto, los objetos clásicos están concebidos a ser algunas veces en el espacio tridimensional de la experiencia humana, estos no son objetos de una percepción en la forma ricamente emotiva de la vida diaria. Es por esta razón principalmente, que los humanistas del siglo XIX y XX, habían creído haber llegado al clímax de las ciencias físicas para representar la realidad, pero el objeto fenoménico es sólo un símbolo de la realidad, así que la realidad por ella misma está por debajo del símbolo, y muy probablemente desconocida. Se puede ver la realización borrosa dinámica de la representación del objeto fenoménico, cayendo en forma última en un espacio pragmático dinámico, luego la posible readaptabilidad de trabajo de los objetos y el espacio, debe de ser una nueva y práctica transliteración para la evaluación de sistemas borrosos dinámicos de la imposibilidad de conocimiento de la realidad que cambia con el tiempo. La pragmática de los conjuntos borrosos dinámicos, debe de ser el vehículo para la transliteración.

De los dos extremos, el empirismo que niega la posibilidad de una ontología de la naturaleza, y el subjetivismo, que busca el significado de la realidad en experiencias noéticas solamente, aparte de una realidad trascendente revelada a través de ellos, se han dividido entre los dos, filósofos y físicos, nos inclinamos a rechazar la forma, y buscaremos una aproximación operacional a la representación de los objetos fenoménicos. Esta corresponde más o menos con los objetivos de una metodología que es compatible con sistemas borrosos dinámicos, en consecuencia de esto llamaremos al objeto fenoménico evaluación, y lo representaremos por un subconjunto borroso

dinámico en un espacio apropiado. Los posibles valores de la evaluación se obtendrán a través de operaciones convenientes por amalgamamiento de los efectos evaluativos de todos los contextos situacionales, que por lo mismo son temporalmente dependientes.

Nuestra inclinación hacia el subjetivismo implica que el estudio de la evaluación del sistema borroso dinámico, debe de incluir el concepto dinámico de significado, a la luz de las experiencias neóticas. Como para describir estas experiencias emplearemos un lenguaje, los tres aspectos de éste: sintaxis, semántica y pragmática, deben de jugar su papel. El espacio que ellos generen, representará el espacio para la evaluación del conjunto borroso dinámico. En el espacio sintáctico, se revelarán lingüística y numéricamente las peticiones concretas e inmediatas de las propiedades físicas, en el dato axiológico de la medición individual. El espacio semántico contiene objetos que son inmanentes puros del contexto subjetivo. Estos objetos tienen significados epistemológicos (formales) de las propiedades del espacio sintáctico. El significado ontológico de las propiedades del espacio sintáctico, son objetos en el espacio pragmático. Los objetos del espacio pragmático, se refieren a una realidad que trasciende objetos inmanentes del espacio semántico. En el espacio pragmático, la realidad en las conveniencias factuales individuales, reveladas y mediadas por un proceso de evaluación, se establece de acuerdo con los parámetros de dependencia del contexto en el tiempo.

Estipulados los espacios anteriores se pueden asumir los siguientes acoplamientos:

1. Las propiedades son manipulables lingüísticamente y viceversa sus objetos, como objetos se definen imprecisamente en términos de las propiedades que cambian en el tiempo.
2. Las propiedades de los objetos individuales se definen en relación con otros objetos.

Luego atribuir una propiedad particular a un objeto es simplemente estipular la manera en la cual el objeto afecta y es afectado por otros objetos, bajo circunstancias particulares que cambian en el tiempo, consecuentemente la representación de las propiedades de un objeto no puede ser única, aparecerán diferentes representaciones de la estipulación de otros objetos que afectan y son afectados por el uno bajo consideración, y también de la estipulación de la manera en la cual se relacionan y estas relaciones cambian con el tiempo.

Además de los problemas epistemológicos en la fundamentación de la teoría de los conjuntos borrosos dinámicos, sumaremos otra, la determinación de la denotación de las entidades axiológicas, como la representación en variables lingüísticas dinámicas de estados ontológicos. A menudo los valores de las variables lingüísticas dinámicas no son claros, y se debe tender a clarificarlos, es aquí donde se tiende a establecer una interpretación pragmática de satisfacción en vez de mostrar la evaluación basada en una representación de variables lingüísticas dinámicas.

Decimos que estamos satisfechos en un contexto situacional dado cuando una evaluación nos da un soporte inductivo para otra, cuando muchos de los contextos en los cuales la primera evaluación es satisfactoria, son contextos en los cuales la segunda evaluación es satisfactoria también. En otras palabras, si se pueden definir algunas clases de condiciones de cercanía entre los contextos $r(t)$ y $s(t)$, donde $r(t)$ es un

contexto dinámico situacional en el cual estamos satisfechos para un tiempo dado por ciertas evaluaciones, luego el más cercano será $r(t)$, el más pragmático.

6.10. RESTRICCIONES BORROSAS DINÁMICAS

Si $x(t)$ es una variable dinámica que toma valores en el universo del discurso $U(t)$, luego una restricción borrosa dinámica $R[x(t)]$ en los valores que se pueden asignar a $x(t)$, es una relación borrosa dinámica en $U(t)$, tal que asignar un valor $u(t)$ a $x(t)$ en el tiempo t , requiera un relajamiento de la restricción expresada por:

$$S = 1 - f_{R[x(t)]} [u(t)]$$

Donde $f_{R[x(t)]} [u(t)]$ es el grado de membresía de $u(t)$ en $f_{R[x(t)]} [u(t)]$

$$x(t) = u(t): f_{R[x(t)]} [u(t)]$$

donde $x(t)$ denota un valor genérico de $x(t)$ en un tiempo t .

Las restricciones borrosas dinámicas nos dan una base para la formulación de las reglas de traslación para proposiciones borrosas dinámicas, esto es, proposiciones que contienen nombres de conjuntos borrosos dinámicos que cambian en el tiempo.

Por una traslación dinámica de una proposición borrosa dinámica, se entiende la representación del significado dinámico de una proposición borrosa como un sistema de ecuaciones relacionales de asignación. La traslación de una proposición borrosa dinámica tiene la forma:

$$\begin{aligned} p(t) \rightarrow R[X_1(t)] &= F_1(t) \\ p(t) \rightarrow R[X_2(t)] &= F_2(t) \\ &\vdots \\ p(t) \rightarrow R[X_n(t)] &= F_n(t) \end{aligned}$$

donde $X_1, \dots, X_n$ son variables que son explícita o implícitamente dependientes del tiempo en $p(t)$, $R[X_1(t)], \dots, R[X_n(t)]$ son las restricciones borrosas dinámicas en estas variables, y $F_1(t), \dots, F_n(t)$ son las relaciones borrosas dinámicas que son asignadas a $R[X_1(t)], \dots, R[X_n(t)]$ respectivamente en un tiempo t .

6.11. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO I

Las reglas de traducción borrosas dinámicas del tipo I son las reglas modificadoras. Estas reglas aseguran que la traducción dinámica de una proposición borrosa dinámica de la forma:

$$p(t) \equiv X(t) \text{ es } mF(t)$$

se expresa por:

$$X(t) \text{ es } mF(t) \rightarrow R[A(X(t))]^T = mF(t)$$

En campo en el cual los conjuntos borrosos dinámicos juegan un papel fundamental es en la teoría del razonamiento borroso dinámico. Un concepto central en el campo es la idea de las variables lingüísticas dinámicas y la relacionada de proposición canónica dinámica. Asumamos que $V_1(t)$ es una variable que puede tomar valores en el conjunto borroso dinámico $X(t)$ que llamaremos la base dinámica. Nuestro conocimiento del valor de la variable se conoce solamente como un valor lingüístico impreciso que cambia en el tiempo en el conjunto dinámico $X(t)$. Usando la teoría de los conjuntos borrosos dinámicos podemos asociar con este valor lingüístico un subconjunto borroso dinámico $T[x(t)]$ tal que para cada elemento $x(t)$ en $X(t)$, $T[x(t)]$ indica el grado en el cual $x(t)$ es compatible con un concepto en un tiempo t . Podemos representar nuestro conocimiento con:

$$V_1(t) \text{ is } A_1(t) \text{ en un tiempo } t$$

Donde $A_1(t)$ es un subconjunto dinámico del universo del discurso $V_1(t)$. A esta declaración la llamaremos una proposición dinámica canónica.

Dada una proposición $V(t)$ es $A(t)$, definimos:

$$\text{Pos}[V(t) \text{ es } A(t)]^T = \max_{x(t)} \{A[x(t)]\}$$

y $\text{Pos}[V(t) \text{ es } A(t)]$ es la medida de posibilidad asociada con la proposición en un tiempo t . Llamaremos a una proposición una tautología dinámica si:

$$A[x(t)]^T = 1 \forall x(t) \text{ en la base dinámica en un tiempo } t.$$

Diremos que una proposición es una tautología absoluta si:

$$A[x(t)]^T = 1 \forall x(t) \text{ en cualquier } t.$$

Introducimos el concepto de una variable dinámica articulada, como una colección de una o más variables atómicas dinámicas distintas. Una proposición canónica dinámica es una declaración de la forma $V(t)$ es $M(t)$, donde $V(t)$ es una variable dinámica articulada y $M(t)$ es un subconjunto borroso dinámico del producto cartesiano de las bases dinámicas de las variables atómicas dinámicas que constituyen la variable dinámica articulada.

6.12. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO II

Las reglas composicionales dinámicas del tipo II pertenecen a la traducción en el tiempo de la proposición $p(t)$ la cual es una composición de proposiciones $q(t)$ y $r(t)$. Las reglas dinámicas de traducción de estos modos de composición son las siguientes: Sea $x(t)$ y $y(t)$ variables que pueden tomar valores en $U(t)$ y $V(t)$ respectivamente, y sea $F(t)$ y $G(t)$ subconjuntos dinámicos de $U(t)$ y $V(t)$. Si:

$$x(t) \text{ is } F(t) \rightarrow \pi_x^T = F(t)$$

$$y(t) \text{ is } G(t) \rightarrow \pi_y^T = G(t)$$

luego:

$$x(t) \text{ is } F(t) \text{ and } y(t) \text{ is } G(t) \rightarrow \pi_{(x,y)}^T = \bar{F}(t) \cap \bar{G}(t)$$

$$x(t) \text{ is } F(t) \text{ or } y(t) \text{ is } G(t) \rightarrow \pi_{(x,y)}^T = \bar{F}(t) \cup \bar{G}(t)$$

$$\text{if } x(t) \text{ is } F(t) \text{ then } y(t) \text{ is } G(t) \rightarrow \pi_{(x,y)}^T = \bar{F}(t) \rightarrow \bar{G}(t)$$

6.13. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO III

Las reglas de cuantificación del tipo III aplicadas a las proposiciones dinámicas de la forma general $Qx(t)$ son $F(t)$, donde Q es un cuantificador borroso, $x(t)$ es una variable que toma valores en $U(t)$ y $F(t)$ es un subconjunto borroso dinámico en $U(t)$. En general un cuantificador borroso es un subconjunto borroso dinámico.

6.14. REGLAS DE TRADUCCIÓN BORROSAS DINÁMICAS DEL TIPO IV

Además de las muchas formas en que se puede cualificar una proposición dinámica $p(t)$, existen tres de particular importancia en el razonamiento borroso dinámico y son: Por el valor lingüístico dinámico de verdad, por el valor lingüístico dinámico probabilístico de verdad, y por el valor dinámico posibilístico de verdad. En lógica borrosa dinámica, el valor de verdad de una proposición dinámica $p(t)$ está definida como una compatibilidad dinámica de una proposición dinámica de referencia $r(t)$ con $p(t)$. Más específicamente, sea $p(t) = x(t) \text{ es } F(t)$, donde F es un subconjunto de $U(t)$, y sea $r(t)$ una proposición dinámica de referencia de la forma $r(t) = x(t) \text{ es } u(t)$, donde $u(t)$ es un elemento de $U(t)$, luego la compatibilidad de $r(t)$ con $p(t)$ se define como:

$$\text{Comp } [x(t) \text{ is } u(t)/x(t) \text{ is } F(t)].$$

6.15. RAZONAMIENTO BORROSO DINÁMICO

Informalmente vamos a entender por razonamiento borroso dinámico el proceso por el cual una posible conclusión imprecisa se deduce de una colección de premisas dinámica imprecisas. Este razonamiento es de naturaleza cualitativa y no cuantitativa, y cae fuera del dominio de aplicación de la lógica clásica.

Es esta capacidad de razonar en términos imprecisos y cambiantes la que nos diferencia de la forma de operar de las máquinas. En esta tesis se describe una nueva dirección que involucra el nuevo concepto de distribución dinámica de posibilidades, el cual nos dará una base natural para la representación del significado dinámico de las proposiciones dinámicas expresado en lenguaje natural. Comenzaremos nuestra discusión introduciendo el concepto de distribución dinámica de posibilidades.

6.16. EL CONCEPTO DE DISTRIBUCIÓN DINÁMICA DE POSIBILIDADES

La consideración básica que permite que nuestro concepto subyaga al razonamiento borroso dinámico, es en naturaleza posibilístico en vez de probabilístico. Para ilustrar lo anterior consideremos la proposición dinámica

$$p(t) = x(t)$$

es un entero en el intervalo $[0,8]$ puede ser posiblemente un valor de $x(t)$ y cualquier entero que no esté en el intervalo no puede ser un valor de $x(t)$. Luego para la proposición dinámica en cuestión

$$\text{Poss}(t)\{x(t) = n\} = 1 \text{ for } 0 \leq n \leq 8$$

y

$$\text{Poss}(t)\{X(t) = n\} = 0 \text{ for } n < 0 \text{ or } n > 8$$

Donde

$$\text{Poss}(t)\{X(t) = n\}$$

Es la abreviación para la posibilidad dinámica de que $x(t)$ asuma el valor n en el tiempo t . Nótese que la distribución dinámica de posibilidad inducida por p es uniforme en el sentido de que los valores de la distribución dinámica de posibilidades son iguales a la unidad para n en $[0, 8]$, y cero en cualquier otro valor en el tiempo t .

Diremos que una proposición dinámica borrosa de la forma $p(t) = x(t)$ es $F(t)$ donde x es una variable que toma valores en el tiempo en un universo del discurso $U(t)$, y $F(t)$ es un conjunto borroso dinámico de $U(t)$, induce una distribución dinámica de posibilidades $\pi_{x(t)}$ que es igual a $F(t)$. Luego

$$\pi_{x(t)} = F(t).$$

En esencia la distribución dinámica de posibilidades de $x(t)$ es un conjunto borroso dinámico que sirve para definir la posibilidad de que x asuma cualquier valor específico en $U(t)$. En forma más concreta si $u(t) \in U(t)$ y

$$\mu_{F(t)}: U(t) \rightarrow [0, 1]$$

es la función dinámica de membresía de $F(t)$, luego la posibilidad dinámica de que

$$x(t) = u(t) \text{ dado } x(t) \text{ es } F(t) \text{ es:}$$

$$\text{Poss}(t)\{X(t) = u(t) \mid X(t) \text{ is } F(t)\} = \mu_{F(t)}[u(t)], u(t) \in U(t).$$

Ya que el concepto de distribución dinámica de posibilidades coincide con el de conjunto borroso dinámico, la distribución dinámica de posibilidades se puede manipular por las reglas de los conjuntos borrosos dinámicos y en forma más particular

por las restricciones borrosas dinámicas.

Si $A(t)$ es un subconjunto borroso dinámico de $U(t)$, luego

$$\pi_{x(t)}^T = \text{Poss}(t)\{x(t) \text{ is } A(t)\} = \text{Sup}_{U(t)}\{\mu_{F(t)}[U(t)] \wedge \mu_{A(t)}[U(t)]\}$$

donde $\mu_{A(t)}$ es la función de membresía de $A(t)$ y $\wedge = \min$.

Para conjuntos borrosos dinámicos arbitrarios $A(t)$ y $B(t)$ de $U(t)$, la medida de posibilidad de la unión de $A(t)$ y $B(t)$ está dada por

$$\pi(A(t) \cup B(t))^T = \pi[A(t)]^T \vee \pi[B(t)]^T$$

donde $\vee = \max$.

La medida de la posibilidad dinámica no tiene la propiedad básica de la aditividad, o sea:

$$\pi[A(t) \cup B(t)]^T \neq \pi[A(t)]^T + \pi[B(t)]^T$$

si $A(t)$ y $B(t)$ son disjuntos.

6.17. ECUACIONES DE ASIGNACIÓN DE POSIBILIDAD DINÁMICA

La razón por la cual el concepto de distribución dinámica de posibilidades juega un papel importante en el razonamiento borroso dinámico se relaciona con nuestra asunción de que una proposición en el lenguaje natural se puede interpretar como la asignación de un subconjunto borroso dinámico a una distribución dinámica de posibilidades. Si $p(t)$ es una proposición dinámica en lenguaje natural podemos decir que $p(t)$ se traduce en ecuación dinámica de asignación

$$p(t) \rightarrow \pi_{(x_1(t), \dots, x_n(t))}^T = F(t)$$

donde $x_1(t), \dots, x_n(t)$ son variables que están explícitas o implícitas en $p(t)$.

$$\pi_{(x_1(t), \dots, x_n(t))}$$

es la distribución dinámica de posibilidades de la variable $x(t) = x_1(t), \dots, x_n(t)$, y $F(t)$ es una relación dinámica borrosa. En este contexto la ecuación de asignación de posibilidades

$$\pi_{(x_1(t), \dots, x_n(t))}^T = F(t)$$

será una traslación dinámica de $p(t)$ e inversamente $p(t)$ será una retraslación que se representa como

$$p(t) \stackrel{T}{\leftarrow} \pi_{(x_1(t), \dots, x_n(t))}^T = F(t)$$

En general una proposición de la forma $p(t) = x(t)$ es $F(t)$, donde $x(t)$ es el nombre de un objeto que cambia en el tiempo o una proposición dinámica no se traslada dentro de

$$p \stackrel{T}{\rightarrow} \pi_{x(t)}^T = F(t)$$

sino dentro de

$$p(t) \stackrel{T}{\rightarrow} \pi_{A[x(t)]}^T = F(t)$$

donde $A[x(t)]$ es un atributo implicado dinámico de $x(t)$.

6.18. PROYECCIÓN Y PARTICULARIZACIÓN

Son dos operaciones que son de particular relevancia para el razonamiento dinámico aproximado. Sea

$$\pi_{(x_1(t), \dots, x_n(t))}$$

una distribución de posibilidades que es una relación dinámica borrosa en

$$\pi_{(x_1(t), \dots, x_n(t))}$$

sea $s = (i_1(t), \dots, i_k(t))$ una subsecuencia dinámica de la secuencia $(1, \dots, n)$, y se denota la secuencia complementaria $s' = (j_1(t), \dots, j_m(t))$. En términos de cada secuencia un k-uplete de la forma

$$[A_{i_1}(t), \dots, A_{j_k}(t)]$$

se puede expresar en forma abreviada como $A_{s(t)}$.

La proyección de

$$\pi_{(x_1(t), \dots, x_n(t))}$$

en

$$U_{s(t)} \stackrel{T}{=} [U_{i_1}(t), \dots, U_{j_k}(t)]$$

es una distribución de posibilidad dinámica expresada por

$$\pi_{x(t)_{s(t)}} \stackrel{T}{=} \text{Pr oy}_{U_{s(t)}} \pi_{(x_1(t), \dots, x_n(t))}$$

y definida por

$$\Pi_{x(t)_{s(t)}} \stackrel{T}{=} \text{Sup}_{U_{s(t)}} \Pi_{x(t)} [u_1(t), \dots, u_n(t)]$$

donde

$$\Pi_{x(t)_{s(t)}}$$

es la función de distribución de posibilidades de

$$\pi_{x(t)s(t)}$$

establecida como el principio dinámico de proyección. La relación entre $x_{[s(t)]}$ y $\pi_{x[s(t)]}$ se puede expresar como: de la distribución dinámica de posibilidades

$$x_{[s(t)]} \text{ y } \pi_{x[s(t)]}$$

la distribución dinámica de posibilidades

$$\pi_{x(t)s(t)}$$

de la subvariable:

$$x_{[s(t)]} = [x_{i_1}(t), \dots, x_{j_k}(t)]^T$$

se puede inferir por proyección:

$$\pi_{(x_1(t), \dots, x_n(t))}$$

en $U[s(T)]$

$$\pi_{x[s(t)]} = \text{Proj}_{U[s(t)]}^T \pi_{(x_1(t), \dots, x_n(t))}$$

Regresando a la operación de particularización, sea

$$\pi_{(x_1(t), \dots, x_n(t))} = F(t)^T$$

que denota la distribución dinámica de posibilidades de la subvariable

$$x_{[s(t)]} = [x_{i_1}(t), \dots, x_{j_k}(t)]^T$$

Informalmente por la particularización de

$$\pi_{(x_1(t), \dots, x_n(t))}$$

se entiende la modificación de

$$\pi_{(x_1(t), \dots, x_n(t))}$$

que resulta de la estipulación de que la distribución de posibilidades de

$$\pi_{x(t)s(t)}$$

es $G(t)$. Más específicamente

$$\pi_{(x_1(t), \dots, x_n(t))} [\pi_{x(t)s(t)} = G(t)] = F(t) \cap \overline{G}(t)^T$$

6.19. EL CONCEPTO DE UNA VARIABLE LINGÜÍSTICA DINÁMICA

El concepto de una Variable Lingüística Dinámica puede ser visto como una herramienta para sistematizar el uso de una sentencia dinámica en lenguaje natural o artificial para el propósito de caracterizar los valores de variables dinámicas y describir

sus interrelaciones. En este papel el concepto de una variable lingüística sirve como una función básica en razonamiento borroso dinámico pero en la representación de valores dinámicos de variables dinámicas y en la caracterización de los valores de verdad de proposiciones borrosas dinámicas.

El concepto básico detrás del concepto de una variable lingüística dinámica es que el significado dinámico de los términos primarios dinámicos es dependiente del contexto mientras que el significado de los conectivos y modificadores no lo es. Un importante aspecto del concepto de una variable lingüística dinámica se relaciona con el hecho de que, en general, los conjuntos dinámicos de términos así como las variables dinámicas no están cerradas bajo ciertas operaciones que se pueden realizar en conjuntos borrosos dinámicos como unión, intersección, producto, etc.

El problema de encontrar un valor lingüístico de $x(t)$ cuyo significado dinámico se aproxima a un subconjunto borroso dinámico dado de $U(t)$ es llamado el problema de la aproximación lingüística dinámica.

CAPÍTULO 7

7.1. CONECTIVOS LOGICOS DINAMICOS Y FORMAS SILOGISTICAS
DINAMICAS

7.2. EL PRINCIPIO DE EXTENSIÓN Y LOS OPERADORES BORROSOS
DINÁMICOS

7.3. SUBCONJUNTOS BORROSOS DE TIPO II

7.4. CONJUNTO DE OPERADORES NO MONÓTONOS

7.5. CUANTIFICADORES LINGÜÍSTICOS

7.6. RAZONAMIENTO SILOGÍSTICO CON CUANTIFICADORES

7.1. CONECTIVOS LOGICOS DINAMICOS Y FORMAS SILOGISTICAS DINAMICAS

Intersección y unión en subconjuntos borrosos Dinámicos

Una cuestión de enorme interés en la teoría de los subconjuntos borrosos dinámicos es la selección del operador apropiado para la unión y la intersección para generalizar el conjunto clásico de operadores al ambiente borroso dinámico. Si consideramos que $A(t)$ y $B(t)$ son dos subconjuntos borrosos dinámicos definidos sobre el conjunto $X(t)$, podemos expresar la conjunción o intersección de estos dos subconjuntos borrosos dinámicos como un subconjunto borroso dinámico $E(t)$ de $X(t)$ denotado por

$$E(t) = A(t) \overset{T}{\cap} B(t),$$

y la disyunción o unión como el subconjunto $F(t)$ de $X(t)$ donde

$$F(t) = A(t) \overset{T}{\cup} B(t).$$

Debemos de hacer varias consideraciones. Una primera es que cualquier modo que elijamos para definir este operador, se debe de reducir al operador de los conjuntos ordinarios cuando los conjuntos borrosos dinámicos se reducen a conjuntos ordinarios precisos. Una segunda consideración es que estos operadores serán asumidos como operadores que sólo dependen de los valores del conjunto en un tiempo t , esto es:

$$E(t) \overset{T}{=} T(A(t), B(t)) \text{ y } F(t) \overset{T}{=} S(A(t), B(t)).$$

Una implicación importante de esta consideración es que varias de las formas de determinar conectivos borrosos dinámicos pueden ser vistos desde el punto de vista de lógica multivaluada dependiente del tiempo. Esta consideración también nos lleva a buscar bajo la teoría de las ecuaciones funcionales.

En el sentido más general un subconjunto borroso dinámico se puede ver como un mapeo de $X(t)$ dentro de algún conjunto $Z(t)$

$$A: X(t) \overset{T}{\rightarrow} Z(t).$$

El conjunto $Z(t)$ es el dominio del grado de membresía dependiente del tiempo. La forma de los operadores de agregación $T(t)$ y $S(t)$ están restringidos por la naturaleza y las propiedades del conjunto $Z(t)$. Podemos notar que en el ambiente de conjuntos precisos, el conjunto $Z = \{0,1\}$. El intervalo más usual para Z es el intervalo unitario, I . Otro conjunto para $Z(t)$ es un conjunto ordenado que se denota por L . La forma más general para $Z(t)$ es un conjunto $P(t)$ el cual es parcialmente ordenado y tiene dos elementos distintos 1 y 0 que son respectivamente el máximo y el mínimo del conjunto.

Se sugiere las operaciones de máximo y mínimo, luego:

$$E(t) \overset{T}{=} \text{Min} [A(t), B(t)], \quad F(t) \overset{T}{=} \text{Max}[A(t), B(t)].$$

Otra consideración en la selección del operador apropiado es el dominio del problema para el cual el conjunto borroso dinámico se usa como modelo. Son dos los dominios fundamentales, el dominio lógico dinámico y el problema de decisión multicriterio. Para capturar la forma del operador usado, se puede sugerir un número de propiedades que son requeridas o deseadas en el operador. Una de las propiedades generalmente esperadas de estos operadores es la conmutabilidad:

$$\begin{aligned}T(A(t), B(t)) &= T(B(t), A(t)) \\S(A(t), B(t)) &= S(B(t), A(t))\end{aligned}$$

Una segunda propiedad que generalmente se requiere es la asociatividad:

$$\begin{aligned}T(A(t), T(B(t), C(t))) &= T(TA(t), (B(t)), C(t)) \\S(A(t), S(B(t), C(t))) &= S(SA(t), (B(t)), C(t))\end{aligned}$$

El satisfacer estas dos condiciones significa que en construir la intersección o unión de una colección de conjuntos borrosos dinámicos el orden en el cual se hace la operación es importante. La asociatividad nos permite una definición de unión e intersección involucrando sólo dos conjuntos y extenderla a cualquier número de conjuntos. En este sentido la consideración de función de punto y asociatividad actúan como agentes simplificantes en la definición de la operación.

Otra propiedad que se requiere usualmente de estos operadores es la de monotisidad. Asumamos

$$A(t) \overset{T}{\geq} A(t') \text{ y } B(t') \overset{T}{\geq} B(t),$$

luego la monotisidad requiere que:

$$\begin{aligned}T(A(t), B(t)) &\overset{T}{\geq} T(A(t'), B(t')) \\S(A(t), B(t)) &\overset{T}{\geq} S(A(t'), B(t'))\end{aligned}$$

La propiedad de monotisidad interactúa con la relación de inclusión de conjuntos borrosos dinámicos. Asumamos que $A(t)$ y $B(t)$ son dos subconjuntos borrosos dinámicos, decimos:

$$A(t) \overset{T}{\subset} B(t)$$

si:

$$A(t) \overset{T}{\geq} B(t) \quad \forall t \in X$$

Si los operadores de intersección y unión satisfacen la condición de monotisidad, luego se puede mostrar que si:

$$A \overset{T}{\subset} A' \text{ y } B \overset{T}{\subset} B'$$

luego

$$A \overset{T}{\cap} B \subset A' \overset{T}{\cap} B', \text{ y } A \overset{T}{\cup} B \subset A' \overset{T}{\cup} B'.$$

El operador min se puede ver como el que nos da el mayor conjunto contenido en cualquiera de los dos otros conjuntos. Luego si:

$$E(t) \overset{T}{=} \text{Min} [A(t), B(t)],$$

E es el mayor conjunto tal que:

$$E \overset{T}{\subset} A, \text{ y } E \overset{T}{\subset} B$$

De manera similar el operador Max se puede ver como el que nos da el menor de los conjuntos contenido en los otros dos. Luego si

$$F(t) \overset{T}{=} \text{Max} [A(t), B(t)],$$

F es el menor conjunto tal que:

$$A \overset{T}{\subset} F, \text{ y } B \overset{T}{\subset} F$$

Si definimos la intersección y la unión por estas dos condiciones sólo nos da los operadores Max y Min.

Cuando el dominio de $Z(t)$ es el intervalo unitario, caracterizan por completo al operador las tres anteriores propiedades, además de la temporalidad y los requerimientos de satisfacer la situación precisa. Un operador $T: [0,1] \times [0,1] \rightarrow [0,1]$ se denomina operador t-norma si:

1. $T(A(t), B(t)) \overset{T}{=} T(B(t), A(t))$, conmutatividad;
2. $T(A(t), T(B(t), C(t))) \overset{T}{=} T(T(A(t), B(t)), C(t))$, asociatividad;
3. $T(A(t), B(t)) \geq T(C(t), D(t))$ si $A(t) \overset{T}{\geq} C(t)$ y $B(t) \overset{T}{\geq} D(t)$;
4. $T(A(t), 1) \overset{T}{=} A(t)$

Podemos ver que la condición (4) junto con la condición (3) implica que:

$$T(0, 1) = T(0, 0) = 0.$$

La condición (1) implica:

$$T(0, 1) = T(1, 0) = 0,$$

y se sigue que:

$$T(1, 1) = 1,$$

y se han capturado las propiedades de la intersección precisa.

Se puede mostrar que el operador min es la mayor de todas las t-normas posibles. Un importante ejemplo de la t-normas es el operador producto:

$$T[A(t), B(t)] = A(t) \cdot B(t)$$

Un operador $S: [0, 1] \times [0, 1] \rightarrow [0, 1]$ se llama un operador co-norma si:

1. $S(A(t), B(t)) = S(B(t), A(t))$.
2. $S(A(t), S(B(t), C(t))) = S(S(A(t), B(t)), C(t))$ asociatividad.
3. $S(A(t), B(t)) \geq S(C(t), D(t))$ si $A(t) \geq C(t)$ y $B(t) \geq D(t)$.
4. $S(A(t), 0) = 0$

Se puede ver fácilmente que:

$$S(1, 1) = S(1, 0) = S(0, 1) = 1, \quad S(0, 0) = 0$$

Esto nos puede mostrar que el operador max es el menor de todas las t-conormas:

$$\text{Max}(A(t), B(t)) \leq S(A(t), B(t)).$$

Se notará que la única distinción entre el operador $T(t)$ y el $S(t)$ es en las condiciones 4. Estas condiciones se pueden ver como las características definitorias del operador respectivo. La condición 4 del operador $T(t)$ implica que el menor argumento es el más influyente en la formulación del “y” como la condición 4 del operador $S(t)$ implica que el mayor argumento es el más influyente en la formulación de la “o”.

Si introducimos el operador de negación de los conjuntos borrosos, se puede obtener una muy interesante relación entre t-normas y t-conormas. Un mapeo $N: [0, 1] \rightarrow [0, 1]$ se llama operador de negación si:

1. $N(1) = 0; N(0) = 1$
2. $N(N(A(t))) = A(t)$. Involución
3. Si $A(t) \geq B(t)$, luego $N(A(t)) \leq N(B(t))$, orden inverso

La definición usual para la negación de conjuntos borrosos cuando $Z = I$ es que:

$$\neg A(t) = 1 - A(t)$$

la cual satisface las condiciones anteriores.

Podemos tener una clase general de negación

$$N(V(t)) = V(t)^*,$$

donde:

$$V^* = \begin{cases} 1, & 0 \leq V \leq a, \\ 1 - \frac{1}{2} \left[\frac{V-a}{b-a} \right]^n & a \leq V \leq b. \\ \frac{1}{2} \left[\frac{c-V}{c-b} \right]^n & b \leq V \leq c \\ 0, & V \geq c \end{cases}$$

Un operador cercanamente relacionado con la negación en conjuntos borrosos dinámicos, es el antónimo. Si asumimos que $A(t)$ es un conjunto borroso dinámico definido en el conjunto X , sea $X^T = [0, M]$, que es un subconjunto de la recta real. El antónimo del subconjunto borroso $A(t)$ denotado por $A^a(t)$ se define por:

$$A^a(t)^T = A(M - x(t))$$

La figura (7.1) muestra un ejemplo que nos permite diferenciar entre la negación y el antónimo:

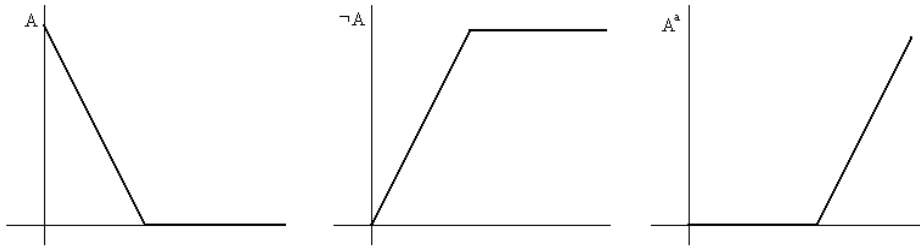


Figura (7.1)

La propiedad de indepotencia requiere que:

$$T(A(t), A(t))^T = A(t), \quad S(A(t), A(t))^T = A(t)$$

Los únicos operadores indepotentes son max y min. En el otro extremo de la indepotencia está la propiedad de Arquímedes que se define como:

$$\forall A(t) \in (0, 1) \quad T(A(t), A(t))^T < A(t) \quad \text{y} \quad S(A(t), A(t))^T = A(t)$$

Si la intersección y la unión son arquímedenas luego se tiene:

$$A(t)^T \cap A(t)^T \subset A(t), \quad A(t)^T \cup A(t)^T \supset A(t).$$

Una propiedad íntimamente relacionada con la anterior es la nilpotencia. Para cualquier secuencia de $a_i \in (0, 1)$ donde $i \in \mathbf{N}$ cualquier operador nilpotente satisface:

$$\begin{aligned} \exists n_0 < +\infty, \quad T(A(t)_1, A(t)_2, \dots, A(t)_{n_0}) &= 0 \\ \exists n_0 < +\infty, \quad S(A(t)_1, A(t)_2, \dots, A(t)_{n_0}) &= 1 \end{aligned}$$

Otra propiedad es la distributividad mutua:

$$\begin{aligned} A \cup (B \cap C) &= (A \cup B) \cap (A \cup C) \\ A \cap (B \cup C) &= (A \cap B) \cup (A \cap C) \end{aligned}$$

7.2. EL PRINCIPIO DE EXTENSIÓN Y LOS OPERADORES BORROSOS DINÁMICOS

Asumamos que $X(t)$, $Y(t)$, $Z(t)$ son tres conjuntos borrosos dinámicos no necesariamente distintos y sea f un mapeo

$$F(t): X(t) \times Y(t) \rightarrow Z(t)$$

El principio de extensión nos lleva a extender este mapeo de ser un mapeo puntual en el espacio $X(t) \times Y(t)$ a ser un mapeo en el espacio $I^{X(t)} \times I^{Y(t)}$. El principio de extensión nos lleva a extender el mapeo de f para tomar subconjuntos borrosos dinámicos como argumentos.

Consideremos $A(t)$ y $B(t)$ como subconjuntos borrosos en $X(t)$ y $Y(t)$ respectivamente, usando el principio de extensión se obtiene:

$$F(t)(A(t), B(t)) = D(t)$$

Donde $D(t)$ es un subconjunto borroso de $Z(t)$ que:

$$D(t) = \bigcup_{x(t), y(t)} \left[\frac{A(x(t)) \cap D(y(t))}{f(x(t), y(t))} \right]$$

Alternativamente podemos definir $D(t)$ para cada $z \in Z$ como:

$$D(z(t)) = \max_{x(t), y(t), f(x(t), y(t))} \left[A(x(t)) \cap D(y(t)) \right]$$

donde se entiende que $D(z(t)) = 0$ si no existen pares $(x(t), y(t)) \in X(t) \times Y(t)$ tales que:

$$f(x(t), y(t)) = z(t).$$

7.3. SUBCONJUNTOS BORROSOS DE TIPO II

El principio de extensión también se emplea para manejar los subconjuntos borrosos de tipo II. Un conjunto borroso del tipo II está definido como un subconjunto borroso dinámico cuya función de membresía es un subconjunto borroso dinámico.

Consideremos que $A(t)$ es un subconjunto borroso de $X(t)$ del tipo II, en este caso tendremos que para cada $x(t) \in X(t)$, el valor de pertenencia de $A(x(t))$ es en sí mismo un subconjunto borroso. En particular asumamos que $A(x(t))$ es un subconjunto borroso en el intervalo unitario. Sea $B(t)$ un subconjunto borroso de la misma clase. El problema ahora es definir el conjunto de operaciones en este conjunto. Asumamos que

$$D(t) = A(t) \overset{T}{\cap} B(t),$$

usando el principio de extensión podemos obtener a $D(t)$ como un subconjunto borroso del tipo II:

$$D(t) = \overset{T}{\cup}_{x(t)} \left[\frac{A(x(t)) \overset{T}{\cap} B(x(t))}{x(t)} \right]$$

que sería:

$$D(x(t)) = \overset{T}{\text{Min}} [A(x(t)), B(x(t))]$$

El problema es el cálculo del mínimo de los dos subconjuntos borrosos $A(x(t))$ y $B(x(t))$. Podemos ver que el mínimo es un mapeo en un tiempo t , $\overset{T}{\text{Min}}: I \times I \rightarrow I$. Usando el principio de extensión podemos extender el operador min de actuar sobre un punto en el espacio $I \times I$ en un tiempo t , a actuar en el subconjunto borroso dinámico $A(x)$ y $B(x)$. Si denotamos $D_{x(t)}(t)$ como el subconjunto borroso dinámico correspondiente al grado de membresía $x(t)$ en $D(t)$, $D(x(t))$, obtenemos:

$$D(z(t)) = \overset{T}{\max}_{x(t), y(t), f(x(t), y(t)) = W(t) \overset{T}{\cap} y(t)} \left[A_{x(t)}(x(t)) \overset{T}{\cap} D_{x(t)}(y(t)) \right]$$

En forma análoga se puede extender la operación unión. Si hacemos $E(t) = A(t) \overset{T}{\cup} B(t)$, luego:

$$E(z(t)) = \overset{T}{\max}_{x(t), y(t), f(x(t), y(t)) = W(t) \overset{T}{\cup} y(t)} \left[A_{x(t)}(x(t)) \overset{T}{\cap} D_{x(t)}(y(t)) \right]$$

Podemos usar el principio de extensión para extender la operación de complemento. Asumamos que $A(t)$ es un subconjunto dinámico del tipo II, luego definimos la negación dinámica

$$\overset{T}{\neg} A(t)$$

de la forma:

$$\overset{T}{\neg} A(t) = 1 - A(t)$$

Si $A(t)$ es un subconjunto borroso dinámico, el uso del principio de extensión nos da:

$$\overset{T}{\neg} A(x(t)) = \frac{A_{x(t)}(z(t))}{1 - z(t)}$$

Un principio cercanamente relacionado con el principio de extensión es llamado el principio de resolución. Este principio nos lleva a extender operaciones cuyos argumentos son conjuntos de subconjuntos borrosos dinámicos. Para introducir esta idea debemos de introducir el concepto de conjuntos borrosos dinámicos de nivel. Asumamos que $A(t)$ es un subconjunto borroso de $X(t)$. El subconjunto borroso dinámico de nivel- α de $X(t)$, denotado $A_{\alpha}(t)$ es un subconjunto ordinario de $X(t)$ tal que:

$$A_{\alpha}(t) = \{x(t) | A(x(t)) \geq \alpha\}$$

Este consiste en todos los elementos cuyo grado de membresía es α al menos. Usando los subconjuntos borrosos dinámicos de nivel se puede representar cualquier subconjunto borroso dinámico $A(t)$ por:

$$A(x(t)) = \sup_{\alpha} \left[\alpha \wedge A_{\alpha}(x(t)) \right]$$

Una representación alternativa usando subconjuntos de nivel es:

$$A(t) = \bigcup_{\alpha} \alpha A_{\alpha}(t)$$

donde $\bigcup$ es la unión y $\alpha A_{\alpha}(t)$ es el subconjunto borroso dinámico cuya función de membresía es $\alpha A_{\alpha}(x(t))$

Asumiendo que X es un conjunto y G es un mapeo $G: 2^{X(t)} \rightarrow Z(t)$. Para extender el mapeo a subconjuntos borrosos dinámicos, obtenemos:

$$G(t): I^{X(t)} \rightarrow Z(t),$$

donde para cualquier subconjunto $A(t)$ y $X(t)$,

$$G(A(t)) = \bigcup_{\alpha} \alpha G(A_{\alpha}(t))$$

Agregación de variables lingüísticas

Los conjuntos borrosos dinámicos son fundamentales en el dominio del razonamiento aproximado dinámico. Un concepto central en este dominio es la idea de un valor lingüístico dinámico y la idea relacionada de una proposición dinámica canónica. Asumamos a $V_1(t)$ como una variable que puede tomar su valor en el conjunto $X(t)$ llamado la base del conjunto. En muchos casos nuestro conocimiento del valor de una variable en lugar de ser precisamente dado como un elemento de $X(t)$, se conoce sólo por como un valor lingüístico impreciso en el tiempo t , basado en el conjunto $X(t)$. A partir de lo anterior podemos representar nuestro conocimiento con declaraciones de la forma $V_1(t)$ es $A_1(t)$, donde $A_1(t)$ es un subconjunto borroso dinámico del universo del discurso de $V_1(t)$. Llamaremos a estas declaraciones proposiciones atómicas canónicas dinámicas. Tales proposiciones indican que el valor de $V_1(t)$ cae en el subconjunto $A_1(t)$.

En forma más general podemos introducir el concepto de variable dinámica articulada. Una variable dinámica articulada es una colección de una o más variables dinámicas atómicas, y ejemplos de estas variables son:

$$V_1(t), (V_1(t), V_2(t)), (V_1(t), V_5(t), V_6(t))$$

Una proposición dinámica canónica en razonamiento aproximado dinámico es una declaración de la forma $V(t)$ es $M(t)$, donde $V(t)$ es una variable dinámica articulada y $M(t)$ es un subconjunto borroso dinámico del producto cartesiano del conjunto base de las variables atómicas que construyen la variable articulada.

Dada la proposición $V(t)$ es $A(t)$, definimos:

$$\text{Pos}_T[V(t) \text{ es } A(t)] = \max_{x(t)}^T [A(x(t))]$$

en donde a $\text{Pos}[V(t) \text{ es } A(t)]$ la llamaremos la medida de la posibilidad dinámica asociada con la proposición en un tiempo t . Diremos que la proposición es normal si $\text{Pos}[V(t) \text{ es } A(t)] = 1$. Si una proposición es normal, luego existe al menos un elemento en el conjunto base totalmente compatible con la proposición.

Una proposición se llama tautológica si: $A(x(t)) = 1 \forall x(t)$ en el conjunto base. Un hecho central en el desarrollo de la teoría del razonamiento aproximado dinámico es la agregación de esta clase de declaraciones dinámicas canónicas. El problema es compuesto por el hecho de que el conjunto base de varias proposiciones puede ser diferente.

Definición

Asumamos a $V_A(t)$ y $V_B(t)$ como dos variables articuladas de los conjunto base $X(t)$ y $Y(t)$. Sea $V_A(t)$ es $M(t)$ y $V_B(t)$ es $N(t)$, dos proposiciones dinámicas. Su conjunción denotada por $V_A(t) \text{ es } M(t) \perp V_B(t) \text{ es } N(t)$, es una proposición $V(t)$ es $P(t)$, donde $V(t)$ es una variable articulada consistente de la unión de las variables atómicas dinámicas que construyen a $V_A(t)$ y $V_B(t)$, y $P(t)$ es un subconjunto borroso del conjunto base $Z(t)$ de $V(t)$ (el producto cartesiano del conjunto base de las variables atómicas dinámicas que construyen $V(t)$). $P(t)$ está definido como tal para cada $z(t) \in Z(t)$.

$$P(z(t)) = M(x(t)) \wedge^T N(y(t))$$

Donde $x(t)$ es el elemento en $X(t)$ que concuerda con $z(t)$ en el dominio que tienen en común en el tiempo t y similarmente $y(t)$ es el elemento que concuerda con $z(t)$ en el dominio que tienen en común en el tiempo t .

Notaremos que en el caso cuando dos variables dinámicas articuladas son las mismas, esta operación es simplemente la operación intersección usual de conjuntos borrosos dinámicos.

$$V_A(t) \text{ es } M(t) \perp_t V_A(t) \text{ es } N(t) = V_A(t) \text{ es } M(t) \cap^T N(t)$$

En el caso cuando $V_A(t)$ y V_B son disociadas, es decir no tienen variables en común, esta operación nos da el producto cartesiano clásico

$$V_A(t) \text{ es } M(t) \perp_t V_B(t) \text{ es } N(t) \stackrel{T}{=} (V_A(t), V_B(t)) \text{ es } M(t) \times N(t).$$

La conjunción de una tautología con una proposición dinámica es una proposición dinámica, y es un hecho que simplemente se reduce a la proposición dinámica original. La conjunción de dos tautologías, es una tautología.

Definición

Asumamos que $V_A(t)$ y $V_B(t)$ son dos variables articuladas tales que $V_B(t)$ contiene todas las variables que están en $V_A(t)$, denotada por

$$V_A(t) \stackrel{T}{\subset} V_B(t).$$

La extensión cilíndrica de la proposición $V_A(t)$ es $M(t)$ a (t) es la proposición $V_B(t)$ es $M^o(t)$, donde

$$V_B(t) \text{ es } M^o(t) \stackrel{T}{=} V_A(t) \text{ es } M(t) \times V_1(t) \text{ es } X_1(t) \times V_2(t) \text{ es } X_2(t) \times \dots \times V_q(t) \text{ es } X_q(t),$$

con $V_1(t), V_2(t), \dots, V_q(t)$ las variables dinámicas atómicas que están en $V_A(t)$ pero no en $V_B(t)$. Las $X_j(t)$ son los respectivos conjuntos borrosos dinámicos base de estas variables. Se puede ver que la extensión cilíndrica dinámica es la conjunción de una proposición con una colección de tautologías y no cambia el contenido de la información de la proposición dinámica original.

La extensión cilíndrica de $V_A(t)$ es $M(t)$ a $V_B(t)$ es $M^o(t)$ se puede expresar en términos de la función de membresía de $M(t)$ y $M^o(t)$ como:

$$M^o(t)(y(t)) \stackrel{T}{=} M(t)(x(t))$$

Donde $x(t)$ es el elemento en el conjunto base de $V_A(t)$ que concuerda con y para la porción de $V_B(t)$ que es $V_A(t)$.

Teorema

Asumamos que $P_1(t), P_2(t)$, y $P_3(t)$ son tres proposiciones dinámicas.

Si $P_3(t) \stackrel{T}{=} P_1(t) \times P_2(t)$, luego $P_3(t) \stackrel{T}{\subset} P_1(t)$ y $P_3(t) \stackrel{T}{\subset} P_2(t)$.

La operación dinámica de contención juega un papel central en los mecanismos dinámicos de inferencia usados en razonamiento aproximado dinámico, en particular el principio dinámico de implicación establece que:

Si tenemos la proposición $V_B(t)$ es $N(t)$ luego se puede inferir cualquier proposición $V_A(t)$ es $M(t)$ tal que $V_B(t)$ es

$$N(t) \stackrel{T}{\subset} V_A(t) \text{ es } M(t).$$

Otra operación que juega un papel significativo en la teoría del razonamiento aproximado dinámico es la proyección.

Definición

Sea $V_A(t)$ es $M(t)$ alguna proposición canónica y sea $V_B(t)$ alguna variable articulada arbitraria. La proyección de la proposición anterior dentro de $V_B(t)$ denotada por

$$\text{proy}_{V_B}^t [V_A(t) \text{ es } M(t)]$$

es una proposición dinámica $V_B(t)$ es $N(t)$, donde $N(t)$ es el subconjunto borroso en el conjunto base de $V_B(t)$ definido por

$$N(z(t)) = \max_{x(t) \in G(t)}^T [M(x(t))],$$

donde $G(t)$ es el subconjunto del conjunto base $V_A(t)$ que contiene a todos los elementos que concuerdan con $z(t)$ en todas las variables que $V_A(t)$ y $V_B(t)$ tienen en común. Notamos que por definición, si $V_A(t)$ y $V_B(t)$ son disociados luego $N(z(t)) = 1$.

Siempre es el caso en el cual $V_A(t)$ es $M(t) \subset V_B(t)$ es $N(t)$, por lo que la operación dinámica de proyección se puede ver como un caso muy especial de contención.

El uso combinado de la conjunción y la proyección dinámica nos da la base de los mecanismos de inferencia en razonamiento aproximado dinámico. Considerando que tenemos una base de conocimiento que consiste en una colección de proposiciones canónicas dinámicas $P_1(t)$, $P_2(t)$, ..., $P_n(t)$. Sea $V(t)$ una variable cuyo valor estamos interesados en conocer basados en este conocimiento, el procedimiento es el siguiente:

1. Conjuntamos las proposiciones $P_1(t)$, $P_2(t)$, ..., $P_n(t)$ para obtener la proposición $P(t)$.
2. Proyectamos $P(t)$ dentro $V(t)$ para encontrar el valor de $V(t)$.

Podemos notar que si $V(t)$ es $A(t)$ es una proposición dinámica, podemos definir no $V(t)$ es a como:

$$V(t) \text{ es } \neg^T A(t) \text{ donde } \neg^T A(t) = 1 - A(z(t))$$

7.4. CONJUNTO DE OPERADORES NO MONÓTONOS

Para definir el conjunto básico de operaciones requerimos que las operaciones fueran monótonas en el sentido de que si:

$$A_1 \subset A_2 \text{ y } B_1 \subset B_2$$

Luego:

$$A_1(t) \overset{T}{\cap} B_1(t) \subset A_2(t) \overset{T}{\cap} B_2(t) \text{ y } A_1(t) \overset{T}{\cup} B_1(t) \subset A_2(t) \overset{T}{\cup} B_2(t)$$

Tratando de desarrollar operaciones de agregación que nos faciliten a hacer razonamiento de sentido común estamos forzados a dar de alta estas condiciones.

Para emular lo anterior definiremos una intersección monótona.

Definición

Asumamos $A(t)$ y $B(t)$ como dos subconjuntos de $X(t)$. Definimos el operador de intersección no monótona η como:

$$\eta(A(t), B(t)) \overset{T}{=} D(t)$$

donde $D(t)$ es un subconjunto borroso dinámico de $X(t)$ tal que:

$$D(x(t)) \overset{T}{=} A(x(t)) \overset{T}{\cap} B(x(t)) \overset{T}{\vee} (1 - \text{Pos}[B(t)|A(t)])$$

Recordemos que:

$$\text{Pos}[B(t)|A(t)] \overset{T}{=} \max_{x(t)} \left[A(x(t)) \overset{T}{\cap} B(x(t)) \right]$$

La primera observación es que la operación no es conmutativa. En lo particular podemos ver que si

$$A(t) \overset{T}{\cap} B(t) = \phi \text{ luego } \eta(A(t), B(t)) \overset{T}{=} A \text{ y } \eta(B(t), A(t)) \overset{T}{=} B(t).$$

Una segunda observación es que el operador no es puntual, esto es, $D(x(t))$ no es justamente una función de $A(x(t))$ y $B(x(t))$, pero por $\text{Pos}(B(t)|A(t))$ depende sólo de $A(t)$ y $B(t)$ para todos los valores de $X(t)$. La tercera observación es que el operador no es monótono. Consideremos:

$$A(t) \overset{T}{=} \{x_1(t), x_2(t), x_3(t), x_4(t)\}, B_1(t) \overset{T}{=} \{x_2(t), x_5(t)\}, B_2 \overset{T}{=} \{x_5(t)\}.$$

En este caso:

$$D_1(t) \overset{T}{=} \eta(A(t), B_1(t)) \overset{T}{=} \{x_2(t)\}, \quad D_2(t) \overset{T}{=} \eta((A(t), B_2(t))) \overset{T}{=} \{x_1(t), x_2(t), x_3(t), x_4(t)\}$$

Por lo que siempre a través $B_2(t) \overset{T}{\subset} B_1(t)$ se tiene $D_1(t) \overset{T}{\subset} D_2(t)$.

Familiarmente podemos notar que si la $\text{Pos}[A(t)|B(t)] \overset{T}{=} 1$ luego $N(A(t), B(t)) \overset{T}{=} A(t) \overset{T}{\cap} B(t)$

Definición

Asumimos que A y B son dos subconjuntos borrosos de X. Definimos la típica unión dinámica no monótona ψ como:

$$\psi(A(t), B(t)) = E(t)$$

donde E(t) es un subconjunto borroso dinámico de X(t) tal que:

$$E(x(t)) = A(x(t)) \wedge (B(x(t)) \wedge \text{Pos}[\neg A(t) \neg B(t)])$$

Podemos ver que el operador dinámico es no conmutativo, no puntual y no monótono. Esto se reduce a una unión ordinaria si

$$\text{Pos}[\neg A | \neg B] = 1.$$

7.5. CUANTIFICADORES LINGÜÍSTICOS

Los cuantificadores lingüísticos son de tres tipos: Incremento, decremento y unimodal.

Un cuantificador de incremento está caracterizado por la relación

$$Q(r_1(t)) \geq Q(r_2(t)), \text{ si } r_1(t) > r_2(t)$$

Estos cuantificadores están caracterizados por valores como: *al menos a en el tiempo t*, *todos en el tiempo*, *muchos en el tiempo*, etc.

Un cuantificador de decremento está caracterizado por la relación

$$Q(r_1(t)) \leq Q(r_2(t)), \text{ si } r_1(t) < r_2(t)$$

Estos cuantificadores están caracterizados por valores como: *algunos en el tiempo t*, *a lo mucho a en el tiempo t*.

Un cuantificador unimodal tiene la propiedad de que:

$$Q(a(t)) \leq Q(b(t)) \leq Q(c(t)) = 1 \geq Q(d(t))$$

para alguna

$$a(t) \leq b(t) \leq c(t) \leq d(t)$$

en un tiempo t

Son útiles para representar términos como *alrededor de q en el tiempo t*.

En muchos casos existen relaciones duales entre cuantificadores dinámicos de incremento y decremento, especialmente para cuantificadores dinámicos proporcionales. Asumamos que $Q_1(t)$ y $Q_2(t)$ son cuantificadores dinámicos definidos en el intervalo unitario. Si $Q_1(t)$ es un cuantificador dinámico de incremento, luego $Q_2(t)$ definido como

$$Q_2(r(t)) \stackrel{T}{=} Q_1(1 - r(t))$$

Se llama el dual dinámico y es un cuantificador dinámico del tipo de decremento.

Un concepto que juega un papel importante en cuantificadores es la idea de la especificidad de un cuantificador dinámico. Un cuantificador como muy cercano a 3 en el tiempo t se ve más específico que el cuantificador alrededor de 3 en el tiempo t . $Q_1(t)$ y $Q_2(t)$ son dos cuantificadores tales que:

$$Q_1(r(t)) \stackrel{T}{\leq} Q_2(r(t)) \quad \forall r(t) \text{ en un tiempo } t,$$

Luego $Q_1(t)$ es más específico que $Q_2(t)$.

Se puede ver que el cuantificador dinámico *para todo* en el tiempo t definido por:

$$\begin{aligned} Q(r(t)) &\stackrel{T}{\leq} 1 \text{ si } r(t) \stackrel{T}{=} 1 \\ Q(r(t)) &\stackrel{T}{=} 0 \text{ si } r(t) \stackrel{T}{\neq} 1 \end{aligned}$$

Es el cuantificador dinámico de incremento más específico. Su dual *ninguno en el tiempo t* está definido por:

$$\begin{aligned} Q(r(t)) &\stackrel{T}{=} 1 \text{ si } r(t) \stackrel{T}{=} 0 \\ Q(r(t)) &\stackrel{T}{=} 0 \text{ si } r(t) \stackrel{T}{\neq} 0 \end{aligned}$$

Es el más específico de todos los cuantificadores dinámicos de decremento.

La negación de *ninguno debe de existir en el tiempo t* , está definida por:

$$\begin{aligned} Q(r(t)) &\stackrel{T}{=} 0 \text{ si } r(t) \stackrel{T}{=} 0 \\ Q(r(t)) &\stackrel{T}{=} 1 \text{ si } r(t) \stackrel{T}{\neq} 0 \end{aligned}$$

y es el menos específico de todos los cuantificadores dinámicos de incremento.

Un concepto que juega un papel significativo en la manipulación de estos cuantificadores lingüísticos es la idea de cardinalidad de un subconjunto borroso dinámico. Consideremos que $A(t)$ es un subconjunto borroso dinámico de $X(t)$; la cuenta sigma dinámica de $A(t)$ denotada por

$$\begin{aligned} &\Sigma \text{cuenta}(A(t)) \\ &\text{está definida como:} \\ &\sum \text{cuenta}(A(t)) \stackrel{T}{=} \sum_i A(x_i(t)) \end{aligned}$$

En algunos casos se usa una cuenta sigma transformada:

$$\sum \text{cuenta}(A(t); f(t)) = \sum_i^T f(t) A(x_i(t)),$$

donde $f: [0, 1] \rightarrow [0, 1]$ está dada por:

$$f(0) = 0, \quad f(1) = 1, \quad f(a(t)) \geq f(b(t)), \text{ si } a(t) > b(t) \text{ en un tiempo } t.$$

Si $A(t)$ y $B(t)$ son subconjuntos borrosos dinámicos, la cuenta relativa sigma dinámica de $A(t)$ y $B(t)$ se define como:

$$\sum \text{cuenta}(A(t)|B(t)) = \frac{\sum \text{cuenta}(A(t) \cap B(t))}{\sum \text{cuenta}(B(t))}$$

Si asumimos la intersección dinámica definida por el operador min, luego:

$$\sum \text{cuenta}(A(t)|B(t)) = \frac{\sum \text{cuenta}(A(x_i(t)) \wedge B(x_i(t)))}{\sum \text{cuenta}(B(x_i(t)))}$$

La $\Sigma \text{Cuenta}(A(t)|B(t))$ es la proporción de elementos en $B(t)$ que también están en $A(t)$. Algunas identidades importantes que se pueden establecer son las siguientes:

1. $\Sigma \text{Cuenta}(A(t) \cap B(t)) + \Sigma \text{Cuenta}(A(t) \cup B(t)) = \Sigma \text{Cuenta}(A(t)) + \Sigma \text{Cuenta}(B(t)).$
2. Si $A(t) \cap B(t) = \emptyset$ luego, $\Sigma \text{Cuenta}(A(t) \cup B(t)) = \Sigma \text{Cuenta}(A(t)) + \Sigma \text{Cuenta}(B(t)).$
3. $\Sigma \text{Cuenta}(A(t)) \vee \Sigma \text{Cuenta}(B(t)) \leq \Sigma \text{Cuenta}(A(t) \cup B(t)) \leq \Sigma \text{Cuenta}(A(t)) + \Sigma \text{Cuenta}(B(t)).$
4. $(0 \vee \Sigma \text{Cuenta}(A(t)) + \Sigma \text{Cuenta}(B(t)) - \text{Cuenta}(X(t))) \leq \Sigma \text{Cuenta}(A(t) \cap B(t)) \leq \Sigma \text{Cuenta}(A(t)) \wedge \Sigma \text{Cuenta}(B(t)).$
5. $\Sigma \text{Cuenta}(A(t)|B(t)) + \Sigma \text{Cuenta}(\neg A(t)|B(t)) \geq 1.$

7.6. RAZONAMIENTO SILOGÍSTICO CON CUANTIFICADORES

Un concepto que juega un papel central es el de *Razonamiento Dinámico Silogístico Borroso*. Como un esquema general de inferencia un silogismo borroso se puede expresar en la forma:

$$\begin{array}{l} Q_1(t) A's(t) \text{ son } B(t) \\ Q_2(t) C's(t) \text{ son } D(t) \\ Q_3(t) C's(t) \text{ son } F(t) \end{array}$$

Donde $Q_1(t)$ y $Q_2(t)$ son cuantificadores borrosos dados, $Q_3(t)$ es un cuantificador borroso que será determinado y $A(t)$, $B(t)$, $C(t)$, $D(t)$, $E(t)$ y $F(t)$ son predicados borrosos dinámicos interrelacionados.

En lo siguiente asumiremos que $A(t)$ y $B(t)$ son subconjuntos borrosos dinámicos de $X(t)$ y $Q(t)$ es un cuantificador lingüístico dinámico absoluto y Q^* es un cuantificador lingüístico dinámico relativo. Por simplicidad indicaremos que la cuenta sigma como simplemente la cuenta y la cuenta sigma relativa como proporción (Prop).

Recordemos que la forma de una proposición en la teoría del razonamiento aproximado es $V(t)$ es $F(t)$. La proposición *hay* $Q(t)$ A 's(t) se puede trasladar a la proposición *Cuenta* $(A(t))$ es $Q(t)$.

Donde Cuenta $(A(t))$ es la variable y $Q(t)$ es el valor en un tiempo t .

También se sugiere que $Q(t)$ A 's(t) son $B(t)$ se traduzca en Cuenta $(A(t) \cap B(t))$ es $Q(t)$, y finalmente sugiere que $Q(t)^* A$'s(t) son $B(t)$, se traduzca en Prop($B(t) | A(t)$) es $Q(t)^*$.

En el razonamiento silogístico dinámico tenemos una o más declaraciones cuantificadas y se está interesado en inferir otras declaraciones. El principio de extensión cuantificado es un medio para hacer este tipo de razonamiento.

Consideremos que se tiene una colección de proposiciones $P_1(t)$, $P_2(t)$, ..., $P_n(t)$, cada una de las cuales es una caracterización borrosa dinámica de una cardinalidad absoluta o relativa expresada como:

$$P_i(t): C_i(t) \text{ es } Q_i(t)$$

En adelante $C_i(t)$ es una variable dinámica correspondiente a alguna cuenta y $Q_i(t)$ es un cuantificador dinámico, absoluto o relativo. El principio de extensión dinámico cuantificado establece que de la colección de premisas $P_1(t)$, $P_2(t)$, ..., $P_n(t)$ se puede concluir una premisa $P(t)$,

$$P(t): C(t) \text{ es } Q(t)$$

donde $C(t)$ es una variable dinámica y $Q(t)$ un cuantificador.

Más específicamente, si

$$C(t) = g^T(C_1(t), C_2(t), \dots, C_n(t))$$

y

$$Q(t) = g^T(Q_1(t), Q_2(t), \dots, Q_n(t))$$

La idea es que si estamos en la posibilidad de encontrar alguna relación funcional entre las C_i 's(t) que representan a $C(t)$, se puede encontrar la forma de $Q(t)$ en un tiempo t .

Desarrollaremos el silogismo intersección/producto dinámico.

$$\begin{aligned} P_1(t): Q_1(t) A's(t) \text{ son } B(t) \\ P_2(t): Q_2(t) (a(t) \text{ y } B(t)) \text{ son } D(t) \\ P(t): (Q_1(t) * Q_2(t)) A's(t) \text{ son } (B(t) \text{ y } D(t)) \end{aligned}$$

En adelante consideraremos que $Q_1(t)$ y $Q_2(t)$ son cuantificadores dinámicos relativos. Empleando las reglas de traslación dinámicas se tiene:

$$\begin{aligned}
P_1(t) &\xrightarrow{T} \text{Prop}(B(t) \mid A(t)) \text{ es } Q_1(t) = C_1(t) \text{ es } Q_1(t) \\
P_2(t) &\xrightarrow{T} \text{Prop}(D(t) \mid A(t) \cap B(t)) \text{ es } Q_2(t) = C_2(t) \text{ es } Q_2(t) \\
P_3(t) &\xrightarrow{T} \text{Prop}(B(t) \cap D(t)) \text{ es } Q(t) = C(t) \text{ es } Q_2(t)
\end{aligned}$$

y por lo mismo:

$$\begin{aligned}
C_1(t) &= \text{Pr op}[B(t) \mid A(t)] = \frac{\sum \text{cuenta}(A(t) \cap B(t))}{\sum \text{cuenta}(A(t))} \\
C_1(t) &= \text{Pr op}\left[D(t) \mid A(t) \cap B(t)\right] = \frac{\sum \text{cuenta}(A(t) \cap B(t) \cap D(t))}{\sum \text{cuenta}(A(t) \cap B(t))} \\
C_1(t) &= \text{Pr op}\left[B(t) \cap D(t) \mid A(t)\right] = \frac{\sum \text{cuenta}(A(t) \cap B(t) \cap D(t))}{\sum \text{cuenta}(A(t))}
\end{aligned}$$

Por lo que:

$$\frac{\sum \text{cuenta}(A(t) \cap B(t) \cap D(t))}{\sum \text{cuenta}(A(t))} = \frac{\sum \text{cuenta}(A(t) \cap B(t) \cap D(t))}{\sum \text{cuenta}(A(t) \cap B(t))} * \frac{\sum \text{cuenta}(A(t) \cap B(t))}{\sum \text{cuenta}(A(t))}$$

se sigue que:

$$C(t) = C_1(t) * C_2(t)$$

y luego:

$$Q(t) = Q_1(t) * Q_2(t)$$

El silogismo dinámico intersección/producto es central en el desarrollo de un número importante de otros silogismos dinámicos. En particular si $B(t) \subset^T A(t)$, se puede mostrar que:

$$\begin{aligned}
P_1(t): & Q_1(t) \text{ A's(t) son } B(t) \\
P_2(t): & Q_2(t) \text{ B's(t) son } D(t) \\
P(t): & (Q_1(t) * Q_2(t)) \text{ A's(t) son } (B(t) \text{ y } D(t))
\end{aligned}$$

Si $B(t) \cap D(t) \subset^T D(t)$ se puede obtener el silogismo dinámico en cadena:

$$\begin{aligned}
P_1(t): & Q_1(t) \text{ A's(t) son } B(t) \\
P_2(t): & Q_2(t) \text{ B's(t) son } D(t) \\
P(t): & \text{al menos } (Q_1(t) * Q_2(t)) \text{ A's(t) son } (B(t) \text{ y } D(t))
\end{aligned}$$

Si además $Q_1(t)$ y $Q_2(t)$ son monótonamente crecientes se obtiene el silogismo dinámico producto.

$$P_1(t): Q_1(t) \text{ A's(t) son } B(t)$$

$$\begin{aligned} P_2(t): Q_2(t) \text{ B's}(t) \text{ son D}(t) \\ P(t): (Q_1(t) * Q_2(t)) \text{ A's}(t) \text{ son C}(t) \end{aligned}$$

Usando el concepto de reversibilidad de la premisa mayor, se desarrolla otro silogismo dinámico llamado RPM. Lo primero que vemos es que RPM indica

$$Q_1(t) \text{ A's}(t) \text{ son B}(t) \overset{T}{\leftrightarrow} Q_1(t) \text{ B's}(t) \text{ son A}(t)$$

La condición en muchos casos es sólo una verdad aproximada. Bajo estas condiciones el silogismo dinámico PRM establece:

$$\begin{aligned} P_1(t): Q_1(t) \text{ A's}(t) \text{ son B}(t) \\ P_2(t): Q_2(t) \text{ B's}(t) \text{ son C}(t) \\ P(t): \text{A lo menos } (0 \overset{T}{\vee} (Q_1(t) + Q_2(t) - 1) \text{ A's}(t) \text{ son (B}(t) \text{ y D}(t)) \end{aligned}$$

Otra forma de silogismo dinámico es el de conjunción consecuencia.y es de la forma:

$$\begin{aligned} P_1(t): Q_1(t) \text{ A's}(t) \text{ son B}(t) \\ P_2(t): Q_2(t) \text{ B's}(t) \text{ son C}(t) \\ P(t): \text{A's}(t) \text{ son (B}(t) \text{ y C}(t)) \end{aligned}$$

donde

$$(0 \overset{T}{\vee} (Q_1(t) + Q_2(t) - 1) \overset{T}{\leq} Q(t) \overset{T}{\leq} Q_1(t) \overset{T}{\wedge} Q_2(t).$$

CAPÍTULO 8

- 8.1. RAZONAMIENTO APROXIMADO DINAMICO
- 8.2. VARIABLES LINGÜÍSTICAS DINÁMICAS
- 8.3. VARIABLES LINGÜÍSTICAS COMPUESTAS Y COMPARADORES BORROSOS DINÁMICOS
- 8.4. DECLARACIONES CONDICIONALES BORROSAS
- 8.5. REPRESENTACIÓN DE LA IMPLICACIÓN DINÁMICA MULTIVALUADA
- 8.6. IMPLICACIONES VALUADAS EN UN INTERVALO
- 8.7. RESULTADOS EN CONDICIONAMIENTO
- 8.8. CALIFICACIÓN DE VERDAD
- 8.9. CALIFICACIÓN DE CERTIDUMBRE Y POSIBILIDAD. EL CASO PRECISO.
- 8.10. REGLAS DE CUALIFICACIÓN DE VERDAD
- 8.11. REGLAS DE CUALIFICACIÓN DE INCERTIDUMBRE
- 8.12. EL MODUS PONENS GENERALIZADO
- 8.13. APROXIMACION LOGICA DEL RAZONAMIENTO APROXIMADO DINAMICO
- 8.14. FACTORES DE CERTIDUMBRE Y LÓGICA MULTIVALUADA
- 8.15. EL PUNTO DE VISTA TEÓRICO POSIBILÍSTICO DE LOS FACTORES DE CERTEZA
- 8.16. UN PARADIGMA DE LOS SISTEMAS DE INFORMACIÓN
- 8.17. EL PROBLEMA DE COMPOSICIONALIDAD
- 8.18. SISTEMAS FORMALES
- 8.19. TIPOS DE ECUACIONES RELACIONALES DINÁMICAS

8.1. RAZONAMIENTO APROXIMADO DINAMICO

El principio del mínimo de especificidad

Sea $U(t)$ un conjunto dinámico que representa el intervalo de la variable dinámica $x(t)$. Una variable dinámica no denota generalmente un objeto preciso capaz de satisfacer o no un predicado como una variable lógica dinámica lo hace, sin embargo empieza de su valor desconocido tomado de algún atributo. Una distribución dinámica de posibilidades $\pi_x(t)$ en $U(t)$ es un mapeo de $U(t)$ al intervalo unitario $[0,1]$ atribuido a la variable $x(t)$. La función $\pi_x(t)$ representa una restricción flexible del valor de $x(t)$ con la siguiente convención:

$$\begin{aligned}\pi_x(t)(u) = 0 & \text{ significa que } x(t) = u \text{ es imposible} \\ \pi_x(t)(u) = 1 & \text{ significa que } x(t) = u \text{ es posible}\end{aligned}$$

La flexibilidad en estas constricciones se modela haciendo $\pi_x(t)(u)$ entre cero y uno para algún valor de u .

La cantidad $\pi_x(t)(u)$ es el grado de posibilidad de que $x(t) = u$, y algunos valores de u serán más posibles que otros. Claramente, si U es el intervalo completo de $x(t)$, al menos uno de los elementos de U será totalmente posible como un valor de $x(t)$, así que:

$$\exists u, \pi_x(t)(u) = 1.$$

La distribución de posibilidades dinámica, si el único conocimiento accesible acerca de $x(t)$ es que $x(t)$ cae en $A(t)$, donde $A(t) \subseteq U(t)$, luego la distribución de posibilidades de $x(t)$ está definida de la forma:

$$\pi_x(t)(u) = \mu_A(t)(u) \quad (8.1)$$

donde $\mu_A(t)$ es la función característica del conjunto de posibles valores de $x(t)$, A . Cuando $\pi_x(t)$ toma valores en el intervalo $[0,1]$ el conjunto A es un conjunto borroso dinámico, y podemos decir que $x(t)$ está en A para expresar que el posible intervalo de valores de la variable A está restringido por el conjunto borroso dinámico A . En otras palabras la función de membresía sirve como una herramienta para representar las distribuciones dinámicas de posibilidades. La ecuación (8.1) no significa que la distribución dinámica de posibilidades sea la misma que la función de membresía, sino que establece que dado que sólo conocemos que $x(t)$ está en A , el grado de posibilidad de que

$$x(t) = u$$

es evaluado por la función de pertenencia $\mu_A(t)(u)$. Este modelo nos lleva a varias formas de conocimiento, específicamente:

Conocimiento completo:

$$A = \{u_0\}, \text{ para alguna } u(t), \text{ tal que } \pi_x(t)(u_0) = 1 \text{ y } \pi_x(t)(u) = 0 \forall u \neq u_0$$

Ignorancia dinámica completa:

$$A \stackrel{T}{=} U \text{ tal que } \pi_x(u) \stackrel{T}{=} 1, \forall u$$

La distribución dinámica de posibilidades cuantifica el conocimiento impreciso o incompleto. $x(t)$ es una variable simplemente valuada cuyo valor no es completamente conocido y sólo se conoce que $x(t)$ está en A , donde $A \neq (u) \forall u \in U$ en un tiempo t . Solamente en este caso se puede interpretar la función de membresía como una distribución dinámica de posibilidades.

Si $\pi_x(t)$ y $\pi_x'(t)$ son tales que $\pi_x(t) < \pi_x'(t)$, se dice que $\pi_x(t)$ es más específico que $\pi_x'(t)$ en el sentido de que no hay valores de u que puedan considerarse menos posibles para $x(t)$, de acuerdo con $\pi_x'(t)$, que para $\pi_x(t)$. Este concepto de especificidad subyace en la idea de que la distribución dinámica de posibilidad no es una propiedad asociada a una variable $x(t)$ por sí misma, sino una propiedad del instrumento de medida o estimación que nos da la información acerca de $x(t)$. En este sentido cualquier distribución de posibilidades $\pi_x(t)$ es provisional en un tiempo t , y puede mejorarse con información, cuando la disponible no es completa. Cuando $\pi_x(t) < \pi_x'(t)$, la información de $\pi_x'(t)$ es redundante y se puede despreciar. Estas características son típicas de la teoría dinámica de posibilidades y no tienen contraparte en la teoría de probabilidades.

En forma más general, cuando la información disponible proviene de varias fuentes que se pueden considerar como confiables, la distribución dinámica de posibilidades con la que contamos es la menos específica que satisface el conjunto de restricciones inducidas por las piezas de información dadas por las diferentes fuentes. Este es el *Principio dinámico de mínima especificidad*. Este tiene el significado siguiente: Dada la proposición $x(t)$ es A , luego cualquier distribución de posibilidades $\pi(t)$ tal que $\pi(t) < \mu_A(t)$ esta de acuerdo con $x(t)$ es A en un tiempo t . Por lo anterior si elegimos $\pi(t) < \mu_A(t)$ para representar nuestro conocimiento acerca de $x(t)$, seremos arbitrarios en la precisión. Luego la ecuación (8.1) se puede adoptar en forma natural y realmente incorpora el principio de mínima especificidad. Cuando varias distribuciones de posibilidades $\pi_x^1(t), \pi_x^2(t), \dots, \pi_x^n(t)$ son factibles para $x(t)$ desde diferentes fuentes, el principio de mínima especificidad cobra la forma:

$$\pi_x(t) = \min_{i=1, \dots, n} \pi_{x(t)}^i \quad (8.2)$$

como una consecuencia de las desigualdades:

$$\pi_x(t) \leq \pi_{x(t)}^i \quad i = 1, \dots, n$$

Nótese que la completa factibilidad de las fuentes considera que:

$$\sup_u \min_{i=1, \dots, n} \pi_{x(t)}^i(u) = 1$$

Por otro lado $\pi_{x(t)}(u) < 1, \forall u$. La subnormalización corresponde a una inconsistencia parcial.

Asumamos que existen dos variables $x(t)$ y $y(t)$ con distribuciones de posibilidades $\pi_x(t)$ y $\pi_y(t)$ respectivamente. El principio de mínima especificidad nos lleva a definir la distribución de posibilidades conjunta $\pi_{x,y}(t)$ como:

$$\pi_{x,y}(t) = \min^T(\pi_x(t), \pi_y(t)) \quad (8.3)$$

como una consecuencia de las desigualdades

$$\pi_{x(t),y(t)}(U, V) \leq \pi_{x(t)}(u), \forall v \quad y \quad \pi_{x(t),y(t)}(U, V) \leq \pi_{y(t)}(v), \forall u$$

Las últimas desigualdades no se toman en cuenta para ligas relacionales entre las variables $x(t)$ y $y(t)$, ya que ellas son válidas para toda v y toda u respectivamente. $\pi_{x,y}(t)$ se dice que es separable y las variables por lo mismo no son interactivas. Si existe una liga desconocida entre $x(t)$ y $y(t)$, la ecuación (8.3) proporciona cotas superiores o grados dinámicos de posibilidad, así las conclusiones deducidas de (8.3) son siempre correctas pero pueden tener poca información en un tiempo t . Desde luego que algunos valores (u, v) pueden considerarse como posibles debido a que no puede ser que:

$$\pi_{x,y}(t) < \min^T(\pi_x(t), \pi_y(t))$$

Este efecto de acotamiento superior esta de acuerdo con el principio de vinculación de Zadeh, que establece que si $x(t)$ está en A luego $x(t)$ está en B , lo que nos da que $A \subseteq B$ o $\mu_A(t) \leq \mu_B(t)$ en un tiempo t .

Dado un estado de conocimiento descrito en términos de distribuciones dinámicas de posibilidades articuladas $\pi_{x_1}(t), \dots, \pi_{x_n}(t)$, es posible obtener deducciones para un problema básico de la forma $(x_1(t), \dots, x_n(t)) \in B$, donde B es un subconjunto borroso ordinario. Debido a la imprecisión del conocimiento disponible, $(x_1(t), \dots, x_n(t)) \in B$ puede ser cierto, falso, o parcialmente conocido en un tiempo t . Se evalúa el estado de incertidumbre por medio de la medida de posibilidades dinámica y la medida de certeza.

La medida de posibilidad $\Pi(t)(B)$ está definida como:

$$\Pi(t)(B) = \sup_{(u_1, \dots, u_n)}^T \pi_{x_1, \dots, x_n}(t)(u_1, \dots, u_n) \quad (8.4)$$

que nos evalúa cual es el límite de extensión de B en el que es consistente con $\pi_{x_1}(t), \dots, \pi_{x_n}(t)$, esto es, ¿con cuál extensión debe de existir alguna eneáda en B que sea compatible con el conocimiento representado por $\pi_{x_1}(t), \dots, \pi_{x_n}(t)$?

La medida dinámica de certeza, también llamada necesidad, $N(B)$ está definida por:

$$N_T(B) = \inf_{(u_1, \dots, u_n) \notin B}^T [1 - \pi_{x_1(t), \dots, x_n(t)}(u_1, \dots, u_n)] \quad (8.5)$$

$N_T(B)$ evalúa con que extensión el conjunto borroso con función de membresía igual a $\pi_{x_1(t), \dots, x_n(t)}$ está contenido en B. La ecuación (8.5) sólo tiene sentido si $\pi_{x_1(t), \dots, x_n(t)}$ está normalizada, de otra forma se tiene que modificar la definición de certeza.

Las medidas dinámicas de posibilidades son conjuntos de funciones que obedecen los siguientes axiomas característicos:

$$\forall I(t), \Pi\left(\bigcup_{i \in I} \right)^T = \sup_{i \in I} \Pi(B_{i(t)}) \quad (8.6)$$

como las medidas de necesidad obedecen el axioma dual

$$\forall I(t), N\left(\bigcap_{i \in I} \right)^T = \inf_{i \in I} N(B_{i(t)}) \quad (8.7)$$

no se permite intercambiar Π y N en (8.6) y (8.7). Si $\pi_{x_1(t), \dots, x_n(t)} = \min_{i=1, \dots, n}^T \pi_{x_i(t)}$, como lo describe el principio dinámico de mínima especificidad, y si $\forall i = 1, \dots, n$, B_i es un subconjunto del universo U_i que es el intervalo de $x_i(t)$, luego tenemos:

$$\Pi(B_1 \times \dots \times B_n) = \min_{i=1, \dots, n}^T \Pi(B_i) = \Pi_i(B_i) = \min_{i=1, \dots, n}^T \Pi_i(B_i)$$

$$N(B_1 + \dots + B_n) = \max_{i=1, \dots, n}^T N(B_i) = N_i(B_i) = \max_{i=1, \dots, n}^T N_i(B_i)$$

donde $\times$ y $+$ denotan el producto cartesiano y el coproducto asociado:

$$(A + B) = A \times \overline{B}$$

donde la barra denota el complemento y $\Pi_i(t)$ y $N_i(t)$ son medidas dinámicas de posibilidad y necesidad basadas en la distribución de posibilidades normalizadas de $x_i(t)$ $\pi_{x_i(t)}$ definidas anteriormente, donde $\Pi(B_i(t))$ es la forma corta de:

$$\Pi(U_1 \times \dots \times B_i \times \dots \times B_n)$$

En forma más general siempre cuando $x_1(t), \dots, x_n(t)$ están fijados la distribución de posibilidades $\pi_{x(t)}$ se puede calcular como la proyección de $\pi_{x_1(t), \dots, x_n(t)}$ en U_i .

$$\pi_{x_i}(u_i) = \Pi(U_1 \times \dots \times U_{i-1} \times U_i \times U_{i+1} \times \dots \times U_n) = \sup_{U_1, \dots, U_{i-1}, U_{i+1}, \dots, U_n}^T \pi_{x_1(t), \dots, x_n(t)}(u_1, \dots, u_n) \quad (8.8)$$

y generalmente $\pi_{x_1(t), \dots, x_n(t)} \leq \min(\pi_{x_1(t)}, \dots, \pi_{x_n(t)})$.

Podemos ahora establecer los pasos de la metodología del razonamiento aproximado. Dado un conjunto de m distribuciones dinámicas de posibilidades $\pi_1(t), \dots, \pi_m(t)$ donde

$\pi_j(t)$ es la restricción borrosa del conjunto de variables $\{x_i(t) \mid i \in I_j\}$, sea x_k una variable cuyo valor se adquiere alrededor, con

$$k \in K = \bigcup_{j=1, \dots, m}^T I_j$$

La distribución dinámica de posibilidades $\pi_k(t)$ se define con base a dos pasos de Computación básicos:

1. Combinación: En virtud del principio dinámico de mínima especificidad, la distribución dinámica de posibilidades general será:

$$\pi(t) = \min_{j=1, \dots, m}^T \pi_j(t)$$

2. Proyección: Se calcula la proyección de $\pi(t)$ en el dominio U_k de $x_k(t)$ usando la ecuación (8.8).

En forma más general la proyección dinámica se puede hacer por el producto Cartesiano de dominios, si es que estamos interesados en encontrar las relaciones entre las variables correspondientes a estos dominios.

Este es el principio dinámico de combinación / proyección. Nótese que la proyección está también de acuerdo con el principio dinámico de mínima especificidad como en la ecuación (10), $\pi_{x_i}(t)(u_i)$ se calcula del valor de mayor posibilidad de la n-eada $(x_1(t), \dots, x_n(t))$ donde:

$$x_i(t) = u_j.$$

Por otra parte, cuando combinamos juntas las distribuciones de posibilidades $\pi_{x_1}(t), \dots, \pi_{x_m}(t)$ donde $m < j \leq n$, $\pi_{x_1}(t), \dots, \pi_{x_m}(t)$ se extiende al universo del discurso $U_1 \times \dots \times U_n$ por medio de extensiones cilíndricas por todos los valores,

$$\pi_{x_1}(t), \dots, \pi_{x_n}(t)(u_1, \dots, u_n) = \pi_{x_1}(t), \dots, \pi_{x_m}(t)(u_1, \dots, u_n),$$

Esta extensión cilíndrica está también de acuerdo con el principio de mínima especificidad, ya que éste produce la mayor distribución dinámica de posibilidades en $U_1 \times \dots \times U_n$, cuya proyección es $\pi_{x_1}(t), \dots, \pi_{x_n}(t)$.

8.2. VARIABLES LINGÜÍSTICAS

Una variable lingüística dinámica $x(t)$ es una variable que se ordena en un conjunto dinámico L de etiquetas simbólicas que usualmente corresponden a categorías en lenguajes naturales que cambian en el tiempo. Una etiqueta L en L puede ser el nombre de un conjunto borroso dinámico de un universo del discurso unidimensional U , a menudo una escala numérica. Sea $x(t)$ una variable cuyo rango es U . Luego por definición, la asignación $X = L \in L$ restringe los valores de $x(t)$ en el conjunto borroso L , por lo que define una distribución de posibilidades

$$\pi_x(t) = \pi_L(t)$$

Un conjunto de etiquetas $L = \{L_1, \dots, L_p\}$ corresponde a una descripción cualitativa de una escala numérica U . Esto es, el conjunto borroso con nombres en L forma una partición borrosa dinámica de U . Esto significa a lo menos que:

$$\forall u(t) \in U, \max_{i=1, \dots, p} \mu_{L_i}(u(t)) > 0 \text{ (condición dinámica de protección)}$$

$$\forall u(t) \in U \forall i \neq j, \min(\mu_{L_i}(t), \mu_{L_j}(u(t))) < 1 \text{ (exclusión dinámica mutua)}$$

Sin embargo a menudo se define una partición borrosa por medio de una condición fuerte:

$$\forall u(t), \sum_{i=1}^p \mu_{L_i}(u(t)) = 1$$

Sea F un conjunto borroso dinámico cuyas etiquetas están en L . El límite dinámico lingüístico $h(t)$ está modelado por medio de dos funciones: un modificador dinámico

$$m_h(t): [0,1] \rightarrow [0,1]$$

y un traductor, una función $q_h(t)$ tal que:

$$\mu_{hF}(t) = m_h(t) \circ \mu_F(t) \circ q_h(t)$$

donde $\circ$ denota la composición de funciones.

Etiquetas como mucho, algunas veces, corresponden a incrementar la precisión

$$m_{\text{mucho}}(t) \circ \mu_F(t) \circ q_{\text{mucho}}(t) \leq \mu_F(t) \text{ y } q_{\text{mucho}}(t) = \text{identidad}$$

o una traslación en U :

$$m_{\text{mucho}}(t) = \text{identidad y } q_{\text{mucho}}(u(t)) = a + u(t)$$

Cuando el universo U está acotado a una escala numérica $[a,b]$ se pueden definir nociones como antónimos en L en términos de simetría con respecto al punto medio de U , llamémoslo:

$$\mu_{\text{ant}(F)}(u(t)) = \mu_F(t)(b + a - u(t))$$

En particular si $[a,b] = [0,1]$, $\mu_{\text{ant}(F)}(u(t)) = \mu_F(t)(1 - u(t))$. Por ejemplo pequeño ahora se puede definir como el antónimo de grande ahora.

Los conectivos lógicos dinámicos como “y” y “o” se definen por medio de operaciones teóricas de conjuntos borrosos dinámicos, especialmente normas t dinámicas para la conjunción y conormas t dinámicas para la disyunción. Con el problema de las variables

lingüísticas dinámicas está relacionado el problema de la aproximación lingüística dinámica. El problema consiste en que dado un conjunto borroso dinámico F en U , y un conjunto dinámico de etiquetas L perteneciente a U , hay que encontrar la mejor etiqueta en L para F . En lo anterior las variables son supuestas como simplemente valuadas, y las etiquetas relacionadas a una variable $x(t)$, corresponden a un conjunto de valores mutuamente exclusivos de $x(t)$. Existen otros tipos de variables que pueden tomar varios valores simultáneamente, y son llamadas conjuntivas en oposición a las simplemente valuadas que son llamadas disyuntivas. El conocimiento incompleto acerca de una variable conjuntiva $x(t)$ se puede modelar por medio de una distribución dinámica de posibilidades en el conjunto potencia de U , por ejemplo, $\pi_x(t)(A)$ es el grado de posibilidad de que el valor de $x(t)$ sea exactamente el elemento de A . Cuando el conocimiento es completo

$$\exists B \subseteq U, \text{ tal que } \pi_x(t)(A) = 1, \text{ si } A = B \text{ y } \pi_x(t)(A) = 0, \text{ si } A \neq B$$

Es claro que en este caso $\mu_B(t)$ no es una distribución dinámica de posibilidades, a sólo que B sea un conjunto borroso dinámico.

Los problemas con variables dinámicas conjuntivas son esencialmente debidos a la imposibilidad práctica de almacenar una distribución dinámica de posibilidades sobre un conjunto dinámico potencia en una computadora. Es muy importante ver herramientas eficientes de representación para representar a $\pi_x(t)$ o al menos encontrar representaciones aproximadas. El conocimiento dinámico conjuntivo incompleto puede ser eficientemente representado por:

1. Pares de subconjuntos (B^*, B^*) de U donde B^* es un conjunto de valores certeramente tomados por $x(t)$. B^* es un conjunto de valores no tomados por $x(t)$ en un tiempo dado.
2. Cuantificadores que describen la cardinalidad de un conjunto de valores formados por una variable dinámica conjuntiva.

8.3. VARIABLES LINGÜÍSTICAS COMPUESTAS Y COMPARADORES BORROSOS

Volveremos a la representación de las constricciones para ligar varias variables. En la literatura se han considerado tres clases de constricciones: compuestas, variables lingüísticas, comparadores borrosos y reglas si ... entonces borrosas.

Las variables lingüísticas dinámicas compuestas corresponden a categorías borrosas dinámicas complejas que involucran varias dimensiones que cambian en el tiempo y pueden expresarse a través de combinaciones jerárquicas de variables lingüísticas dinámicas simples, relacionadas vía conectivos borrosos dinámicos que puede ser operaciones lógicas.

Otro tipo muy común de relaciones borrosas dinámicas son los comparadores borrosos dinámicos que expresan igualdades dinámicas aproximadas o desigualdades dinámicas fuertes. Una igualdad aproximada E es una relación dinámica binaria borrosa que es reflexiva,

y simétrica

$$\mu_E(u(t), v(t)) = \mu_E(v(t), u(t))$$

$$\mu_E(u(t), v(t)) = \mu_E(v(t), u(t)).$$

Una manera simple de modelar una igualdad dinámica aproximada entre dos variables $x(t)$ y $y(t)$ es usar un índice de distancia $d(t)$ y un intervalo borroso dinámico M tal que

$$\mu_M(t)(0) = 1$$

y $\mu_M(t)$ no se incrementa en $[0, +\infty]$. Así definimos:

$$\pi_{x,y}(u(t)) = \mu_E(u, v) = \mu_M(d(u, v)). \quad (8.9)$$

Nótese que ya que $d(u(t), v(t)) \geq 0$ sólo es útil la parte positiva de M .

Usando $d(u(t), v(t)) = |u - v|$ nos permite comparar órdenes de magnitud de variables numéricas. Volteando M en sus complementos:

$$\overline{M}(\mu_{\overline{M}} = 1 - \mu_M)$$

en la ecuación (14), modelamos una expresión borrosa que expresa que $x(t)$ difiere de $y(t)$.

Una relación borrosa dinámica R tal que:

$$\mu_R(u(t), v(t)) = 1 - \mu_M(\max(0, u(t) - v(t)))$$

modela más o menos formas de la desigualdad

$$x(t) > y(t)$$

donde

$$\pi_{x,y}(u(t), v(t)) = \mu_R(u(t), v(t)).$$

Los comparadores borrosos dinámicos de órdenes de magnitud relativos se pueden obtener remplazando $d(u(t), v(t))$, por $u(t)/v(t)$ y escogiendo M tal que

$$\mu_M(t) = 1 \text{ y } \mu_M(u(t)) = \mu_M(1/u(t)),$$

en la ecuación (8.9), asumiendo $u(t) \geq 0$ y $v(t) \geq 0$. Luego modelamos una relación borrosa dinámica de vecindad V_0 . El radio en el cual u está en la vecindad de v es:

$$\mu_{V_0}(u(t), v(t)) = \mu_M(u(t)/v(t))$$

De manera similar una relación de insignificancia es de la forma:

$$\mu_{N_e}(u(t), v(t)) = \mu_{V_0}^T(u(t) + v(t), v(t)) = \mu_M^T(1 + u(t) / v(t))$$

que expresa que $u(t)$ es insignificante en frente de $v(t)$ si y sólo si $u(t) + v(t)$ está en la vecindad de $v(t)$.

8.4. DECLARACIONES CONDICIONALES BORROSAS

La más extensa de las variedades de declaraciones borrosas dinámicas que involucran varias variables, es la regla “si... luego” con predicados borrosos dinámicos. Llamaremos una regla dinámica elemental de la forma “Si $x(t) \in A$ luego $(t) \in B$ ”, corresponde a una distribución dinámica de posibilidades de la forma:

$$\pi_{x(t), y(t)}(u(t), v(t)) = \mu_{A \rightarrow B}^T(u(t), v(t)) = I(\mu_A u(t), \mu_B v(t))$$

donde I es una implicación en lógica multivaluada. Ha habido gran discusión alrededor de la elección de la operación implicación I . Varios puntos de vista se pueden adoptar para tomar esta elección: un punto de vista puramente algebraico, una aproximación de deducción orientada, y una investigación dentro de la semántica de las reglas borrosas. Nosotros tomaremos el punto de vista puramente algebraico.

8.5. REPRESENTACIÓN DE LA IMPLICACIÓN DINÁMICA MULTIVALUADA

Existen básicamente tres clases de operaciones de implicación dinámica:

1. Las basadas en la implicación $(p \rightarrow q)$ se define como $(\neg p \wedge q)$, y es de la forma

$$I(a(t), b(t)) = S(n(a(t)), b(t)) \text{ (S-implicaciones)}$$

donde S es una conorma dinámica para disyunciones multivaluadas y n es una negación fuerte.

2. Aquellas basadas en la idea de que la implicación refleja un ordenamiento parcial de las proposiciones. La principal clase de estas implicaciones esta basada en un concepto de residuación en las estructuras de redes con una operación de semigrupo T que puede ser una conjunción

$$I(a(t), b(t)) = \sup \{c \in [0, 1] \mid T(a(t), c(t)) \leq b(t)\} \text{ (R implicaciones)},$$

donde T es una norma dinámica triangular, y R es un residuador. Otra implicación como:

$$\begin{cases} I(a(t))^T = 1 & \text{si } a(t) \leq^T b(t) \\ 0 & \text{otro valor} \end{cases}$$

pertenece a esta familia sin derivarse de la residuación.

3. Las basadas en el punto de vista borroso dinámico de si ...en t luego ...en t entonces, que nos llevan a la forma:

$$I(a(t), b(t))^T = S(n(a(t)), T(a(t), b(t))) \quad (\text{Implicación QL})$$

donde S es la conorma dinámica, n es una negación fuerte, y T es el dual-n de S. Esta forma aparece con un par de reglas borrosas:

$$\begin{aligned} &\text{Si } x(t) \in^T A \text{ luego } y(t) \in^T B \\ &\text{por lo tanto si } x(t) \notin^T A \text{ luego } y(t) \notin^T C \end{aligned}$$

se interpreta directamente como una restricción borrosa en (x(t), y(t)) de la forma:

$$(x(t), y(t))^T \in (A \times B) \vee (x(t), y(t))^T \in (\neg A \times C). \quad (8.10)$$

La interpretación de las reglas borrosa vía la ecuación (8.10) se ha adoptado en control borroso como un punto de vista universal.

Desde un punto de vista axiomático, se requieren las propiedades siguientes para una operación I de 2 plazas que represente una operación de implicación:

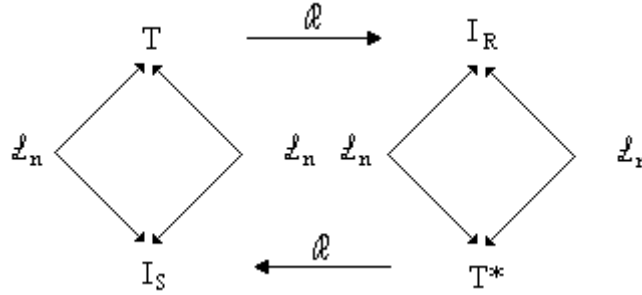
- I₁. si $a(t) \leq^T a'(t)$ luego $I(a(t), b(t))^T \geq I(a'(t), b(t))^T$
- I₂. si $b(t) \geq^T b'(t)$ luego $I(a(t), b(t))^T \geq I(a(t), b'(t))^T$
- I₃. $I(0, b(t))^T = 1$ (la falsedad implica cualquier cosa en un tiempo t)
- I₄. $I(1, b(t))^T = b(t)$ (la tautología no justifica algo en un tiempo t)
- I₅. $I(a(t), b(t))^T \geq b(t)$, es la contraparte numérica de $(t) \rightarrow (p(t) \rightarrow q(t))$.
- I₆. $I(a(t), a(t))^T = 1$, principio dinámico de identidad.
- I₇. $I(a(t), I(b(t), c(t)))^T = I(b(t), I(a(t), c(t)))^T$, principio dinámico de intercambio.
- I₈. $I(a(t), b(t))^T = 1$ si y sólo si $a(t) \leq^T b(t)$. La implicación define un orden.
- I₉. $I(a(t), b(t))^T = I(n(b(t)), n(a(t)))^T$, para alguna negación fuerte n. Ley de contraposición.

Sea T una norma triangular, y sea L_n y R_n la transformación de T en una implicación tipo S y R respectivamente.

$$\begin{aligned} I_S(t) &= L_n(T(t)) \text{ con } I_S(a(t), b(t))^T = n(T(a(t), n(b(t)))) \text{ y } I_R = R(T(t)) \text{ tal que } I_R(a(t), \\ &b(t))^T = \sup \{c \in [0, 1], T(a(t)c(t)) \leq^T b(t)\} \end{aligned}$$

Notando que $L_n \circ L_n(T) = T$ y que $L_n \circ R_n(T)$, produce una operación de conjunción T* que no es más grande en general que una norma dinámica triangular. Se puede

probar que el siguiente diagrama conmuta cuando T es continua y Arquimedean, o la operación mínima:



Esto es que

$$T = L_n \circ R \circ L_n \circ R(T),$$

por lo tanto si

$$L_n(I_R(a(t), b(t))) = T^*(a(t), b(t))$$

luego

$$L_n(I_R(n(b(t)), n(a(t)))) = T^*(b(t), a(t)) \neq T^*(a(t), b(t)).$$

La operación

$$T^* = L_n \circ R(T),$$

se puede llamar una pseudoconjunción dinámica asociada a T .

8.6. IMPLICACIONES VALUADAS EN UN INTERVALO

Turksen notó que las formas conjuntivas y disyuntivas normales no coinciden en lógicas multivaluadas, cuando se usan para calcular la $I(a(t), b(t))$ en un intervalo. El problema es estar seguro de que el intervalo es lo suficientemente angosto para ser informativo y suficientemente ancho para cubrir varias implicaciones. En el caso de las formas disyuntivas dinámicas (FDD) y conjuntivas dinámicas (FCD) normales se tiene:

$$\begin{aligned} \text{(FCD)} \quad (p(t) \rightarrow q(t)) &= \neg p(t) \vee q(t) \\ \text{(FDD)} \quad (p(t) \rightarrow q(t)) &= (p(t) \wedge q(t)) \vee (\neg p(t) \wedge q(t)) \vee (\neg p(t) \wedge \neg q(t)) \end{aligned}$$

lo que nos lleva a las siguientes contrapartes multivaluadas para la implicación dinámica:

- Una implicación dinámica tipo S para (FCD) $(a(t) \rightarrow b(t))$
- $(\text{FDD}) \quad (a(t) \rightarrow b(t)) = S(T(a(t), b(t)), T(n(a(t), b(t)), T(n(a(t)), n(b(t))))$

La motivación principal de la aproximación del intervalo valuado es que generalmente:

$$(\text{FDD}) \quad (a(t) \rightarrow b(t)) \leq (\text{FCD}) \quad (a(t) \rightarrow b(t))$$

para muchas familias de normas y conormas triangulares que son n -duales.

También los presupuestos si ... luego... se pueden expresar por implicaciones valuadas en intervalos. Un modelo complicado para si $x(t)$ es A luego y es B borroso dinámico, se hace por medio de reglas que pueden expresarse por medio de una clase especial de conjuntos borrosos dinámicos del tipo 3. Consideremos un conjunto borroso en $U \times V$ con grados de membresía que son descritos como un conjunto borroso de intervalos valuados en el tiempo. Particularmente consideramos cantidades de la forma $\pi_{B/A}(u(t), \pi(v(t)))$ en lugar de $\pi_{y/x}(u(t), v(t))$, donde $\pi_{B/A}$ depende de A y B y es valuada en un intervalo y $\pi(v(t))$ es el grado de posibilidad de que $y(t) = v(t)$.

8.7. RESULTADOS EN CONDICIONAMIENTO

Otros resultados significativos del modelado de las reglas si ... luego ..., es la distinción entre la condicional y la implicación material. Como es conocido en teoría de la probabilidad $P(B|A) \neq P(\neg A \vee B)$, y que una probabilidad asociada con una declaración condicional refiere más a menudo a una probabilidad que a una condicional lógica $A \xrightarrow{T} B = \neg A \vee B$. Sea $\pi_{x(t), y(t)}$ una distribución de posibilidades conjunta. Una condición natural para la definición de una distribución de posibilidades condicional, deriva de los siguientes requerimientos:

$$\begin{aligned} & \text{Posibilidad } (x(t) \stackrel{T}{=} u(t) \wedge y(t) \stackrel{T}{=} v(t)) \stackrel{T}{=} \text{Posibilidad } (y(t) \stackrel{T}{=} v(t)) \\ & \text{dado} \\ & x \stackrel{T}{\bullet} u(t) \stackrel{T}{\bullet} \text{Posibilidad } (x(t) \stackrel{T}{=} u(t)) \end{aligned} \quad (8.11)$$

Si la ecuación (8.11) es válida para medidas de posibilidad dinámica, el resultado en ecuaciones funcionales forzan a

$$\begin{aligned} & \bullet \text{ a ser un producto si } \bullet \text{ es estrictamente monotónica} \\ & ((a(t) \stackrel{T}{<} b(t) \text{ y } c(t) \stackrel{T}{\geq} d(t))) \text{ o } (a(t) \stackrel{T}{\geq} b(t) \text{ y } c(t) \stackrel{T}{>} d(t)) \text{ implica } a(t) \stackrel{T}{\bullet} b(t) \stackrel{T}{>} c(t) \stackrel{T}{\bullet} d(t). \end{aligned}$$

Por lo tanto cuando $\bullet$ es estrictamente isotónico

$$(a(t) \stackrel{T}{>} b(t) \text{ y } c(t) \stackrel{T}{>} d(t)) \rightarrow a(t) \stackrel{T}{\bullet} b(t) \stackrel{T}{>} c(t) \stackrel{T}{\bullet} d(t),$$

luego la operación dinámica mínima también se permite. Como una consecuencia se define la distribución condicional dinámica de posibilidades por:

$$\pi_{x(t), y(t)}(u(t), v(t)) \stackrel{T}{=} \pi_{x(t)/y(t)}(u(t), v(t)) * \pi_{x(t)}(u(t))$$

donde $\pi_{x(t)}$ es la proyección de $\pi_{x(t), y(t)}$ con $*$ = producto o mínimo. Cuando $\bullet =$ producto, $\pi_{x(t)/y(t)}$ corresponde a una distribución dinámica condicional, en términos de una distribución dinámica de posibilidades valuada por intervalos.

$$\pi_{y(t)/x(t)}^H(u(t), v(t)) = \pi_{x(t)y(t)}^T(u(t), v(t)) \text{ si } \pi_{x(t)y(t)}^T(u(t), v(t)) < \pi_{x(t)}^T(u(t)) = [\pi_{x(t)y(t)}^T(u(t), v(t)), 1] \text{ si } \pi_{x(t)y(t)}^T(u(t), v(t)) = \pi_{x(t)}^T(u(t)) \quad (8.12)$$

La ecuación (8.12) expande el intervalo dinámico de soluciones admisibles de punto valuado. Para poder seleccionar una distribución dinámica de posibilidades de la ecuación (8.12), debemos notar que $\pi_{x(t)/y(t)}$ debe de ser normal, más específicamente, $\forall u, \exists v, \pi_{x(t)y(t)}(v(t), u(t)) = 1$ ya que por definición $\pi_{x(t)/y(t)}$ produce una distribución de posibilidades en $y(t)$ cuando se conoce $x(t)$. Siempre es posible satisfacer este requerimiento debido a que:

$$\sup_u \pi_{x(t), y(t)}^T(u(t), v(t)) = \pi_x^T(u(t)).$$

Se puede seleccionar_(t)

$$v_u \in V(u) = \{v \mid \pi_{x(t), y(t)}^T(u, v) = \pi_x^T(u)\}$$

y definir

$$\forall u, (\pi_{x(t)/y(t)} u, v) = \begin{cases} \pi_{x(t), y(t)}^T & \text{si } v \neq V_u \\ 1 & \text{en otro lado} \end{cases} \quad (8.13)$$

Esta definición tiene la ventaja de continuidad con respecto a las variaciones de $\pi_{x(t), y(t)}$ y satisface los requisitos mínimos de información de distancia. La elección de v_u permanece arbitraria de forma tal que pueden ser candidatas varias distribuciones de posibilidades.

Otra aproximación es tomar el principio dinámico del mínimo de la ecuación (8.12) y escoger:

$$(\pi_{x(t)/y(t)} u, v) = \begin{cases} \pi_{x(t), y(t)}^T & \text{si } v \notin V_u \\ 1 & \text{en otro lado} \end{cases} \quad (8.14)$$

La pregunta ahora es como figurar cuando una regla borrosa del tipo “si $x(t)$ es A , luego y es B ” nos lleva a una distribución condicional de posibilidades en el sentido de las ecuaciones (8.11 a 24). ¿Dada una representación de una declaración condicional del tipo $\mu_A \rightarrow \mu_B$ y una operación de combinación $*$, es $\mu_A * \mu_A \rightarrow_B^T$ una distribución de posibilidades adjunta? Esto nos lleva a la identificación de una π_x con μ_A y $\pi_{x, y}$ con $\mu_A * \mu_A \rightarrow_B^T$ en (8.11) y sugiere que μ_A es la proyección de $\mu_A * \mu_A \rightarrow_B^T$ en el intervalo de x .

$$\sup_u \mu_A(u(t)) * I(\mu_A(u(t)), \mu_B(v(t))) = \mu_A(u(t))$$

La traslación dinámica de declaraciones que son explícitamente calificadas como parcialmente ciertas o no completamente seguras, ya que cambian con el tiempo. En esta sección se revisa la representación de este conocimiento, considerando a las declaraciones borrosas dinámicas como aquellas en las cuales un cuantificador dinámico está implícito, y llamaremos disposiciones dinámicas.

8.8. CALIFICACIÓN DE VERDAD

Una declaración sólo puede ser cierta con respecto a otra declaración en un tiempo definido, lo nos lleva a tratar con declaraciones parcialmente ciertas reconstruyendo las declaraciones verdaderas que subyacen en ella en un tiempo dado, de acuerdo con la equivalencia siguiente:

$$“x \text{ es } A \text{ es } \tau\text{-cierta} \Leftrightarrow \exists B, “x \text{ es cierta y } \mu_B = f(\mu_\tau, \mu_A)”$$

donde τ es un valor dinámico borroso de verdad cuyo significado como “al menos cierto hoy”, “no muy cierto mañana”, etc.

Para ser más precisos es necesario definir el grado de verdad de una declaración borrosa “ $x(t)$ es A ” dado que “ $x(t)$ es B ” en este tiempo. Cuando $B = \{u_0\}$ el valor de verdad de “ $x(t)$ es A ” es simplemente $\mu_A(u_0)$. En forma más general el grado de verdad de “ $x(t)$ es A ” será el conjunto borroso dinámico τ definido vía el principio de extensión:

$$\mu_\tau(t) = \begin{cases} \sup \left\{ \mu_B(u(t)) \mid \mu_A(u(t)) = r \right\} & \text{si } \mu_A^{-1}(t) \neq \emptyset \\ 0 & \text{si } \mu_A^{-1}(t) = \emptyset \end{cases} \quad (8.15)$$

τ es igual a $\mu_A(B)$, el grado de membresía definido en el conjunto borroso dinámico A , de un elemento $x(t)$ imprecisamente localizado, restringido por $\mu_{B(t)}$. La función de membresía de $\mu_A(B)$ se denotará por $\mu_{A|B}$ por simplicidad. La ecuación (8.15) sugiere que si “ $x(t)$ es A ”, $\mu_A(t)(B)$ es cierto dado que “ $x(t)$ es B ” es cierto. Inversamente, dado que “ $x(t)$ es A ” es τ -cierto”, el conjunto borroso dinámico B , al que “ $x(t)$ es A es τ -cierto, dado que $x(t)$ es B ”, es una solución de la ecuación (26), donde son dados τ y μ_A . El principio de mínima especificidad nos lleva a considerar la gran solución:

$$\mu_B(u(t)) = \mu_\tau(\mu_A(u(t))) \quad (8.16)$$

La ecuación (8.16) permite que los conjuntos borrosos dinámicos de $[0,1]$ sea interpretados en términos de valores de verdad:

- “ $x(t)$ es cierto” debe de ser equivalente a “ $x(t)$ es A ” de forma tal que el conjunto borroso de $[0,1]$ con función de membresía $\mu(t) = r$ pueda ser llamado *cierto en un tiempo t* .
- “ $x(t)$ es falso” debe de ser equivalente a “ $x(t)$ es $\bar{A}$ ” con $\mu_{\bar{A}} = 1 - \mu_A$ y luego el conjunto borroso $[0,1]$ con función de membresía $\mu(t) = 1 - t$ se llama *falso en un tiempo t* .
- “ $x(t)$ es A es desconocido” debe de ser equivalente a “ $x(t)$ es U ” donde U es el dominio de $x(t)$ y luego, el conjunto $[0,1]$ corresponde directamente al caso de un valor desconocido de verdad.

Los valores usuales de 1 y 0 se recuperan cuando A y B son claros, digamos:

$$\begin{array}{ll} \text{Si } B \subseteq A, \text{ luego } \mu_A(B) = \{1\} & \text{(cierto)} \\ \text{Si } B \subseteq \bar{A} \text{ luego } \mu_A(B) = \{0\} & \text{(falso)} \\ \text{mientras que si } B \cap A \neq \emptyset \text{ y } B \cap \bar{A} \neq \emptyset, \mu_A(B) = \{0,1\} & \text{(desconocido)} \end{array}$$

Si sólo B es borroso dinámico, $\mu_{A(t)}(B)$ tiene su soporte en $\{0,1\}$, así que si A no es borroso dinámico hace que la declaración “x(t) es A es τ -cierto” no tenga sentido, ya que τ es un subconjunto borroso dinámico de $[0,1]$ que representa un valor de verdad intermedio, a menos que interpretemos “x(t) es A” como “x(t) no está tan lejano de A”. Si A es claro, “no tan lejano en este momento” es borroso dinámico.

Similarmemente la ecuación (8.16) no se puede aplicar a un conjunto borroso cuando τ es un valor intermedio preciso.

$$\tau = \{t\}, t \in [0,1],$$

luego la ecuación (8.16) puede obtener un valor de x preciso a ser recuperado,

$$x \in \{u \mid \mu_A(u) = t\},$$

que algunas veces se reduce a un simple valor.

La paradoja producida por valores precisos de verdad asociados a proposiciones borrosas dinámicas se resuelve en la práctica por el hecho de que nadie está en la posibilidad de volver a proveer un valor borroso dinámico. Por ejemplo esto es medio cierto ahora se puede traducir en una forma borrosa dinámica por $r = 0.5$. En forma más general nos permite interpretar declaraciones como: “x(t) es A, es r-cierto” $t \in [0,1]$, usando un valor borroso de verdad $\bar{r}$, tal que $\mu_i(t) = 1$, μ_i es continua como:

$$\begin{array}{ll} \mu_i(t) = r'/r & \text{si } 0 \leq r' \leq r \\ \mu_i(t) = (1 - r')/(1 - r) & \text{si } 1 \geq r' \geq r \end{array}$$

Esta familia interpola entre los valores continuos de verdad “cierto” y “falso”. “x(t) es A, es r-cierto” indica que los valores más pausibles de x(t) son aquellos que $\mu_A(x(t))$ es grande, o “es cierto que $\pi_x = \mu_A$ ”. La elección entre un significado y otro es cosa de contexto. Esta ambigüedad persiste cuando se cambia verdad por τ -verdad, donde τ puede ser un valor intermedio. En ambos casos el modelo determina de que la declaración dinámica cualificada de verdad, es equivalente a la declaración borrosa dinámica

$$\text{“x(t) es B” con } \mu_B = \mu_\tau \circ \mu_A.$$

El problema crucial es la traducción de los cualificadores de verdad del lenguaje natural en conjuntos borrosos dinámicos adecuados de un intervalo unitario. La interpretación que

$$\pi_x^T = \mu_A^T \text{ de "x(t) es A es cierto en un tiempo t"}$$

nos permite ver a “cierto en un tiempo dado” como un valor borroso dinámico. Sobre los valores de verdad diferentes de verdad, tendemos a olvidar los valores borrosos dinámicos de verdad tales como el soporte de B, con

$$\mu_B^T = \mu_\tau^T \circ \mu_A^T,$$

que estrictamente caen en el soporte de A que son valores precisos. La otra interpretación $\mu_A(x(t))$ es alto hoy se refiere a valores de verdad que especifican a A por restringir los valores de $x(t)$ a niveles de corte de A, por ejemplo.

Sea $B \subseteq^T A$. Esto es lo que pasa cuando se llaman los valores de verdad como muy cierto ayer, con $\mu_{\text{muy cierto}}(r(t)) = r^2(t)$, $\forall r \in [0,1]$. Luego por definición “x(t) es A es muy cierto” es equivalente a “x(t) es muy A”. En lenguaje natural no es claro que expresiones tales como “muy cierto ahora”, “absolutamente cierto mañana”, que a menudo se interpretan como $\tau = 1$ en la literatura borrosa, tienen significados por sí mismos si consideramos que nada es más cierto que falso en un tiempo t.

Otros ejemplos de valores borrosos dinámicos de verdad consideran el dual de la ecuación (26) cambiando sup en inf y usando 1 en lugar de 0 cuando

$$\mu_A^{-1}(r) = \phi,$$

por ejemplo el complemento del valor borroso de verdad de “x(t) es A” dado que “x(t) no es B”. Estos dos valores de verdad están de acuerdo con cualquier otro, cuando μ_A está dentro de la inclusión:

$$\bar{\mu}_A(B) \subseteq^T \mu_A(B)$$

Otra cantidad de interés es $\mu_{\bar{A}}(B)$ que es el antónimo de $\mu_A(B)$, en la escala [0,1].

Sea:

$$\mu_{\bar{A}|B}(r) = \mu_{A|B}^T(1-r)$$

Este antónimo expresa el hecho de que dado “x(t) es b”, “x(t) es A es τ ” es equivalente a “x(t) no es A es el ant(τ)”.

8.9. CALIFICACIÓN DE CERTIDUMBRE Y POSIBILIDAD. EL CASO PRECISO.

Consideremos como Zadeh nociones locales tanto a la verdad como la certeza dinámicas y la posibilidad. Dado que “x(t) es B”, por ejemplo

$$\pi_{x(t)}^T = \mu_B$$

la posibilidad dinámica (respecto a la certeza) de la declaración borrosa dinámica “x(t) es A” refleja la consistencia (respecto a la vinculación lógica dinámica) de A con B. Cuando A es preciso estas nociones son capturadas por los grados de posibilidad $\Pi(A)$ y de necesidad $N(A)$ obtenidos de $\pi_x = \mu_B$.

A la inversa, si una declaración incierta de la forma “x(t) es A es cierta en grado α ”, donde A es precisa, si encontramos la declaración verdadera que subyace “x(t) es B” nos lleva a resolver para μ_B la ecuación funcional:

$$N(A) = \inf_{u \notin A}^T [1 - \mu_B(u(t))]^T = \alpha \quad (8.17)$$

la cual por el principio de mínima especificidad nos lleva a:

$$\mu_B(u(t)) = \max^T(\mu_A(u(t)), 1 - \alpha) \quad (8.18)$$

Nótese que en (8.17) la igualdad se puede cambiar en una desigualdad $N(A) \geq \alpha$ sin que se cambie el resultado, así que la declaración dinámica cualificadora de verdad se debe de interpretar como “x(t) es A es a lo menos cierta con grado de verdad α en un tiempo t”. Cuando $\alpha = 1$ (Certeza total) la ecuación (28) nos lleva a que $A = B$, y cuando $\alpha = 0$, la ecuación (8.18) nos lleva a que $B = U$ (total ignorancia).

En el caso de declaraciones de imposibilidad cualificada como “x(t) es A es al menos posible en grado β ” es equivalente a decir “x(t) no es A es a lo menos posible en grado $1 - \beta$ ” debido a que

$$N(\overline{A})^T = 1 - \Pi(A)$$

Más difícil es tratar declaraciones del tipo “x(t) es A es al menos β posible”. Para esto debemos de resolver la ecuación:

$$\Pi(A) = \sup_{u \in A}^T \mu_B(u(t))^T \geq \beta \quad (8.19)$$

Si aplicamos el principio dinámico de mínima especificidad llegamos a nada, pero $\mu_B(u) = 1 \forall u$, la información no informante nos dice algo no trivial.

Para salirnos de este problema tenemos dos soluciones:

- Interpretamos la declaración como una *meta distribución dinámica de posibilidades*, que será el conjunto:

$$A_\beta^* = \left\{ \pi \left| \sup_{u \in A}^T \pi(u(t)) \geq \beta \right. \right\}$$

- Interpretando “x(t) es A es al menos posible en grado β en el tiempo t” como “todos los elementos de A son valores posibles de x(t), al menos con grado β en el tiempo t”, lo que nos lleva a establecer en x la siguiente restricción:

$$\forall u(t), \pi_{x(t)}(u(t)) \geq \min(\mu_A(u(t)), \beta) \quad (8.20)$$

Nótese que (8.20) implica (8.19) pero no inversamente, así que cuando $\pi_{x(t)}$ satisface (31), es adentro de A_β^π . La ecuación (8.30) se puede escribir en forma de conjunto valuado como

$$\mu_B(u) = [\min(\mu_A(u), \beta), 1],$$

donde B es un subconjunto borroso dinámico de intervalo valuado. Si las ecuaciones (8.30) y (8.19) se usan para interpretar la declaración de incertidumbre precisa “x(t) es A es exactamente posible en grado α en un tiempo t”, vista como “x(t) es A es al menos posible en grado α en un tiempo t”, se obtiene una declaración equivalente “x(t) es B $^\alpha$ ” tal que:

$$\mu_{B^\alpha}(u(t)) = \min(\alpha, \mu_A(u(t)), \max(\alpha, 1 - \mu_A(u(t)))$$

8.10. REGLAS DE CUALIFICACIÓN DE VERDAD

Cuando A y B son conjuntos borrosos, “si x(t) es A luego y es B” tiene el significado “las más de las x(t) son A, las más de las y = f(x(t)) son B”, que describe el hecho de que hay una función del supA al supB tal que la imagen de A a través de f esté contenida en

$$B(f(A)) \subseteq B,$$

que se lee como:

$$\forall u, \mu_A(u) \leq \mu_B(f(u))$$

que ha sido propuesta como una función borrosa dinámica de un conjunto borroso dinámico A a un conjunto borroso dinámico B. Esta ecuación, introduciendo la norma triangular T:

$$\forall u(t), \sup [\mu_A(u(t)) \mid f(u(t)) = v(t)] \leq \mu_B(v(t))$$

8.11. REGLAS DE CUALIFICACIÓN DE INCERTIDUMBRE

Existen situaciones cuando la regla borrosa dinámica significa “las más de las x(t) $\in$ A y lo más confiable es que y(t) $\in$ B”. Una traslación simple de este tipo de reglas es la desigualdad:

$$\forall u, \mu_A(u(t)) \leq g_{u(t)}^T(B)$$

donde $g_{u(t)}(B)$ evalúa el grado de confianza de la declaración “y es B” cuando $x = u$. La función $g_{u(t)}$ puede ser cualquier clase de medida de incertidumbre acorde con la confianza expresada en la regla.

La interpretación semántica de las reglas borrosas posibilita a seleccionar una función de implicación de acuerdo con el significado de la regla. Para algunas de estas reglas, el principio de mínima especificidad es suficiente para derivar la distribución de posibilidades que subyace bajo la regla.

8.12. EL MODUS PONENS GENERALIZADO

El modus ponens generalizado es una forma de inferencia en el que se involucran predicados borrosos dinámicos. En una forma simple se puede plantear de la siguiente forma:

- De la regla: “si $x(t)$ es A luego $y(t)$ es B”
- $y(t)$ el hecho: “ $x(t)$ es A”
- se puede inferir que: “ $y(t)$ es B”

donde B' se calcula de A' , A y B. En la práctica la regla se expresa como una fórmula lógica, por medio de la operación f que es generalmente una implicación, y que se interpreta en el sentido semántico como una distribución condicional dinámica de posibilidades

$$\pi_{y(t)|x(t)} = f(\mu_{A(t)}, \mu_{B(t)}).$$

Luego B' se obtiene por medio del principio de combinación/proyección bajo la forma:

$$\mu_{B'}(v(t)) = \sup_u^T \min^T [\mu_{A'}(u(t)), f(\mu_A(u(t)), \mu_B(v(t)))]$$

donde $f(a(t), b(t)) = I(a(t), b(t))$, implicación o $T(a(t), b(t))$, una norma triangular.

Para seleccionar la operación f , de forma tal que B' satisface intuitivamente propiedades satisfactorias de acuerdo con la semejanza entre A y A' , se han sugerido los siguientes axiomas:

1. $B' \supseteq^T B$
2. $A_1' \supseteq^T A_2' \Rightarrow B_1' \supseteq^T B_2'$ (monotonicidad)
3. $A' = \neg A \Rightarrow B' = V$ (De no A se sigue la ignorancia)
4. Simetría entre modus ponens y modus tollens:
Si $B' = A' \circ_m^T f(A, B)$, luego $\neg A' = \neg B' \circ_m^T (f(A, B))$
5. $B' \subset^T B$, si $A' \subsetneq^T A$

Otro conjunto de axiomas pueden ser:

6. $A \overset{T}{\circ}_m f(A,B) \overset{T}{=} B$ Coincidencia con el modus ponens clásico
7. Siempre (i) $[muy A] \overset{T}{\circ}_m f(A,B) \overset{T}{=} B$ donde $\mu_{muy A} \overset{T}{=} (\mu_A)^2$, o (ii) $[muy A] \overset{T}{\circ}_m f(A,B) \overset{T}{=} muy B$.
8. $[más o menos A] \overset{T}{\circ}_m f(A,B) \overset{T}{=} más o menos B$, donde $\mu_{más o menos B} \overset{T}{=} \mu_B^{1/2}$.
9. Siempre (i) $\neg A \overset{T}{\circ}_m f(A,B) \overset{T}{=} V$ o (ii) $\neg A \overset{T}{\circ}_m f(A,B) \overset{T}{=} \neg B$

Para la formulación dada por una operación dinámica de combinación m que deriva una función de implicación I tal que de nuevo

$$A \overset{T}{\circ}_m I(A,B) \overset{T}{=} B.$$

Por supuesto los resultados son similares y se obtienen de resolver la ecuación funcional

$$\sup m[\mu_A(u(t)), I(\mu_A(u(t)), \mu_B(v(t)))] \overset{T}{=} \mu_B(v(t)) \quad (33)$$

donde m es la incógnita, o la ecuación dinámica relacional borrosa:

$$\sup m[\mu_A(u(t)), \mu_R(u(t), v(t))] \overset{T}{=} \mu_B(v(t)) \quad (34)$$

donde μ_R es la incógnita. La ventaja de la formulación de la ecuación (33) es que presenta la necesidad de especificar las propiedades de la operación dinámica de combinación m , en términos de que inferencia dinámica se produce. Se cuenta con cuatro propiedades:

- $M_1 = m(a, I(a,b)) \overset{T}{\leq} b$ es 1 del modus ponens
- $M_2 = m(1,1) \overset{T}{=} 1$ es el modus ponens
- $M_3 = m(0,1) \overset{T}{=} 0$ de acuerdo con 3 del modus ponens y M_1
- $M_4 = a \overset{T}{\leq} a' \Rightarrow m(a,b) \overset{T}{\leq} m(a',b)$ de acuerdo con 2 del modus ponens.

La formulación de la ecuación (34) insiste en el principio dinámico de combinación/proyección (la composición $\sup$ - m) es el bloque de construcción básico del proceso de razonamiento dinámico y deriva la implicación como una aplicación del principio de mínima especificidad.

Variantes de la deducción borrosa

Existe una contradicción entre tres requerimientos:

- El principio de mínima especificidad: Nos dice que si $\pi_{x(t)} \overset{T}{=} \mu_A$ y $\pi_{y(t)|x(t)} \overset{T}{=} I(\mu_A, \mu_B)$, luego la menor distribución dinámica específica de posibilidades en $(x(t), y(t))$ es $\min(\mu_A, I(\mu_A, \mu_B))$, forzando $m = \min$.

- La coincidencia entre el modus ponens clásico y el modus ponens dinámico cuando $A' = A$
- El empleo de varias implicaciones bien portadas para modelar las reglas borrosas dinámicas.

La única forma de eliminar las contradicciones entre los tres planteamientos es eliminar el axioma 6 del modus ponens dinámico y aceptar que en algunas situaciones, de A y una regla borrosa dinámica, donde el predicado borroso dinámico A aparece en la parte de condición dinámica, B no se sigue necesariamente.

8.13. APROXIMACION LOGICA DEL RAZONAMIENTO APROXIMADO DINAMICO

Claramente, si los predicados borrosos dinámicos están presentes en la parte condicional de una regla y si ésta es precisa por tener información disponible, el alcance en el cual una condición elemental se satisface, se interpreta naturalmente en términos del grado dinámico de membresía de la información correspondiente al conjunto borroso dinámico que representa al predicado dinámico. Cuando sólo se tiene disponible información incompleta, el patrón borroso dinámico genera problemas para estimar cual es el alcance, y si es cierto que la condición se satisface en el tiempo. Esta última situación ocurre en forma particular cuando las conclusiones borrosas dinámicas son permitidas en las reglas, y las reglas se usan en cascada en el proceso de razonamiento dinámico.

Vamos a ver los puntos de combinación y propagación usando operaciones dinámicas totalmente conectivas composicionales. Empezando con la regla de producción “si $p(t)$ luego $q(t)$ ” donde $p(t)$ es una condición compuesta de la forma “ $p_1(t)$ y ... $p_n(t)$ ”:

1. El grado de satisfacción asociado a $p(t)$, dijamos $d(p(t))$, se supone que es de la forma: $d(p(t)) = d(p_1(t)) * d(p_2(t)) * \dots * d(p_n(t))$;

2. La evaluación del grado $d(q(t))$ asociado a $q(t)$, se ve como el problema llamado problema disociado: de $d(p(t))$ y del grado asociado a la regla $d(p(t) \xrightarrow{T} q(t))$ donde $d(p(t) \xrightarrow{T} q(t)) = I(d(p(t)), d(q(t)))$

y I alguna función conectiva de implicación multivaluada, se encuentra $d(q(t))$;

3. La fusión de grados $d^1(q(t))$, ... , $d^k(q(t))$ obtenida de k diferentes cadenas de inferencia se logra por medio de alguna combinación de funciones C bajo la forma: $C(d^1(q(t)), \dots, d^k(q(t)))$.

Asumamos que todos los grados $d(p(t))$ considerados pertenecen al intervalo real $[0,1]$ por simplicidad. Los candidatos naturales para la operación $*$ son las operaciones llamadas normas triangulares dinámicas, como

$$* = \min, * = \text{producto}, a * b = \max(a + b - 1, 0).$$

Sin embargo no sólo se pueden considerar operaciones conjuntivas como promedios para la operación * cuando se consideran algunos efectos compensatorios o deseables entre los grados de satisfacción de condiciones elementales.

8.14. FACTORES DE CERTIDUMBRE Y LÓGICA MULTIVALUADA

La propagación se puede ver como un problema de separación, digamos, vemos para el límite inferior de $d(q(t))$, conociendo los límites inferiores de $d(p(t))$ y de

$$d(p(t) \xrightarrow{T} q(t)).$$

Eso corresponde a una extensión del esquema del modus ponens en lógica multivaluada. Del límite superior en $d(p(t))$ y del límite inferior en

$$d(p(t) \xrightarrow{T} q(t)) \text{ o en } d(\neg p(t) \xrightarrow{T} \neg q(t)),$$

se puede obtener un límite superior de $d(q)$. Esto corresponde a una extensión del esquema del modus tolens. Consideremos una función de implicación de la forma

$$d(p \xrightarrow{T} q) = 1 - \Delta(d(p), 1 - d(q))$$

donde Δ es una norma triangular continua que modela una operación de conjunción; es tal que

$$d(p \rightarrow q) \xrightarrow{T} d(\neg p \rightarrow \neg q)$$

si Δ es simétrica. Tenemos un primer patrón de propagación:

$$\begin{array}{l} d(p \xrightarrow{T} q) \geq a \\ d(q \xrightarrow{T} p) \geq a' \\ b \leq d(p) \leq B \\ \hline 1 - I_R(b, 1-a) \leq d(q) \leq I_R(1-B, 1-a') \end{array} \quad (8.20)$$

donde I_R es la función de implicación definida por residuación de Δ , llamémosla:

$$I_R(r, t) = \sup \{s, s \in [0,1], \Delta(s, r) \leq t\}$$

El intervalo que contiene a $d(q)$ de acuerdo con el patrón I puede ser vacío. Este es el caso por ejemplo si $a = a' = 1$ y $b = B = 0.05$. Esto significa que en este cálculo,

$$d(p), (p \xrightarrow{T} q), \text{ y } d(q \xrightarrow{T} p)$$

no son independientes, por ejemplo, algunos valores de $d(p)$ son prohibidos. Dados

$$d(p \rightarrow q) \text{ y } d(q \rightarrow p).$$

Por contraste, con

$$d(p \rightarrow q) = I_R(d(p), d(q))$$

tenemos dos patrones de inferencia distintos porque I_R no es igual a su recíproco:

$$\begin{array}{l} d(p \rightarrow q) \geq a \\ d(q \rightarrow p) \geq a' \\ b \leq d(p) \leq B \end{array} \quad (8.21)$$

$$\Delta(b, a) \leq d(q) \leq I_R(a', B)$$

y

$$\begin{array}{l} d(p \rightarrow q) \geq a \\ d(\neg p \rightarrow \neg q) \geq a' \\ b \rightarrow d(p) \leq B \end{array} \quad (8.22)$$

$$\Delta(b, a) \leq d(q) \leq I_R(a', B)$$

Se puede verificar que el límite inferior de $d(q)$ en (8.21) y (8.22) nunca serán mayores que las cotas superiores, de forma que el intervalo que contiene a $d(q)$ nunca es vacío.

Las cotas superiores e inferiores que se ven debajo de los patrones (8.20)-(8.22) sugieren una solución a la combinación de problemas de la forma:

$$\max_{i=1,\dots,k} d_{i*}(q(t)) \leq \min_{i=1,\dots,k} d_i^*(q(t))$$

para definir las cotas finales para la conclusión q , donde $d_{i*}(q(t))$ y $d_i^*(q(t))$ denotan los límites inferior y superior dados por la derivación i .

Cuando el límite inferior resultante llega a ser mayor que el límite inferior resultante, después de los pasos de la combinación, se interpreta como una contradicción en la base de conocimiento. Nótese que los patrones (8.21) o (8.22) junto con esta combinación de pasos lleva la conclusión q_j al límite inferior.

$$d(q(t)) \geq \max_{i=1,\dots,k} \Delta(b_i(t), a_{ij}(t))$$

donde $a_{ij} \leq d(p_i \rightarrow q_j)$ es el factor de certeza asociado a la regla “si p_i luego q_j ” y $b_i \leq d(p_i)$ es el factor de certeza asociado al hecho p_i . La parte derecha de la desigualdad se puede ver como una composición t-morma max de una relación borrosa (a_{ij}) , con un conjunto borroso (b_i) .

8.15. EL PUNTO DE VISTA TEÓRICO POSIBILÍSTICO DE LOS FACTORES DE CERTEZA

Para clarificar el significado de los números asociados a la formulación lógica en los patrones de inferencia deductivos, se han asumido conocidos ciertos límites para el valor de la medida de incertidumbre de proposiciones que forman la base de conocimiento. Esta medida de incertidumbre puede ser una medida de probabilidad, de posibilidad, de necesidad, etc. Esto significa que los grados no son composicionales con respecto a los conectivos lógico. En el caso de una medida de necesidad N que por definición satisface la propiedad

$$N(p \overset{T}{\wedge} q) = \overset{T}{\min} (N(p), N(q)),$$

se pueden establecer los siguientes patrones:

$$\begin{array}{c} N(p \overset{T}{\rightarrow} q) \geq a \\ N(q \overset{T}{\rightarrow} p) \geq a' \\ b \leq N(p) \leq B \\ \hline \min(b, a) \leq N(q) \leq 1 \text{ si } a' \leq B, \text{ y } B \text{ en otro valor} \end{array} \quad (8.23)$$

Se ve claro que el patrón (8.23) es un caso especial del (2) donde $d = N$, $\Delta = \overset{T}{\min}$ y I_R es la implicación dinámica de Gödel.

Con las medidas de incertidumbre se tienta uno a reemplazar la implicación material por condicionamiento en los patrones. Una medida condicionada de posibilidad $\Pi(\cdot | p(t))$ se define como que satisface la relación

$$\Pi(p(t) \overset{T}{\wedge} q(t)) = \overset{T}{\min} (\Pi(q(t) | p(t)), \Pi(p(t))).$$

El principio de mínima especificidad nos lleva a definir $\Pi(\cdot | p(t))$ como la mayor solución de la ecuación, digamos:

$$\Pi(q(t) | p(t)) = \begin{cases} 1 & \text{si } \Pi(q(t)) \overset{T}{\wedge} \Pi(p(t)) \\ \Pi(q(t) \overset{T}{\wedge} \Pi(p(t))) & \text{en otro valor} \end{cases}$$

La condición dual de la medida de necesidad es simplemente

$$N(q(t) | p(t)) = 1 - \Pi(\neg q(t) | p(t)),$$

luego cambiando la implicación dentro del condicionamiento en el patrón (4) se produce la misma conclusión. Por lo tanto, se puede obtener un fuerte patrón de inferencia usando la ecuación de condicionamiento y la definición de una medida de posibilidad

$$\Pi(p(t) \vee q(t)) = \max^T (\Pi(p(t)), \Pi(q(t))):$$

$$\Pi(q) = \max^T (\min^T \Pi(q(t) | p(t)), \Pi(p(t))), \min^T \Pi(q(t) | \neg p(t)), \Pi(\neg p(t))).$$

Este patrón es más fuerte que la versión condicionada de (4) y está en completo acuerdo con el principio de combinación/proyección de Zadeh.

8.16. UN PARADIGMA DE LOS SISTEMAS DE INFORMACIÓN

La frecuente confusión que existe en la relación entre la verdad y la incertidumbre en la literatura del razonamiento aproximado, aparentemente se debe a la carencia de un paradigma dedicado a interpretar la verdad parcial y los grados de incertidumbre en una estructura simple, aunque la distinción entre los dos conceptos se halla realizado mucho tiempo antes. Tal paradigma se puede inspirar por el punto de vista del sentido común, como la compatibilidad entre una declaración y la realidad. Esta primitiva definición de verdad ha sido criticada por los filósofos, sin embargo puede ser levemente modificada, si cambiamos la debatible palabra “realidad” por la frase qué sabemos acerca de la realidad, en un tiempo dado. Luego calculando el valor de verdad de una declaración S en un tiempo t, se puede estimar su conformidad con una descripción D que es lo que se conoce del estado actual del hecho. Como una consecuencia la evaluación de la verdad es un procedimiento de apareamiento de un patrón semántico en un tiempo. Se pueden encontrar la siguientes situaciones:

1. *Declaraciones dinámicas precisas con información completa.* S es una declaración bien definida sin ningún predicado o cuantificador vago, y D es una descripción dinámica precisa (completa) de estado actual del hecho. Siempre D es compatible con S y S es verdad, o D no es compatible con S y S es falsa. Esta es la situación que se maneja en lógica clásica.
2. *Declaraciones borrosas con información completa.* En este caso la conformidad de S con respecto a D llega a ser materia de grado porque S contiene predicados o cuantificadores dinámicos vagos, y el estado actual del hecho puede ser por ejemplo una frontera. Por ejemplo la declaración S para evaluar “Hoy Juan es alto” si sabemos que hoy Juan mide 1.75 m, es una situación en la cual no podemos decir si es completamente falsa, o completamente cierta. Luego lo que se puede llamar un grado de verdad $\tau(S|D) \in [0,1]$ se puede asociar a la declaración S, con un valor de verdad de $\tau(S|D) = \mu_{\text{altura}(t)}(1.75)$, donde $\tau(S|D) = 1$ significa que S es totalmente cierto y $\tau(S|D) = 0$ significa que S es falsa. Esta situación tiene que ver con la lógica multivaluada y en especial con la de funcionales de valor.
3. *Declaración precisa con información incompleta.* En este caso la conformidad de S con respecto a D no describe el estado actual del hecho. Por ejemplo, pueden ser d, d' dos estados distintos de hechos pero compatibles con D, de tal forma que d es compatible con S, pero d' es compatible con no S'. Por lo anterior el valor de verdad de S, que es siempre cierto o verdadero se puede desconocer. Se puede asociar un grado de posibilidad $\Pi(S)$ a S que por convención $\Pi(S) = 0$, significa $\tau(S|D) = 0$, S es falso y $\Pi(S) = 1$ significa que $\tau(S|D)$ es desconocida. S es verdad se expresa

también por $\Pi(\text{no } S) = 0$, o en forma equivalente por: $N(S) = 1 - \Pi(\text{no } S) = 1$, donde $N(S)$ se interpreta como un grado de certeza. Es pertinente hacer notar que este tipo de convenciones son naturales cuando D es un subconjunto ordinario de estados de facto. Cuando D es un subconjunto borroso, Π y N admiten grados y son medidas típicas de posibilidad y necesidad. Cuando D es una distribución entonces el grado de incertidumbre asociado a S es un único grado de probabilidad $P(S) = 1 - P, \text{ no } S$. Luego en la situación de un conocimiento incompleto, la incertidumbre asociada a la declaración S toma la forma de uno de dos números, dependiendo de la naturaleza del conocimiento disponible. Es claro que estos números no son grados de verdad, sino que sólo nos reflejan el estado de la creencia acerca de la verdad o falsedad de la declaración.

4. *Declaraciones borrosas dinámicas con información incompleta.* En este caso la verdad es materia de grado y puede ser mal conocida. En otras palabras para cada valor de verdad $\alpha = \tau(S | D)$ representa el grado de conformidad de la declaración borrosa S con un estado preciso del facto d compatible con D , corresponde al grado de incertidumbre de que S tenga un valor de verdad α , que refleja la incertidumbre de que d sea un estado de facto cierto. Esta es la situación más compleja. Cuando S y D , ambos se pueden expresar como conjuntos borrosos $\tau(S | D)$ es un valor de verdad borroso.

8.17. EL PROBLEMA DE COMPOSICIONALIDAD

Un importante resultado de la anterior interpretación de la verdad es que los grados de incertidumbre no pueden ser composicionales para todos los conectivos. Específicamente no existen operaciones $\perp$ y $*$ en $[0,1]$, no hay funciones de negación f tales que las siguientes identidades valgan simultáneamente para todas las proposiciones S_1, S_2, S donde g permanece para una función valuada en $[0,1]$ que se destina para estimar la incertidumbre:

1. $g(\text{no } S) = f(g(S))$
2. $g(S_1 \wedge S_2) = g(S_1) * g(S_2)$
3. $g(S_1 \vee S_2) = g(S_1) \perp g(S_2)$

Más específicamente 1-3 vinculan $\forall S, g(S) = 0$ o, $g(S) = 1$. Por ejemplo estamos en el caso determinístico donde una declaración es siempre cierta o falsa. Se permiten otras formas de composicionalidad, por ejemplo:

$$\Pi(S_1 \vee S_2) = \max(\Pi(S_1), \Pi(S_2)),$$

en teoría de posibilidades, pero generalmente

$$\Pi(S_1 \wedge S_2) \geq \min(\Pi(S_1), \Pi(S_2)).$$

La igualdad

$$\Pi(S_1 \wedge S_2) = \min(\Pi(S_1), \Pi(S_2))$$

es válida sólo para proposiciones que pertenecen a variables no interactivas. La anterior

imposibilidad que resulta es una nueva forma de afirmar un hecho bien conocido, que el intervalo $[0,1]$ no se puede equipar con una estructura de Álgebra Booleana.

Este resultado se base en la consideración que las proposiciones a evaluar no son borrosas y por lo mismo pertenecen al álgebra Booleana. Por contraste, los valores de verdad de proposiciones borrosas se pueden considerar composicionales cuando pueden ser evaluados precisamente, cuando no existe incertidumbre que perturbe la información. Esto por el hecho de que conjuntos cerrados de proposiciones borrosas no son álgebras booleanas más largas. Por ejemplo si usamos $\max$, $\min$, $1 - (\cdot)$ para expresar la disyunción, conjunción y negación de proposiciones borrosas equipando conjuntos de estas proposiciones con una estructura distributiva de trama, que es una estructura más fuerte que el álgebra Booleana y que es compatible con el intervalo unitario. Algunas veces argumentos contra los conjuntos borrosos confían en tópicos composicionales. Estos argumentos se pueden basar en la fuerte consideración de que el álgebra de proposiciones a ser evaluada es Booleana o que los grados de verdad intentan modelar incertidumbre.

La presencia o ausencia de reglas composicionales es un criterio para distinguir entre lógica de grados de verdad que maneja proposiciones vagas bajo información completa, y lógica de incertidumbre que maneja proposiciones clásicas bajo información incompleta.

El hecho de que la información imprecisa e incierta, que se representa en la teoría de las distribuciones de posibilidades se pueda algunas veces expresar lingüísticamente en términos de predicados vagos, que se representan por conjuntos borrosos, no nos permite pensar en que exista una confusión entre grados de posibilidad, que reflejan incertidumbre, y grados de verdad. De acuerdo con Zadeh, la información “Juan es alto”, se puede representar por la distribución de posibilidades $\pi_{x(t)}$ donde $\pi_{\text{altura (Juan)}}$ está definida por:

$$\forall u \pi_x(u(t)) = \mu_A(u(t)) \quad (1)$$

para nuestro ejemplo:

$$\pi_{\text{altura (Juan)}}(u(t)) = \mu_{\text{altura}}(u(t)).$$

Esto no significa que el grado de posibilidad y el grado de verdad son la misma cosa. Desde luego $\pi_{\text{altura (Juan)}}$ estima en que grado es posible que la proposición “altura (Juan) = u” sea verdad conociendo la información disponible, mientras $\mu_{\text{altura}}(u(t))$ estima que tan verdadero es que alguna persona cuya altura sea precisamente u, sea alto. Luego como grado de posibilidad y grado de verdad respectivamente $\pi_{\text{altura (Juan)}}(u(t))$ y $\mu_{\text{altura}}(u(t))$ ni se refieren a la misma, ni al mismo estado de conocimiento. Lo que (1) expresa es que si conocemos que Juan es alto, los más o menos valores de esta altura son precisamente los más o menos compatibles con el predicado vago de altura. La asignación (1) expresa que el grado de posibilidad d que Juan sea alto sea u, es proporcional a la compatibilidad de u con alto, cuando se conoce que Juan es alto, por lo que se puede escribir en forma más precisa:

$$\pi_{\text{altura (Juan)}}(u(t) \mid \text{altura}) = \mu_{\text{altura}}(u(t)),$$

que nos enfatiza la naturaleza no simétrica de la igualdad (1).

8.18. SISTEMAS FORMALES

Clásicamente un sistema lógico consiste de:

1. La definición de un lenguaje; una gramática que produce fórmulas bien formadas.
2. *Una semántica*, definida por medio de interpretaciones que mapean cada término básico del lenguaje in un elemento del dominio dependiendo de la interpretación, y cada fórmula bien formada en un valor de verdad de un conjunto L de verdades. Un subconjunto de L, que se llama el conjunto de valores de verdad designado, tal que la fórmula para la cual cualquier interpretación asociada a un valor de verdad designado, se llama tautología. El conjunto de valores de verdad designado, algunas veces se reduce al elemento tope de L.
3. *Una sintaxis*. Un sistema formal compuesto por axiomas y reglas de inferencia. Una fórmula derivable de este sistema formal se llama teorema y su derivación una prueba. Con respecto a la semántica, una sintaxis es sonora si todos los teoremas son tautologías y completa si todas las tautologías son teoremas.

8.19. TIPOS DE ECUACIONES RELACIONALES DINÁMICAS

En nuestra discusión consideraremos un conjunto borroso dinámico $x(t)$ en $\mathbf{X}$ y se especifica una relación borrosa dinámica $R(t)$ en $\mathbf{X} \times \mathbf{Y}$ en el proceso de composición. Los diferentes tipos de operadores dinámicos de composición que emplearemos en las ecuaciones relacionales dinámicas son los siguientes:

1. La composición dinámica Max-Min.
2. La composición dinámica Min-Max.
3. La ecuación relacional borrosa dinámica adjunta.

En la composición dinámica Max-Min combinamos un conjunto borroso dinámico $X(t) \in F[\mathbf{X}(t)]$ con una relación borrosa dinámica $R(t) \in F[\mathbf{X}(t) \times \mathbf{Y}(t)]$ para obtener un conjunto borroso dinámico $\mathbf{Y} \in F[\mathbf{Y}(t)]$ que representamos como:

$$\mathbf{Y}(t) = \mathbf{X}(t) \circ R(t)$$

Con la función de membresía dada por el máximo en un tiempo t sobre el mínimo respectivo en el mismo tiempo t, o sea:

$$Y(y(t)) = \max_{x \in X} \left[\min(X(x(t)), R(x(t), y(t))) \right] \vee \left[X(x(t)) \wedge R(x(t), y(t)) \right]$$

La composición dinámica Min-Max es la ecuación relacional dinámica dual de la anterior por lo que se puede escribir la función de membresía de la forma:

$$Y(y(t)) = \min_{x \in X} \left[\max(X(x(t)), R(x(t), y(t))) \right] \bigwedge_{x \in X} \left[X(x(t)) \bigvee R(x(t), y(t)) \right]$$

Los dos tipos de relaciones dinámicas se pueden generalizar usando tanto las normas dinámicas s como la t como conectivos dinámicos, lo que nos da:

$$Y(y(t)) = \max_{x \in X} \left[(X(x(t)) \mathbin{\mathbf{t}} R(x(t), y(t))) \right] \bigvee_{x \in X} \left[X(x(t)) \mathbin{\mathbf{t}} R(x(t), y(t)) \right]$$

y

$$Y(y(t)) = \min_{x \in X} \left[(X(x(t)) \mathbin{\mathbf{s}} R(x(t), y(t))) \right] \bigwedge_{x \in X} \left[X(x(t)) \mathbin{\mathbf{s}} R(x(t), y(t)) \right]$$

Como el mínimo en un tiempo t es la máxima norma t dinámica, sea:

$$a(t) \mathbin{\mathbf{t}} b(t) \leq \min[a(t), b(t)]$$

para toda norma t dinámica, y el máximo es el más pequeño, por lo que tenemos las siguientes cotas que caracterizan el intervalo de todos los posibles resultados de la composición:

$$\bigvee_{x \in X} \left[X(x(t)) \mathbin{\mathbf{t}} R(x(t), y(t)) \right] \leq \bigvee_{x \in X} \left[X(x(t)) \mathbin{\mathbf{t}} R(x(t), y(t)) \right]$$

$$\bigwedge_{x \in X} \left[X(x(t)) \mathbin{\mathbf{s}} R(x(t), y(t)) \right] \leq \bigwedge_{x \in X} \left[X(x(t)) \mathbin{\mathbf{s}} R(x(t), y(t)) \right]$$

En la ecuación de relación borrosa dinámica adjunta se introduce el pseudo complemento borroso dinámico. Para t normas dinámicas continuas, podemos definir el siguiente operador borroso dinámico:

$$a(t) \mathbin{\mathbf{\diamond}} b(t) = \sup \left(c(t) \in [0,1] \mid a(t) \mathbin{\mathbf{t}} c(t) \leq b(t) \right)$$

Para las t normas que se especifican como un mínimo se tiene al siguiente operador:

$$a(t) \mathbin{\mathbf{\nabla}} b(t) = \sup \left(c(t) \in [0,1] \mid a(t) \mathbin{\mathbf{\wedge}} c(t) \leq b(t) \right) = \begin{cases} 1 & \text{si } a(t) \leq b(t) \\ b(t) & \text{si } a(t) > b(t) \end{cases}$$

BIBLIOGRAFÍA

BIBLIOGRAFIA

1. Bezdek J. (de.) *Analysis of Fuzzy Information, Vol 1: Mathematics and Logic*. Boca Raton, FL: CRC press. 1987.
2. Bezdek J. (de.) *Analysis of Fuzzy Information, Vol II: Artificial Intelligence and Decision Systems*. Boca Raton, FL: CRC press. 1987.
3. Bezdek J. (de.) *Analysis of Fuzzy Information, Vol III: Applications in Engineering and Science*. Boca Raton, FL: CRC press. 1987.
4. Dubois, D and H. Prade. *Fuzzy Sets and Systems: Theory and applications*. New York. Academic Press. 1980.
5. Esragh, F. and E.H. Mamdani. *A general approach to linguistic approximation*. Int. J. Man-Mach. Stud., **11**, 101-519.
6. Gupta M. and T. Yamakawa. (eds.) *Fuzzy computing: Theory, Hardware and Applications*. Amsterdam: North-Holland. 1988.
7. Gupta M. and T. Yamakawa. (eds.) *Fuzzy computing: Theory, Hardware and Applications*. Amsterdam: North-Holland. 1988.
8. Gupta M. and T. Yamakawa. (eds.) *Fuzzy Logic in Knowledge-Based Systems, Decision and Control*. Amsterdam: North-Holland. 1988.
9. Lee C. *Fuzzy Logic in Control Systems: Fuzzy Controllers. Part I*. IEEE Trans. Syst. Man and Cybern. **20** 404-418. 1990
10. Lee C. *Fuzzy Logic in Control Systems: Fuzzy Controllers. Part II*. IEEE Trans. Syst. Man and Cybern. **20** 419-433. 1990
11. Negoita C. and D.A. Ralescu. *Applications of Fuzzy Sets to System Analysis*. Stuttgart: Birkhäuser. 1975.
12. Novák V. *Fuzzy Sets and Their Applications*. Bristol: Adam Hilger. 1989.
13. Novák V. *On the Syntactico-semantical completeness of First-order Fuzzy Logic. Part I. Syntactical aspects*. Kybernetika **26**. 1990.
14. Novák V. *On the Syntactico-semantical completeness of First-order Fuzzy Logic. Part II. Main results*. Kybernetika **26**. 1990.
15. Novák V. and W. Pedrycz. *Fuzzy Sets and T-norms in the Light of Fuzzy Logic*. Int. J. Man-Machine Stud. **29** 113-127. 1988.
16. Pavelka J. *On Fuzzy Logic: Part I*. Zeit. Math. Logik Gurndl. Math. **25** 45-52. 1979.
17. Pavelka J. *On Fuzzy Logic: Part II*. Zeit. Math. Logik Gurndl. Math. **25** 119-134. 1979.
18. Pavelka J. *On Fuzzy Logic: Part III*. Zeit. Math. Logik Gurndl. Math. **25** 447-464. 1979.
19. Sanches A. and L. A. Zadeh. *Approximate Reasoning in Intelligent Systems. Decision and Control*. Oxford: Pergamon Press. 1987.
20. Zadeh L.A. *Fuzzy Sets*. Inf. Control, **8**, 338-353.
21. Zadeh L.A. *Quantitative Fuzzy Semantics*. Inf. Sci. **3**, 159-176. 1973.
22. Zadeh L.A. *The Concept of a Linguistic variable and its Applications to Approximate Reasoning I*. Inf. Sci. **8**, 199-257. 1975.
23. Zadeh L.A. *The Concept of a Linguistic variable and its Applications to Approximate Reasoning II*. Inf. Sci. **8**, 301-357. 1975.
24. Zadeh L.A. *The Concept of a Linguistic variable and its Applications to Approximate Reasoning III*. Inf. Sci. **9**, 43-80. 1975.
25. Zadeh L.A. *A computational approach to Fuzzy Quantifiers in Natural Languages*. Comp. Math. Appl. **9**, 149-184. 1983.