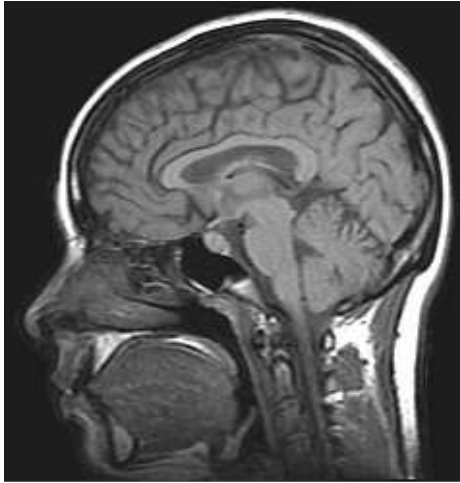




# 2008

## Modelos Híbridos de Inteligencia Computacional aplicados en la Segmentación de Imágenes de Resonancia Magnética



Universidad Nacional de Mar del Plata, Facultad de Ingeniería  
Tesis del Doctorado en Ingeniería, orientación Electrónica  
Ing. Gustavo Javier Meschino

Directora: Dra. Emilce Moler  
Codirectora: Dra. Virginia Ballarin



RINFI es desarrollado por la Biblioteca de la Facultad de Ingeniería de la Universidad Nacional de Mar del Plata.

Tiene como objetivo recopilar, organizar, gestionar, difundir y preservar documentos digitales en Ingeniería, Ciencia y Tecnología de Materiales y Ciencias Afines.

A través del Acceso Abierto, se pretende aumentar la visibilidad y el impacto de los resultados de la investigación, asumiendo las políticas y cumpliendo con los protocolos y estándares internacionales para la interoperabilidad entre repositorios



Esta obra está bajo una [Licencia Creative Commons Atribución- NoComercial-CompartirIgual 4.0 Internacional](https://creativecommons.org/licenses/by-nc-sa/4.0/).

***"En tanto las leyes de la matemática se refieren a la realidad, no son ciertas;  
en tanto son ciertas, no se refieren a la realidad."***

***"La mayoría de las ideas fundamentales de la ciencia  
son esencialmente sencillas y, por regla general,  
pueden ser expresadas en un lenguaje comprensible para todos."***

*Albert Einstein (1879-1955), científico estadounidense de origen alemán.*

# Tabla de Contenidos

---

|                                                                            |           |
|----------------------------------------------------------------------------|-----------|
| <b>TABLA DE ACRÓNIMOS.....</b>                                             | <b>7</b>  |
| <b>TABLA DE FIGURAS.....</b>                                               | <b>8</b>  |
| <b>1. INTRODUCCIÓN.....</b>                                                | <b>13</b> |
| 1.1. Motivación y presentación del problema .....                          | 13        |
| 1.2. Fundamento del sistema propuesto .....                                | 16        |
| 1.3. Antecedentes de la problemática a abordar.....                        | 17        |
| 1.3.1. Segmentación manual .....                                           | 19        |
| 1.3.2. Segmentación a través de métodos de procesamiento de imágenes ..... | 19        |
| 1.3.2.1. Umbralamiento .....                                               | 19        |
| 1.3.2.2. Métodos basados en gradiente .....                                | 20        |
| 1.3.2.3. Métodos basados en el histograma de intensidades.....             | 20        |
| 1.3.2.4. Contornos Activos – Modelos Deformables.....                      | 20        |
| 1.3.2.5. Crecimiento de Regiones .....                                     | 21        |
| 1.3.2.6. Morfología Matemática .....                                       | 21        |
| 1.3.3. Segmentación a través de Reconocimiento de Patrones.....            | 22        |
| 1.3.3.1. Métodos de clasificación supervisados.....                        | 24        |
| 1.3.3.2. Métodos de clasificación no supervisados.....                     | 27        |
| 1.3.4. Segmentación con técnicas de Inteligencia Computacional .....       | 28        |
| 1.3.4.1. Redes Neuronales .....                                            | 29        |
| 1.3.4.2. Sistemas basados en Lógica Difusa .....                           | 30        |
| 1.3.4.3. Algoritmos Genéticos.....                                         | 32        |
| 1.3.4.4. Sistemas basados en Agentes .....                                 | 33        |
| 1.4. Objetivos y aportes de esta tesis.....                                | 33        |
| 1.5. Trabajos previos propios .....                                        | 35        |
| 1.6. Estructura de la tesis .....                                          | 36        |
| <b>2. IMÁGENES DE RESONANCIA MAGNÉTICA.....</b>                            | <b>39</b> |
| 2.1. Introducción .....                                                    | 39        |
| 2.2. Fundamentos de la Resonancia Magnética.....                           | 41        |
| 2.2.1. Bases físicas .....                                                 | 41        |
| 2.2.2. Proceso de excitación .....                                         | 43        |
| 2.2.3. Secuencias de lectura de la señal .....                             | 47        |

|           |                                                                                     |           |
|-----------|-------------------------------------------------------------------------------------|-----------|
| 2.2.3.1.  | Saturación - Recuperación .....                                                     | 49        |
| 2.2.3.2.  | Inversión - Recuperación.....                                                       | 51        |
| 2.2.3.3.  | Secuencia de pulso spin-eco .....                                                   | 53        |
| 2.3.      | Distorsiones en las imágenes .....                                                  | 60        |
| 2.4.      | El rol de la RM en la detección de enfermedades cerebrales.....                     | 61        |
| <b>3.</b> | <b>LÓGICA DIFUSA .....</b>                                                          | <b>67</b> |
| 3.1.      | Introducción .....                                                                  | 67        |
| 3.2.      | Historia .....                                                                      | 68        |
| 3.3.      | Conceptos de Lógica Difusa .....                                                    | 69        |
| 3.3.1.    | Imprecisión .....                                                                   | 69        |
| 3.3.2.    | Conjuntos Difusos.....                                                              | 71        |
| 3.3.3.    | Operaciones entre Conjuntos Difusos.....                                            | 75        |
| 3.4.      | Toma de decisiones con Lógica Difusa .....                                          | 76        |
| 3.4.1.    | La Lógica Difusa y la modelado de la decisión.....                                  | 78        |
| 3.4.2.    | Expresiones difusas y diferentes lógicas .....                                      | 79        |
| 3.4.3.    | Modificadores.....                                                                  | 83        |
| 3.4.3.1.  | Concentración .....                                                                 | 83        |
| 3.4.3.2.  | Dilatación .....                                                                    | 84        |
| 3.4.4.    | Lógica Difusa Compensatoria (LDC).....                                              | 84        |
| 3.4.5.    | Modelización de expresiones .....                                                   | 86        |
| <b>4.</b> | <b>PROCESAMIENTO DE IMÁGENES DE RESONANCIA MAGNÉTICA: MÉTODO<br/>PROPUESTO.....</b> | <b>90</b> |
| 4.1.      | Introducción .....                                                                  | 90        |
| 4.2.      | Determinación de predicados .....                                                   | 91        |
| 4.3.      | Cuantificación de los valores de verdad de los predicados difusos .....             | 94        |
| 4.4.      | Implementación y formalización del método .....                                     | 100       |
| 4.5.      | Optimización con Algoritmos Genéticos .....                                         | 103       |
| 4.5.1.    | Sistemas híbridos.....                                                              | 104       |
| 4.5.2.    | Sistema de optimización propuesto.....                                              | 107       |
| 4.6.      | Esquema del método propuesto .....                                                  | 114       |
| 4.7.      | Validación del método.....                                                          | 115       |
| 4.7.1.    | Medidas de similitud .....                                                          | 115       |
| 4.7.1.1.  | Coeficiente de Tanimoto.....                                                        | 117       |
| 4.7.1.2.  | Coeficiente de Exactitud .....                                                      | 119       |

|                    |                                                      |            |
|--------------------|------------------------------------------------------|------------|
| 4.7.1.3.           | Porcentaje de error en la clasificación .....        | 120        |
| 4.7.1.4.           | Coeficiente de Dice .....                            | 120        |
| 4.7.1.5.           | Matrices de Confusión .....                          | 121        |
| 4.7.2.             | Imágenes de prueba provenientes de simulaciones..... | 121        |
| 4.7.3.             | Imágenes de prueba reales .....                      | 123        |
| 4.7.3.1.           | Conjunto de imágenes de prueba #1 .....              | 124        |
| 4.7.3.2.           | Conjunto de imágenes de prueba #2 .....              | 125        |
| <b>5.</b>          | <b>RESULTADOS .....</b>                              | <b>127</b> |
| 5.1.               | Con imágenes simuladas .....                         | 127        |
| 5.1.1.             | Sin distorsión .....                                 | 127        |
| 5.1.1.1.           | Medidas de calidad .....                             | 130        |
| 5.1.1.2.           | Matrices de confusión.....                           | 131        |
| 5.1.2.             | Con distorsión.....                                  | 132        |
| 5.1.3.             | Con los operadores MAX y MIN .....                   | 137        |
| 5.1.4.             | Comparaciones con otros métodos .....                | 140        |
| 5.1.5.             | Análisis del TC mediante test estadístico .....      | 143        |
| 5.2.               | Con imágenes reales.....                             | 146        |
| 5.2.1.             | Conjunto de imágenes #1.....                         | 147        |
| 5.2.1.1.           | Mediciones de calidad de la segmentación .....       | 147        |
| 5.2.1.2.           | Comparación con otros métodos de segmentación .....  | 150        |
| 5.2.2.             | Conjunto de imágenes #2.....                         | 150        |
| 5.2.2.1.           | Mediciones de calidad de la segmentación .....       | 150        |
| 5.2.2.2.           | Comparación con otros métodos de segmentación .....  | 153        |
| 5.3.               | Discusión.....                                       | 154        |
| <b>6.</b>          | <b>CONCLUSIONES .....</b>                            | <b>157</b> |
| <b>APÉNDICE A:</b> | <b>ALGORITMOS GENÉTICOS.....</b>                     | <b>164</b> |
|                    | Introducción.....                                    | 164        |
|                    | Nociones de Algoritmos Genéticos .....               | 166        |
|                    | Codificación .....                                   | 166        |
|                    | Población Inicial.....                               | 166        |
|                    | Creación de una nueva generación .....               | 167        |
|                    | Detención del Algoritmo .....                        | 167        |
|                    | Diversidad de la población .....                     | 168        |

|                                                         |            |
|---------------------------------------------------------|------------|
| Aptitud de los individuos .....                         | 169        |
| Selección .....                                         | 169        |
| Opciones para la reproducción .....                     | 170        |
| Mutación.....                                           | 171        |
| Caso especial – Cruzamiento sin mutación .....          | 172        |
| Caso especial – Mutación sin cruzamiento .....          | 172        |
| Aplicación.....                                         | 172        |
| <b>APÉNDICE B: INTERFAZ GRÁFICA DE DESARROLLO .....</b> | <b>176</b> |
| <b>APÉNDICE C: CONJUNTOS DIFUSOS TIPO 2 .....</b>       | <b>181</b> |
| <b>AGRADECIMIENTOS.....</b>                             | <b>184</b> |
| <b>REFERENCIAS.....</b>                                 | <b>187</b> |

---

# Tabla de Acrónimos

---

|      |   |                                                                     |
|------|---|---------------------------------------------------------------------|
| AG   | = | Algoritmo Genético                                                  |
| FCM  | = | <i>Fuzzy-C-Means</i>                                                |
| FID  | = | Free Induction Decay                                                |
| IC   | = | Inteligencia Computacional                                          |
| INU  | = | <i>Intensity Non-uniformity</i> , No uniformidades en la intensidad |
| IRM  | = | Imágenes de Resonancia Magnética                                    |
| LCR  | = | Líquido Cefalorraquídeo ( <i>Cerebro Spinal Fluid</i> )             |
| LD   | = | Lógica Difusa                                                       |
| MB   | = | Materia Blanca ( <i>White Matter</i> )                              |
| MG   | = | Materia Gris ( <i>Gray Matter</i> )                                 |
| PD   | = | <i>Proton Density</i> , Densidad de Protones                        |
| PDLC | = | Predicados Difusos y Lógica Compensatoria, el método propuesto      |
| RM   | = | Resonancia Magnética                                                |
| RN   | = | Redes Neuronales                                                    |
| SOM  | = | <i>Self Organizing Maps</i> , Mapas Autoorganizados                 |
| TAC  | = | Tomografía Axial Computada                                          |
| TC   | = | <i>Tanimoto Coefficient</i> , Coeficiente de Tanimoto               |
| TSK  | = | Sistema de inferencia difusa de <i>Takashi, Sugeno, Kang</i> .      |

---

# Tabla de Figuras

---

|            |                                                                                                                        |    |
|------------|------------------------------------------------------------------------------------------------------------------------|----|
| Figura 1:  | Métodos de segmentación de Imágenes. ....                                                                              | 18 |
| Figura 2:  | Diferentes tecnologías para Diagnóstico por Imágenes.....                                                              | 40 |
| Figura 3:  | Núcleos de Hidrógeno con sus vectores alineados aleatoriamente en todas las direcciones. ....                          | 43 |
| Figura 4:  | Efecto apreciable en los núcleos de Hidrógeno ante la presencia de un fuerte campo magnético externo.....              | 43 |
| Figura 5:  | Vector de magnetización y precesión en espiral.....                                                                    | 45 |
| Figura 6:  | Tiempo de relajación $T_1$ .....                                                                                       | 45 |
| Figura 7:  | Tiempo de relajación $T_2$ .....                                                                                       | 48 |
| Figura 8:  | Tiempo de relajación $T_2^*$ .....                                                                                     | 48 |
| Figura 9:  | Secuencia saturación-recuperación. ....                                                                                | 50 |
| Figura 10: | Secuencia inversión-recuperación. ....                                                                                 | 52 |
| Figura 11: | Secuencia inversión-recuperación. (Reproducida de [Blink, 2004]) .....                                                 | 53 |
| Figura 12: | Variación de la magnitud del vector de magnetización longitudinal (Reproducida de [Coussement, 2000]).....             | 53 |
| Figura 13: | Secuencia de pulso spin-eco (Reproducida de [Blink, 2004]). ....                                                       | 54 |
| Figura 14: | Imagen de densidad de protones (PD).....                                                                               | 57 |
| Figura 15: | Imagen pesada en $T_2$ . ....                                                                                          | 57 |
| Figura 16: | Imagen pesada en $T_1$ . ....                                                                                          | 58 |
| Figura 17: | Diferentes formas de funciones de pertenencia.....                                                                     | 74 |
| Figura 18: | Función de pertenencia para definir el conjunto difuso “Gris oscuro” aplicado a la variable “Intensidad de Gris” ..... | 75 |
| Figura 24: | Histogramas de frecuencias relativas de intensidades de gris.....                                                      | 95 |
| Figura 25: | Funciones de pertenencia preliminares.....                                                                             | 96 |

|                                                                                                                                         |     |
|-----------------------------------------------------------------------------------------------------------------------------------------|-----|
| Figura 26: Función de pertenencia tipo “Z” y sus parámetros.....                                                                        | 97  |
| Figura 27: Función de pertenencia tipo “S” y sus parámetros.....                                                                        | 98  |
| Figura 28: Función de pertenencia tipo gaussiana asimétrica y sus parámetros. ....                                                      | 99  |
| Figura 33: Progreso típico de entrenamiento de un Algoritmo Genético. ....                                                              | 111 |
| Figura 34: Funciones de pertenencia antes y después de la optimización realizada por el Algoritmo Genético.....                         | 113 |
| Figura 35: Esquema de los pasos de diseño y optimización del sistema. ....                                                              | 114 |
| Figura 36: Esquema de las consideraciones necesarias para el cálculo de diferentes medidas de similitud. ....                           | 116 |
| Figura 37: Ejemplo de imágenes binarias para el cálculo de diferentes medidas de calidad de una segmentación.....                       | 117 |
| Figura 38: Coeficientes de Tanimoto.....                                                                                                | 119 |
| Figura 39: Imágenes pesadas en $T_1$ con distintos niveles de distorsión. ....                                                          | 123 |
| Figura 40: Resultados obtenidos en el corte #30 de una imagen simulada, sin distorsiones (ruido y no-uniformidades en intensidad).....  | 128 |
| Figura 41: Resultados obtenidos en el corte #60 de una imagen simulada, sin distorsiones (ruido y no-uniformidades en intensidad).....  | 128 |
| Figura 37: Resultados obtenidos en el corte #90 de una imagen simulada, sin distorsiones (ruido y no-uniformidades en intensidad).....  | 129 |
| Figura 43: Resultados obtenidos en el corte #120 de una imagen simulada, sin distorsiones (ruido y no-uniformidades en intensidad)..... | 129 |
| Figura 44: Resultados obtenidos en el corte #150 de una imagen simulada, sin distorsiones (ruido y no-uniformidades en intensidad)..... | 130 |
| Figura 40: Resultado obtenido con nivel de ruido de 3% y de INU de 20%. ....                                                            | 133 |
| Figura 41: Resultado obtenido en el caso de máximo nivel de ruido y de INU disponible en simulación. ....                               | 134 |

|                                                                                                                                       |     |
|---------------------------------------------------------------------------------------------------------------------------------------|-----|
| Figura 47: Coeficientes de Tanimoto obtenidos con el método propuesto para el líquido cefalorraquídeo (LCR).....                      | 136 |
| Figura 48: Coeficientes de Tanimoto obtenidos con el método propuesto para la materia gris (MG).....                                  | 136 |
| Figura 49: Coeficientes de Tanimoto obtenidos con el método propuesto para la materia blanca (MB).....                                | 137 |
| Figura 50: Coeficientes de Tanimoto obtenidos con los operadores MAX – MIN para el líquido cefalorraquídeo (LCR). ....                | 138 |
| Figura 51: Coeficientes de Tanimoto obtenidos con los operadores MAX – MIN para la materia gris (MG).....                             | 139 |
| Figura 52: Coeficientes de Tanimoto obtenidos con los operadores MAX – MIN para la materia blanca (MB). ....                          | 139 |
| Figura 44: Representación gráfica de los Coeficientes de Tanimoto obtenidos a través de diferentes métodos en imágenes simuladas..... | 143 |
| Figura 54: Gráficos de caja que muestra los promedios de los Coeficientes de Tanimoto (calculados en los tres tejidos).....           | 145 |
| Figura 55: Corte #59 de una imagen real (Conjunto #1).....                                                                            | 147 |
| Figura 56: Corte #79 de una imagen real (Conjunto #1).....                                                                            | 148 |
| Figura 57: Corte #99 de una imagen real (Conjunto #1).....                                                                            | 148 |
| Figura 59: Corte #80 de una imagen real (Conjunto #2).....                                                                            | 151 |
| Figura 60: Corte #90 de una imagen real (Conjunto #2).....                                                                            | 151 |
| Figura 61: Corte #100 de una imagen real (Conjunto #2).....                                                                           | 152 |
| Figura 63: Funciones de pertenencia tipo 2 preliminares, a ser posteriormente optimizadas. ....                                       | 160 |
| Figura 64: Resultado preliminar para una Resonancia Magnética de hombro.....                                                          | 161 |
| Figura 19: Posible representación del cromosoma que identifica los números 13 y 3 en un Algoritmo Genético. ....                      | 166 |

|                                                                                                                               |     |
|-------------------------------------------------------------------------------------------------------------------------------|-----|
| Figura 20: Diferentes métodos de generación de hijos para una nueva población en un Algoritmo Genético. ....                  | 167 |
| Figura 21: Selección estocástica uniforme. ....                                                                               | 170 |
| Figura 22: Posible configuración de la proporción de individuos en una nueva generación generados por diferentes métodos..... | 171 |
| Figura 23: Representación esquemática que resume el funcionamiento iterativo de un Algoritmo Genético. ....                   | 173 |
| Figura 61: Interfaz gráfica utilizada para pruebas.....                                                                       | 177 |
| Figura 30: Interfaz gráfica para visualización de resultados. ....                                                            | 178 |
| Figura 31: Representación de los valores de verdad obtenidos. ....                                                            | 179 |
| Figura 32: Imágenes de los tejidos asignados a cada píxel. ....                                                               | 179 |
| Figura 65: Representación gráfica de un conjunto difuso de tipo 2. ....                                                       | 182 |

---

# 1. Introducción

# 1. Introducción

---

## 1.1. Motivación y presentación del problema

---

Los sistemas de procesamiento de imágenes digitales constituyen actualmente una herramienta casi indispensable en la práctica de la medicina moderna. Los sistemas de adquisición muestran un desmesurado crecimiento que se incrementa día a día. Sin embargo, la evolución de los equipos a veces no es reflejada en el proceso de interpretación de imágenes. Por lo tanto, es de fundamental importancia el desarrollo de nuevos paradigmas y metodologías con el fin de disponer un espectro importante de opciones para el procesamiento y la visualización de las mismas.

Una de las tareas más importantes en el análisis de imágenes médicas es la segmentación, entendiéndose como tal al proceso de particionarlas según sus componentes estructurales más importantes en regiones homogéneas con respecto a alguna de sus características, como textura o intensidad [Pham et al., 2000]. Un método de segmentación busca una partición tal que las regiones obtenidas correspondan a distintas estructuras anatómicas o regiones de interés de la imagen.

Una segmentación precisa es requisito indispensable para gran cantidad de aplicaciones, como cálculo de volúmenes de ciertos tejidos y su posterior representación tridimensional, terapia de radiación, planes de cirugía, detección de tejidos anormales. Una vez realizada la segmentación, la información puede usarse por los especialistas para comparar volúmenes, morfologías y características de los tejidos con estudios normales u otras regiones en la misma imagen. Así pueden determinarse parámetros de normalidad con la idea de detectar patologías y asistir a las decisiones en diagnóstico y terapia [Moler, 2003, Courchesne et al., 2000].

Esta tesis surge como producto del trabajo sistemático con imágenes digitales en conjunto con un grupo interdisciplinario de médicos, integrado por especialistas en Diagnóstico por Imágenes, patólogos, un traumatólogo, un anestesiólogo y un psiquiatra.

De la interacción con los médicos surgen pautas concretas para interpretar las imágenes, que pueden provenir de equipos distintos. En conjunto con los patólogos se trabajó con imágenes microscópicas de biopsias de médula ósea, en las que la segmentación adecuada permitió obtener medidas para evaluar presencia o grado de ciertos desórdenes metabólicos [Moler, 2003]. Con el traumatólogo y los especialistas en imágenes se estudiaron **imágenes de resonancia magnética** (IRM) de hombro, con el psiquiatra y los mismos especialistas se estudiaron IRM de cerebro. Con el anestesiólogo se trabajó con imágenes de tomografía de pulmón [Tusman et al., 2006].

De esta experiencia se fueron adquiriendo diversas técnicas y métodos utilizados por los expertos en la interpretación de imágenes. Su conocimiento puede representarse en un conjunto de predicados o sentencias que ayudan a identificar los diversos componentes de las imágenes. Por ejemplo: en biopsias de médula ósea “la celularidad es una zona de textura rugosa”, o más específicamente en lo que hace a esta tesis, en IRM “el líquido se ve hipointenso o negro en la imagen T1 y blanco o hiperintenso en la imagen T2” (las imágenes T1, T2 y PD son propias de las IRM y serán explicadas en el Capítulo 2). Así surge la idea de implementar estas consideraciones en un sistema que, a partir del conocimiento, efectúe un procesamiento objetivo de las imágenes.

Dado que los conceptos involucrados (en los ejemplos anteriores: “rugosa”, “hiperintenso”, “hipointenso”, etc.) son esencialmente subjetivos e imprecisos, es inmediato pensar en sistemas basados en Inteligencia Computacional y en particular en **Lógica Difusa** (LD) como herramienta principal.

Paralelamente, se fueron adquiriendo conocimientos sobre el enfoque axiomático de la Lógica Multivaluada como una extensión de la Lógica de Predicados tradicional. El esquema de toma de decisiones que ofrece este paradigma fue extendido para utilizarse en el procesamiento de imágenes. En lo que respecta a este tema, se exploraron las posibilidades de este enfoque conjuntamente con un experto en el estudio de Lógica Multivaluada.

Como aplicación preliminar se tomaron IRM de cerebro, en las cuales se pretendió discriminar diferentes tipos de tejidos. En la primera fase de diseño del sistema se consideró el conocimiento de los expertos expresado lingüísticamente en

forma de predicados. En la segunda fase se optimizaron ciertos parámetros del mismo, agregando la información que contenía una imagen segmentada previamente, aunque fuera parcialmente, que presentara píxeles prototípicos correctamente identificados de cada uno de los tejidos a detectar.

El paradigma elegido para la fase de optimización se basa en los llamados Sistemas Híbridos. Los sistemas que consideran este paradigma son de especial interés en el ámbito de la Inteligencia Computacional. Proponen la utilización simultánea y complementaria de dos o más técnicas, como por ejemplo la combinación de LD con **Algoritmos Genéticos** (AG), sistemas denominados Difuso-Genéticos (*Genetic Fuzzy Systems*).

Los AG han demostrado ser una herramienta apropiada para este propósito y han sido ampliamente utilizados en la optimización de sistemas de inferencia difusos basados en reglas, de tipo Mamdani o Sugeno, con resultados alentadores [Pal and Pal, 2003].

Los AG son una técnica de búsqueda utilizada para hallar soluciones aproximadas en problemas de optimización de parámetros. Constituyen una clase particular de algoritmos evolutivos que usan técnicas inspiradas en la biología, tales como herencia, mutación, selección y cruzamiento, como se explicará detalladamente. Son útiles en problemas que requieren procesos de búsqueda eficientes y efectivos [Goldberg, 1989].

Los posibles enfoques de optimización de un sistema basado en LD son muy variados y podrían orientarse a:

- la optimización de las funciones de pertenencia,
- la selección de predicados,
- la optimización de los predicados,
- la incorporación de modificadores en los predicados.

Estos enfoques no son exhaustivos, pero constituyen los fundamentales para este sistema basado en análisis de predicados. Podrían abordarse de a uno por separado o simultáneamente.

## 1.2. Fundamento del sistema propuesto

---

Esta tesis presenta un marco de trabajo (*framework*) general para segmentación de imágenes, el que comprende las siguientes etapas:

- Extracción del conocimiento de los expertos, etapa denominada “elicitación del conocimiento”.
- Construcción de un modelo compuesto por predicados lingüísticos que describan las componentes de la imagen que se desea segmentar.
- Elección de características de la imagen y conjuntos difusos que permitan la cuantificación del valor de verdad de los predicados definidos en el paso anterior.
- Optimización del sistema con AG a partir de píxeles prototipos correctamente segmentados (o clasificados).
- Procesamiento de nuevas imágenes provenientes de la misma población estadística con la que se ha determinado el modelo.

Aunque no es sencillo llegar al modelo lingüístico adecuado y optimizado, una vez obtenido éste, el esquema de procesamiento no involucra cálculos matemáticos complejos, por lo que los tiempos de procesamiento son relativamente cortos.

Si bien la idea es presentar un marco de trabajo general, esta tesis se ha desarrollado fundamentalmente para IRM de cerebro. La principal ventaja de las IRM es la discriminación de diferentes tipos de tejidos para una posterior cuantificación de las mismas y de esta manera asistir en el diagnóstico de diferentes patologías. La segmentación de este tipo de imágenes es un requerimiento constante en medicina.

Pero una de las dificultades que se presenta en las IRM es que hay un gran solapamiento entre las intensidades de gris que presentan diferentes tejidos. Es por eso que un enfoque con técnicas que trabajen con la modelización de la vaguedad parecen ser adecuadas, como es el caso de la que se propone en este trabajo. En este sentido, la LD ofrece un esquema de trabajo adecuado y ha sido empleada en este contexto [Denkowski et al., 2004], conjugando la ventaja de implementar conceptos inciertos con la posibilidad de manejar sentencias en lenguaje natural. La modelación

de la vaguedad se logra a través de variables lingüísticas, lo que permite aprovechar el conocimiento de los expertos, al contrario de lo que ocurre en otros métodos más cercanos a las cajas negras y exclusivamente basados en datos, como son, por ejemplo, las **Redes Neuronales** (RN). Se puede lograr un sistema que puede ser más fácilmente comprendido y aceptado por los profesionales que no dominan conceptos complejos de matemática, como es el caso de los médicos, quienes finalmente podrían ser los usuarios.

Se abordó el problema como uno correspondiente a la disciplina de asistencia en la toma de decisiones. En esta aplicación la decisión a tomar sería a qué tejido asignar cada píxel de la imagen entre los que son factibles de presentarse. Para la implementación de los operadores lógicos difusos se utilizan las operaciones que sugiere la **Lógica Difusa Compensatoria** (LDC), que ha demostrado ser altamente eficiente en este contexto.

La elección de IRM de cerebro para probar el sistema se fundamenta en que se dispone de gran cantidad de datos provenientes de simulaciones que permiten una evaluación exhaustiva de los resultados en condiciones variadas y con una sistematización que sería muy difícil de obtener en estudios reales, en los que por supuesto están involucrados pacientes. Igualmente se presentan los resultados en imágenes reales, evaluados por especialistas en imágenes.

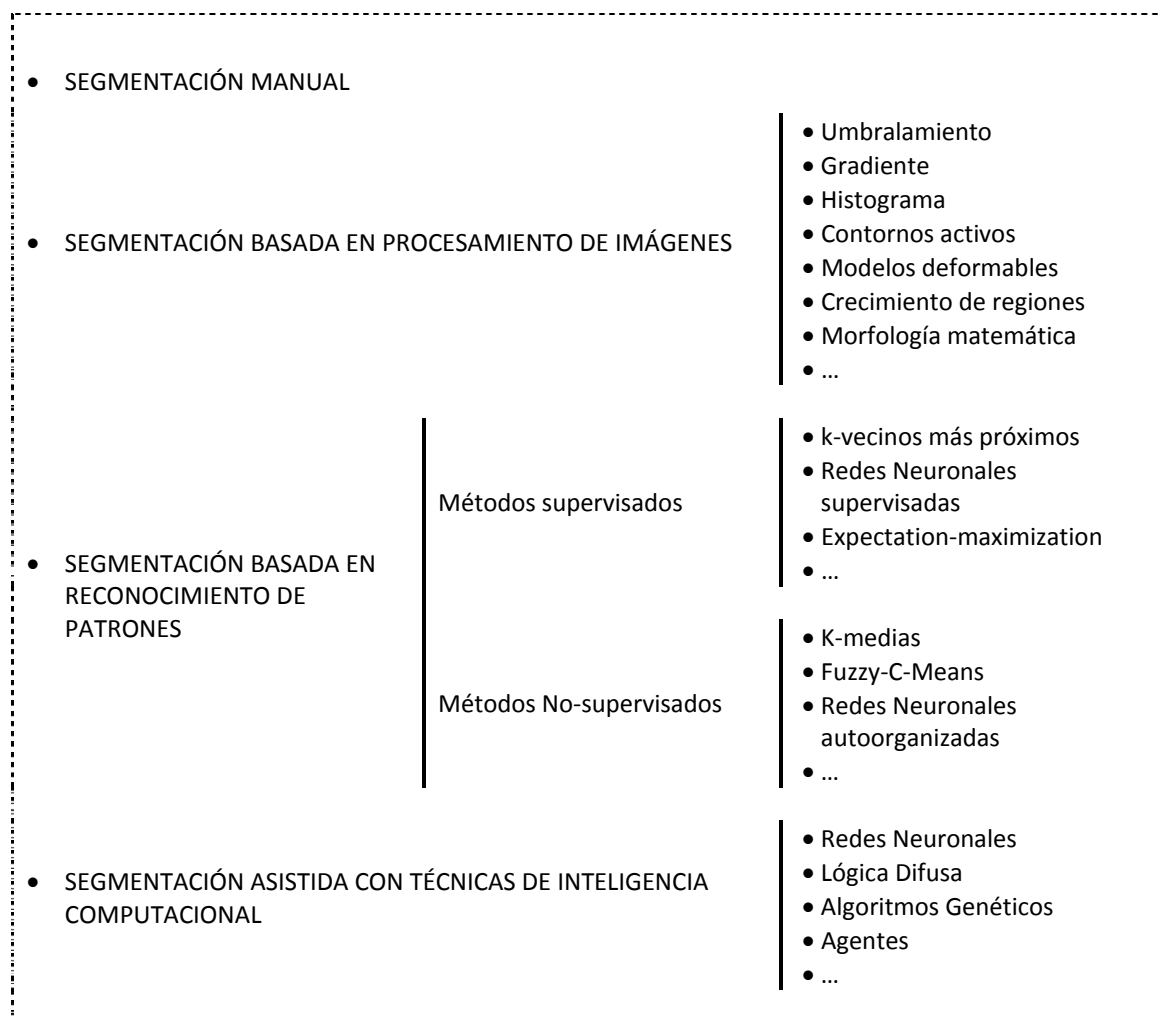
El sistema también se ensayó con algunas IRM de hombro y se presentan algunas pautas de cómo puede ser modificado para contemplar el procesamiento de otro tipo de imágenes.

### 1.3. Antecedentes de la problemática a abordar

Continuamente se proponen nuevas técnicas para lograr la discriminación de diferentes tipos de tejidos (o sustancias), cada una de ellas con sus ventajas y limitaciones. En general podría afirmarse que los diferentes métodos tienen en común el reconocimiento de diferentes tejidos mediante la interpretación de las imágenes que el equipo de **Resonancia Magnética** (RM) entrega.

Se realiza a continuación una reseña de los métodos que se han presentado, abordando específicamente el problema de la segmentación de IRM, procurando ubicarlos en una apropiada clasificación de las técnicas empleadas. Esta clasificación es compleja porque los mejores resultados se obtienen generalmente a través de la combinación de técnicas provenientes de diferentes paradigmas.

Pueden hallarse un interesante resumen de métodos de segmentación en general en [Pham et al., 2000] y en particular para el reconocimiento de tejidos en IRM en [Liew and Yan, 2006]. Basándose en estos trabajos, la Figura 1 muestra una posible categorización de los métodos que serán posteriormente descriptos.



**Figura 1: Métodos de segmentación de Imágenes.**

Se muestra una posible categorización de los algoritmos descriptos en la bibliografía, según los paradigmas de los que provienen.

### 1.3.1. Segmentación manual

La segmentación manual de imágenes es una tarea que, dependiendo la complejidad de las mismas, puede ser sumamente tediosa y requerir gran cantidad de tiempo, hasta varias horas por cada caso a analizar.

Las segmentaciones así obtenidas suelen ser subjetivas, dependientes del operador y sus resultados no son siempre repetibles. En algunos casos se requiere un gran conocimiento anatómico para no cometer errores. De todas maneras, una adecuada segmentación realizada por especialistas es una tarea necesaria para la evaluación de métodos de clasificación automáticos o semiautomáticos.

En segmentación de tejidos cerebrales en IRM es fundamental remover lo que no corresponda a la región del cerebro (principalmente cráneo y meninges). Aún esta tarea que pareciera ser sencilla puede llevar un tiempo considerable. La segmentación de los tejidos debe hacerse corte por corte y puede requerir más de dos horas llegar a un resultado.

### 1.3.2. Segmentación a través de métodos de procesamiento de imágenes

#### 1.3.2.1. Umbralamiento

El umbralamiento (*thresholding*) consiste en separar una imagen para convertirla en una o más imágenes binarias teniendo en cuenta rangos de niveles de gris. Es una técnica básica que no suele ser efectiva en sí misma pero se utiliza en diversas etapas de otros métodos. Existen métodos automáticos para hallar los umbrales, basados en intensidades locales o en conectividad de los píxeles [Sahoo et al., 1988].

No existen en la bibliografía trabajos que den cuenta de la aplicación directa de este método en la segmentación de IRM de cerebro debido a las limitaciones que conlleva la superposición de intensidades de grises que corresponden a los diferentes tejidos.

### 1.3.2.2. Métodos basados en gradiente

El gradiente de una imagen tiene propiedades similares a las de la derivada primera en una función unidimensional. Buscando los máximos en la imagen gradiente se podrán encontrar discontinuidades en la imagen original, lo que estará asociado a bordes de objetos de interés. La operación gradiente es fundamental para otros métodos de segmentación.

En [Jimenez-Alaniz et al., 2006] se combina la segmentación de regiones por medio de un clasificador bayesiano con la detección de bordes por medio de los valores normalizados del gradiente, para la segmentación de IRM de cerebro.

### 1.3.2.3. Métodos basados en el histograma de intensidades

Mediante el análisis del histograma de intensidades pueden hallarse valores críticos de gris para proceder a umbralamientos en varios niveles (*multithresholding*) [Glasbey, 1993]. Esta información es utilizada también con otros fines que no necesariamente sean la elección de valores umbrales [Bonnet et al., 2002] y ha sido también aplicada en imágenes color [Delon et al., 2005].

En IRM de cerebro se ha utilizado la información del histograma para complementar otras técnicas, como se verá, por ejemplo, en la siguiente sección.

### 1.3.2.4. Contornos Activos – Modelos Deformables

Los modelos deformables parten de una curva (2D) o volumen (3D) expresado de manera paramétrica, que es adaptado iterativamente hasta lograr representar un objeto o área de interés de la imagen. Una vez obtenidos los parámetros, éstos son útiles también para representar el objeto o región obtenida [McInerney and Terzopoulos, 1996].

En [Li et al., 2006] se presenta una aplicación de un modelo de contornos activos para segmentación de IRM de cerebro, demostrando ser robusta y eficiente. El método presentado es optimizado extrayendo información del histograma. Se requiere un pre-procesamiento para minimizar el ruido.

En [Angelini et al., 2007] se propone un método para IRM basado en un modelo deformable que se adapta según medidas de homogeneidad, sin requerir información

*a priori*, con el fin de discriminar los diferentes tejidos del cerebro. Se compara su desempeño con otras técnicas (utilizando modelos del Markov y algoritmos de clasificación, umbralamientos óptimos, etc.) mostrando una mejora. Se tiene en cuenta la necesidad de disminuir distorsiones y el método requiere la implementación numérica de ecuaciones diferenciales, con sus problemas y limitaciones.

#### **1.3.2.5. Crecimiento de Regiones**

Este método se basa en la extracción de una parte de la imagen que comparte características similares: textura, intensidades de gris, etc. Se parte de una pequeña región (semilla) y el proceso consiste en ir incorporando nuevas regiones en tanto cumplan un criterio de similitud [Fan et al., 2005].

Se han utilizado estos algoritmos en IRM para segmentar regiones de tumores, vasos sanguíneos, ventrículos o diferentes partes anatómicas, pero no resultan métodos adecuados cuando se trata de identificar diferentes tejidos [Chulho and Dong Hun, 2007].

#### **1.3.2.6. Morfología Matemática**

Desde hace ya varios años se estudia la aplicación de una teoría basada en el álgebra de conjuntos denominada Morfología Matemática [Serra, 1992]. Mediante operaciones no lineales básicas (erosión, dilatación) y sus combinaciones es posible resaltar o atenuar componentes de la imagen similares a una forma geométrica dada permitiendo analizar su forma, tamaño, orientación y superposición [Dougherty, 1993].

En IRM de cerebro se han aplicado filtros secuenciales alternativos, basados en las operaciones morfológicas, para la extracción de la región del encéfalo, quitando cráneo y meninges. Las aplicaciones básicas de estos filtros son la atenuación del ruido y la extracción selectiva de objetos en la imagen [Pastore et al., 2006]. Sin embargo estos métodos no son adecuados para la identificación de tejidos.

En Morfología Matemática se define la Transformada Watershed. Esta transformada realiza una partición de la imagen en regiones mediante la inundación de su gradiente visto como un relieve topográfico [Vincent and Soille, 1991], partiendo de los mínimos locales. A diferencia de los métodos convencionales, la partición de la

imagen se realiza en base a la similitud y discontinuidad de los niveles de gris de los píxeles simultáneamente [Soille and Pesaresi, 2002].

Si en lugar de comenzar la inundación por todos los mínimos se efectúa solamente a partir de un conjunto de regiones denominadas “marcadores”, se producen tantas regiones como marcadores se hayan definido [Serra, 1992].

En imágenes con texturas finas (como lo son las IRM) su gradiente posee poco contraste y definición resultando en regiones que no delimitan fielmente las regiones de interés. Sin embargo, se ha utilizado en IRM de cerebro para la detección de los ventrículos y para delimitación de tumores. En el reconocimiento de tejidos ha dado resultados aceptables con el agregado de información proveniente de atlas probabilísticos para la definición de marcadores [Grau et al., 2004].

### 1.3.3. Segmentación a través de Reconocimiento de Patrones

Un paradigma muy utilizado para la segmentación consiste en el agrupamiento o clasificación de los píxeles mediante diversas técnicas. Los píxeles quedan caracterizados por vectores numéricos que los identifican con algún criterio, siendo sus componentes los denominados característicos o descriptores. Estos vectores ingresan al algoritmo de clasificación.

En IRM se han propuesto diversos clasificadores aplicados a las intensidades de grises presentes solamente en un tipo de imagen (por ejemplo, solamente la imagen pesada en  $T_1$  o solamente la imagen pesada en  $T_2$ <sup>1</sup>), caso denominado “monoespectral” o combinando las intensidades en más de un tipo de imagen (por ejemplo  $T_1$ ,  $T_2$  y PD), caso “multiespectral”.

En [Bezdek et al., 1993] se efectúa una reseña de nueve algoritmos básicos de clasificación empleados específicamente en IRM.

Los **característicos**, también denominados descriptores, se extraen de las imágenes a segmentar. La selección de un buen conjunto de descriptores es una etapa clave en cualquier proceso de clasificación, en este caso para lograr una segmentación adecuada.

---

<sup>1</sup> La definición de estos tipos de imágenes se darán con detalle en el Capítulo 2.

### **Característicos basados en intensidad de gris**

Este enfoque suele utilizarse cuando se dispone de más de una imagen, como el caso de la RM, donde es posible construir un vector con más de una intensidad de gris que representa a cada píxel. Mediante estos vectores se procede a la clasificación de todos los píxeles. En imágenes color se han utilizado los vectores formados por las componentes de color (RGB, HSV, etc.) para alimentar a un clasificador.

### **Característicos basados en descriptores de texturas**

El concepto de textura ha sido ampliamente utilizado en los últimos años. La textura se define como una repetición de un patrón básico, que se ha denominado “texel”, que puede ser caracterizado por diferentes descriptores matemáticos. Se han presentado descriptores basados en características morfológicas o estadísticas (cómo se disponen espacialmente las intensidades de gris), frecuenciales (a través de la transformada de Fourier o de Wavelets), fractales (calculando la dimensión fractal de una región) [Chen et al., 1993] y otros [Todd and Buf, 1993]. Estos valores se arreglan en forma de vector para proceder a su clasificación.

Se han utilizado característicos de texturas en segmentación de IRM de cerebro en [Kovalev et al., 2001] con el fin de hallar patologías. Por ser una técnica de gran costo computacional es que ha sido superada por otras en el caso particular de este tipo de imágenes.

La Morfología Matemática también ha sido utilizada para obtener característicos de las texturas y proceder a su posterior clasificación [Ballarin, 2001]. Esta técnica funciona en IRM, pero no ha sido evaluada su robustez.

### **Característicos basados en atlas**

Algunas técnicas se basan en atlas anatómicos, que constituyen información *a priori* construida en base a casos previamente segmentados y que se disponen para estos fines. Su utilización requiere inicialmente un proceso de registración.

La registración es un proceso de transformación de un conjunto de datos en un determinado sistema de coordenadas, necesario para poder comparar o integrar los datos obtenidos de diferentes fuentes. En este caso, la imagen a segmentar debe ser registrada con la imagen respectiva del atlas a utilizar.

En recientes aplicaciones se emplean algoritmos simples [Schwarz and Kasperek, 2007] o algoritmos basados en la conectividad de los píxeles [Zhou and Bai, 2007] entrenados con los datos provenientes de atlas, haciendo especial hincapié en el proceso de registración.

### **Campos aleatorios de Markov**

La teoría de Modelos de Campos Aleatorios de Markov (MRF, *Markov Random Fields*) y de Modelos Ocultos de Markov (HMM, *Hidden Markov Models*) ha sido investigada desde los años 80 [Li, 1995a] pero su utilización en segmentación fue posterior [Zhang et al., 2001]. Los HMM son modelos probabilísticos de la probabilidad conjunta de una colección de variables aleatorias con sus observaciones y estados. En estos modelos, la información espacial de una imagen se codifica mediante restricciones contextuales de píxeles vecinos. Imponiendo esas restricciones, se espera que píxeles vecinos tengan la misma etiqueta o intensidades similares. Esto se logra caracterizando las influencias mutuas entre píxeles usando distribuciones condicionales de tipo MRF. Los parámetros del modelo suelen estimarse mediante el algoritmo EM. Se han aplicado en IRM cerebro [Held et al., 1997] y continúa su utilización con nuevas variantes [Pyun et al., 2007, Ibrahim et al., 2006].

#### **1.3.3.1. Métodos de clasificación supervisados**

Los algoritmos de clasificación supervisados requieren un conjunto de píxeles de ejemplo previamente identificados o etiquetados (también llamado “conjunto de entrenamiento”). Sus parámetros internos son modificados, generalmente en forma iterativa, mediante la influencia de los vectores de descriptores de los píxeles de entrenamiento y sus respectivas etiquetas.

#### **K-Nearest Neighbours (kNN)**

Este método requiere una determinada cantidad de píxeles previamente etiquetados, denominados prototipos. Cuando un nuevo píxel debe ser clasificado, se compara con los prototipos a través de una medida de distancia, generalmente euclidiana. Se evalúan las clases a la que corresponden los K prototipos más cercanos y se asigna al nuevo dato la clase que más apareció entre ellos.

Se han propuesto modificaciones para hacer este algoritmo más eficiente en el procesamiento de imágenes. En [Warfield, 1996] se sugiere una “transformada distancia” para calcular y evaluar los k-vecinos más próximos a los datos de entrenamiento, compuestos por 250 píxeles de cada clase. Se aplica en IRM de cerebro y se muestra una mejora en la eficiencia comparando con el algoritmo original.

Cocosco describe un método totalmente automático, robusto y adaptivo para la clasificación de tejidos en estudios de cerebro 3D de RM. El proceso parte de un conjunto de datos de entrenamiento que es optimizado utilizando una técnica de “podado” (prunning). Primeramente se reduce la fracción de vóxeles incorrectamente detectados usando un árbol de grafos. Luego, las muestras correctas alimentan un clasificador kNN para clasificar la imagen 3D completa. No se requiere conocimiento *a priori* de las distribuciones de intensidades de los tejidos (procedimiento no paramétrico), lo que le confiere al proceso una gran robustez [Cocosco et al., 2003].

### **Redes Neuronales Supervisadas**

Puede encontrarse una gran cantidad de trabajos basados en Redes Neuronales (RN) entre los años 1995 y 1999, si bien su utilización no ha cesado hasta el presente. Una red neuronal es un arreglo de procesadores básicos, las “neuronas”, que se encuentran interconectados por pesos, las “sinapsis”, adaptables iterativamente para reflejar un conjunto de datos de entrada y sus respectivas salidas. Las RN necesitan una primera etapa de entrenamiento, generalmente costosa computacionalmente, para luego ser consultadas para obtener salidas ante nuevas entradas, etapa mucho menos exigente computacionalmente. Durante la etapa de entrenamiento debe cuidarse que la red no “aprenda” demasiado bien los datos de entrenamiento para no perder su capacidad de generalización [Egmont-Petersen et al., 2002]. El problema de hallar una arquitectura apropiada a un determinado problema no es sencillo y sólo se dispone de reglas heurísticas. La arquitectura puede ser optimizada por métodos adicionales como por ejemplo, los AG, que son también utilizados en esta tesis.

La capacidad de generalización de diferentes tipos de RN con el fin de la clasificación de los píxeles según sus características ha sido ampliamente utilizada, entrenando las redes con ejemplos y consultándola para el resto de la imagen o para otras imágenes estadísticamente similares para el caso de IRM [Song et al., 2007].

En [Zijdenbos et al., 2002] se utilizan RN con seis entradas, una capa de neuronas escondida y una salida. Se entrenan con imágenes de referencia y mapas de probabilidades *a priori* de los tejidos cerebrales. La salida de la red (una única neurona) indica la presencia o no de una lesión. No se obtiene una segmentación de los tejidos.

### **Expectation – Maximization / Mixturas Gaussianas**

El algoritmo “Expectation – Maximization” (EM) es un método de dos pasos aplicable cuando una parte de los datos es observable. El primer paso (E) asume que el modelo de mixturas gaussianas actual es correcto y halla la probabilidad de que cada dato pertenezca a cada gaussiana. El segundo paso (M) mueve las gaussianas para maximizar su verosimilitud. Es decir, cada gaussiana “toma” los puntos que el paso anterior le asignó probabilísticamente. El algoritmo EM es un método que se usa frecuentemente para ajustar una mezcla de modelos gaussianos a un conjunto de datos. Si bien se asocia este método a la clasificación, en realidad consiste en una etapa preliminar en la que se obtienen funciones densidad de probabilidad que caracterizan a los datos.

Para la segmentación el cerebro, el modelo queda expresado en base a las siguientes premisas: “hay 3 clases de tejidos: materia gris, materia blanca y líquido cefalorraquídeo”, “los tejidos tienen una distribución gaussiana” y “las no-uniformidades de la intensidad se representan por un campo multiplicativo”. Así aplicado, el algoritmo asigna iterativamente las clases de tejidos a los píxeles (paso “E”) y un campo multiplicativo (paso “M”). La salida del algoritmo es una clase para cada píxel y una estimación del campo de no-uniformidad. El algoritmo puede ser inicializado con una distribución *a priori* de las clases mediante un atlas [Wells et al., 1996] y se ha utilizado combinado con otros métodos [Ashburner and Friston, 2005].

La determinación del modelo tiene un costo computacional importante en comparación con otros algoritmos.

En [Pohl et al., 2004] se presenta una clasificación jerárquica de las estructuras cerebrales. Como ejemplo, se propone primero la segmentación en regiones de cerebro y no-cerebro y luego la segmentación del cerebro en sus diferentes tejidos,

mediante la aplicación del algoritmo EM, incluyendo la información de un atlas de probabilidades *a priori* de las clases.

### 1.3.3.2. Métodos de clasificación no supervisados

En los algoritmos de clasificación no supervisados los píxeles se auto-agrupan iterativamente, sin una especificación previa de clases o etiquetas. En este caso los parámetros internos del sistema de clasificación se adaptan a través de la optimización de una medida de calidad de la representación que se pretende.

#### K-medias

Este método requiere conocer *a priori* la cantidad de clases en que se separarán los datos. Al comenzar se eligen centros de *cluster* iniciales, pertenecientes al espacio de los datos a clasificar (elegidos de los propios datos o aleatoriamente). Se asignan los datos a cada *cluster* según la distancia a su centro sea la mínima. Luego los centros de *cluster* son modificados en sucesivas iteraciones promediando las componentes de los datos que corresponden a cada uno. Cuando los centros no cambian su posición el algoritmo converge y finaliza.

Se han propuesto numerosas modificaciones al algoritmo para adaptarlo a diversas aplicaciones [Jian, 2005]. En [Abrás and Ballarín, 2005] se presenta una modificación que tiene en cuenta la cantidad de veces que aparecen determinados patrones de grises constituidos por las intensidades en T1, T2 y PD en IRM.

#### Fuzzy-C-Means

Este algoritmo es similar al K-medias pero asigna a cada dato una pertenencia difusa a los *clusters*, en el rango 0 a 1. Un algoritmo de clasificación difuso no necesariamente utiliza conceptos de conjuntos difusos tales como funciones de pertenencia y operaciones difusas [Baraldi and Blonda, 1999]. El algoritmo **Fuzzy-C-Means** (FCM) minimiza una función de costo basada en las distancias difusas de los datos a los centros de *cluster*, modificando iterativamente la ubicación de los centros.

Se han propuesto modificaciones específicas para segmentar IRM con este algoritmo [Siyal and Yu, 2005, Kang et al., 2008, Szilágyi et al., 2007, Salvado et al., 2007]. La mayoría de ellos lo utilizan como esquema de segmentación inicial para

luego continuar el procesamiento con otros procesos [Clark et al., 1994]. En [Supot et al., 2007] se sugiere un método para acelerar la convergencia de este algoritmo.

Yang y otros autores [Yang et al., 2007] presentan un método de Cuantificación Vectorial que denominan “fuzzy-soft LVQ” que se desempeña satisfactoriamente en la segmentación de regiones en IRM. Se parte del FCM y se incorpora la información de este algoritmo en otro del tipo LVQ (*Learning Vector Quantization*). Se debe indicar la cantidad de tejidos o sustancias (*clusters*) a detectar. Si bien es un algoritmo automático (los valores iniciales de centros de *clusters* se seleccionan arbitrariamente), cada tipo de estudio implicaría cambiar la cantidad de *clusters* a separar y los tipos de tejido a los que correspondería cada *cluster*. Los tejidos son separados en grupos y un experto debe etiquetar manualmente los grupos.

### **Redes Neuronales Autoorganizadas**

A diferencia de las RN supervisadas, este tipo de redes intenta autoorganizar los datos según su similitud, según un criterio de distancia entre ellos. Una de las redes más conocida son los **mapas autoorganizados de Kohonen** (SOM, *Self Organizing Maps*) [Kohonen, 2001]. Los datos se muestran en una grilla bidimensional que guarda la misma topología que el conjunto de datos multidimensional y que puede ser utilizada para agrupar los datos y obtener así una clasificación según sus similitudes naturales. Se han aplicado en el procesamiento de imágenes para obtener pseudocolores [Meschino et al., 2006] y también específicamente en la segmentación de IRM de cerebro [Tian and Fan, 2007].

### **1.3.4. Segmentación con técnicas de Inteligencia Computacional**

Las técnicas de Inteligencia Computacional pueden contribuir en varios aspectos: optimizan la clasificación de los píxeles, agregan conocimiento del experto, guían los contornos activos, etc. Numerosas técnicas de clasificación surgen del paradigma de la Inteligencia Computacional, de las cuales ya se han citado algunas.

La combinación de dos o más técnicas de IC da lugar a los denominados “modelos híbridos”, de gran aplicación en diversos campos y en particular en la

segmentación de imágenes [Castellanos and Mitra, 2000]. Por ser éste un tema de suma importancia en esta tesis, se abordará posteriormente con más detalle.

#### 1.3.4.1. Redes Neuronales

Ya se ha presentado el uso de las RN como método de clasificación tanto supervisado como no supervisado. Diferentes tipos de redes suelen ser combinadas entre sí y con otros algoritmos.

Di Bona [Di Bona et al., 2003] presenta un enfoque para la clasificación de densidades de tejidos cerebrales en estudios tridimensionales basado en RN jerárquicas. Un grupo de neurólogos seleccionó estudios normales y patológicos para probar el método. En una primera etapa se utiliza un SOM para cada una de las características de los píxeles (esto permite que cada SOM pueda ser mejorado sin afectar las otras características). Luego se utiliza una red multicapa para la clasificación final. Los característicos utilizados en la clasificación son los valores de gris con sus posiciones en el contexto 3D, diferencias de contraste con respecto al valor medio de gris y valores de gradientes. Si bien constituye una interesante aplicación, el entrenamiento de todas las RN involucradas requiere una considerable complejidad en su implementación y un costo computacional importante.

Shen [Shen et al., 2005] aplica una arquitectura de tipo red neuronal para optimizar un parámetro de la metodología propuesta: el grado de “atracción” entre píxeles vecinos, que depende de sus características y ubicación. Se propone una extensión del algoritmo FCM que mejora su desempeño. Se prueba el algoritmo en imágenes simuladas con diferentes intensidades de ruido para demostrar su robustez. La función de costo es modificada *ad hoc*. No emplea ningún tipo de conocimiento *a priori* ni requiere intervención de un experto en imágenes.

En un reciente trabajo [Song et al., 2007] se propone el uso de una Red Neuronal Probabilística que ha sido modificada para segmentar imágenes similares a las que se presentan en este trabajo. Se utilizan también SOM para la estimación de funciones densidad de probabilidad. Finalmente se aplica un mecanismo de etiquetado basado en la regla de Bayes, basándose en los vectores prototipo que se obtienen en el SOM. Se prueba el método con diferentes niveles de distorsión en imágenes simuladas

y reales. El método mejora el desempeño de otros con los que es comparado, si bien su complejidad en cuanto a los entrenamientos de las redes que requiere hace que sea exigente computacionalmente.

#### 1.3.4.2. Sistemas basados en Lógica Difusa

La LD es un paradigma que extiende la lógica tradicional, permitiendo valores de verdad de sentencias y pertenencias a conjuntos con rangos continuos en  $[0, 1]$ . Será tratada con detalle más adelante. Permite traducir el lenguaje natural en un algoritmo que utiliza un modelo matemático.

Los modelos más utilizados son llamados de sistemas de inferencia difusa (FIS, *fuzzy inference systems*), basados en un conjunto de reglas de tipo IF – THEN (SI – ENTONCES). Los diferentes FIS existentes priorizan separadamente la interpretabilidad y la exactitud del modelo.

En el **modelo difuso de Mamdani** [Mamdani, 1974] las reglas se diseñan en base al conocimiento experto. Las reglas relacionan diferentes conceptos asociados a variables “de entrada” por medio de conjuntos difusos con los asociados a variables “de salida”, siendo su estructura, para la regla  $j$ -ésima, suponiendo una sola variable de salida:

$$R_j: IF x_1 \text{ is } CE_1^{(j)} AND x_2 \text{ is } CE_2^{(j)} AND \dots AND x_n \text{ is } CE_n^{(j)} THEN y \text{ is } CS^{(j)},$$

donde  $x_i$  son las variables de entrada,  $y$  es la variable de salida,  $CE_i^{(j)}$  son conjuntos difusos de la variable de entrada  $i$  utilizados en la regla  $j$ ,  $CS^{(j)}$  es un conjunto difuso de la variable de salida utilizado en la regla  $j$  y  $n$  es la cantidad de entradas.

En estos sistemas se definen etapas del procesamiento:

- la *fuzzificación*, en la cual los valores de las entradas (utilizadas en los antecedentes) se transforman en pertenencias a los distintos conjuntos difusos,
- la *implicación*, en la cual se determina un conjunto difuso como salida de cada regla,
- la *agregación*, en la cual se combinan los conjuntos salida de todas las reglas en un solo conjunto difuso,

- y la *defuzzificación*, en la que se convierte el conjunto difuso final en un valor para la variable de salida.

En principio, los FIS de Mamdani se diseñan únicamente a partir del conocimiento de un experto de campo y no a partir de un conjunto de datos de entrada – salida, si bien se han propuesto variantes en este sentido. Su principal fortaleza es la interpretabilidad del modelo generado [Alcalá et al., 2006].

El **modelo difuso de Sugeno** (también conocido como de Takagi–Sugeno–Kang, TSK) [Takagi and Sugeno, 1985, Sugeno and Kang, 1988] está basado en reglas cuyo consecuente no es una variable lingüística sino una función lineal de las variables de entrada. La regla *j*-ésima presenta la siguiente estructura:

$$R_j: IF x_1 \text{ is } CE_1^{(j)} \text{ AND } x_2 \text{ is } CE_2^{(j)} \text{ AND } \dots \text{ AND } x_n \text{ is } CE_n^{(j)} \\ THEN y_j = f_j(x_1, x_2, \dots, x_n),$$

donde  $x_i$  son las variables de entrada,  $y_j$  es el valor de salida de la regla  $j$ ,  $CE_i^{(j)}$  son conjuntos difusos de la variable de entrada  $i$  utilizados en la regla  $j$ ,  $n$  es la cantidad de entradas y  $f_j$  es una función matemática de las variables de entrada.

Cada regla describe un subespacio difuso del espacio de entrada-salida a través de las funciones  $f_j(x_1, x_2, \dots, x_n)$ . La característica de estos subespacios es que la pertenencia de un dato a ellos es un valor real en el intervalo  $[0,1]$  definido a través de funciones de pertenencia a los conjuntos difusos. Para determinarlos se pueden utilizar distintos métodos, generalmente de agrupamiento difuso. El valor de salida  $y$  se obtiene como un promedio pesado de las salidas  $y_j$  según el valor de verdad de los antecedentes.

Los FIS de Sugeno se determinan mediante operaciones sobre un conjunto de datos de entrenamiento (pares entrada – salidas esperadas). Su principal fortaleza es la exactitud para modelar el conjunto de datos y la capacidad de generalización [Alcalá et al., 2006].

Se han aplicado sistemas que utilizan reglas difusas con diversos fines de segmentación, por ejemplo, la extracción del cerebro completo de la IRM [Hata et al., 2000].

En [Chang et al., 2002] se propone un algoritmo que comienza con un conjunto de reglas difusas para definir los distintos tejidos, incorporando funciones de pertenencia adaptivas según los histogramas de grises. Los píxeles que no pudieron ser clasificados pasan por una segunda etapa basada en el algoritmo FCM al que se le propone una modificación. Los resultados se presentan únicamente en base a evaluación visual de expertos y no se ha realizado un estudio comparativo con medidas objetivas de la calidad de la segmentación obtenida.

Denkowski y colaboradores [Denkowski et al., 2004] crearon un sistema difuso de segmentación de IRM que consiste en dos etapas principales: primero, un análisis de histograma basado en la función de pertenencia  $S$  (que se explicará más adelante) y la entropía de Shannon; luego, la clasificación de los píxeles a través de un sistema de inferencia difuso basado en reglas.

Kobashi y demás autores [Kobashi et al., 2006] utilizaron LD para determinar de manera semiautomática los límites de los lóbulos cerebrales en IRM. Un operador debe indicar superficies iniciales aproximadas para los lóbulos, las que son suavizadas y deformadas mediante contornos activos y modelos de superficie. Ambos modelos son asistidos por una base de reglas difusas que ha sido elaborada por expertos con el fin de describir los límites de las diferentes regiones de interés a segmentar. Utilizan las operaciones MIN y MAX para los conectivos AND y OR respectivamente y la técnica de inferencia de Mamdani.

#### **1.3.4.3. Algoritmos Genéticos**

Los AG son un método de optimización que intentan minimizar una función de costo por medio de una búsqueda dirigida por operadores genéticos como son el cruzamiento y la mutación. Se explican en detalle en el *Apéndice A* y forman una parte esencial del método propuesto en este documento.

Se han empleado AG para guiar las iteraciones y automatizar métodos basados en crecimiento de regiones [Chun and Yang, 1996] y para seleccionar el conjunto de entrenamiento en el método kNN [Kuncheva, 1995].

En [Sasikala et al., 2006] se utilizan AG con el fin de optimizar una nueva función objetivo del algoritmo FCM que tiene en cuenta.

Más recientemente se han utilizado AG en [Jinn-Yi and Fu, 2008] para poder automatizar la cantidad de grupos en que se quiere realizar una segmentación de una imagen multiespectral.

Se han aplicado para solucionar el problema de la elección del mejor conjunto de características (*feature selection*) en un sistema de clasificación o dividir las características en subconjuntos [Rokach, 2008].

#### **1.3.4.4. Sistemas basados en Agentes**

El paradigma de "agentes" aborda el desarrollo de entidades que puedan actuar de forma autónoma y razonada. La Inteligencia Computacional intenta construir dichas entidades como si fueran "inteligentes".

En [Germond et al., 2000] se combina un método basado en agentes, un modelo deformable, crecimiento de regiones y un detector de bordes.

Se han utilizado también sistemas basados en agentes para complementar el algoritmo FCM y guiar las iteraciones del crecimiento de regiones [Haroun et al., 2006].

Richard [Richard et al., 2004] utiliza el concepto de agentes inteligentes, los que son adaptados dinámicamente dependiendo de su posición en la imagen, su relación topográfica y la información radiométrica disponible. En su trabajo presenta una comparación con otras técnicas utilizando imágenes simuladas y reales.

Se han aplicado también métodos que provienen de otras disciplinas, como el descubrimiento del conocimiento, para objetivos más específicos, como es la discriminación de diferentes tipos de patologías [Siromoney et al., 2000].

### **1.4. Objetivos y aportes de esta tesis**

---

Según lo expuesto, existen numerosos enfoques para la segmentación de imágenes y en particular de RM de cerebro.

Sin embargo, el hecho de que siguen proponiéndose nuevos paradigmas indica que ninguno de ellos es autosuficiente ni ha solucionado el problema en su totalidad.

Los métodos para detectar tejidos cerebrales hallados en la bibliografía suelen ser específicamente desarrollados para una aplicación particular y por consiguiente no pueden ser adaptados para funcionar exitosamente con otros tipos de imágenes. Por otra parte, excepto en algunos trabajos de los últimos años, no se hace un análisis del desempeño de los métodos propuestos en condiciones de ruido y otras distorsiones, o bien se presentan resultados sólo en imágenes simuladas, los cuales no siempre pueden ser extrapolados a imágenes reales. Aún en el caso de presentar resultados en imágenes reales, no se asegura la adaptabilidad del método a imágenes obtenidas en más de un equipo.

Los trabajos presentados que utilizan LD suelen tratar con sistemas de inferencia difusa y no con simples predicados. Para realizar las operaciones matemáticas suelen emplear los operadores MAX y MIN para implementar los conectivos AND y OR respectivamente.

El **objetivo principal** de la presente tesis es: el diseño, desarrollo, validación y comparación de un sistema de clasificación y segmentación de los distintos tejidos cerebrales en IRM de cerebro, basado en predicados difusos, que:

- incorpore conocimiento de los especialistas que pueda ser expresado en los predicados;
- utilice la menor cantidad de información *a priori* y que no se base en atlas anatómicos;
- se adapte a imágenes provenientes de diversos equipos de RM;
- sea robusto al ruido y la **no-uniformidad en la intensidad** (INU, *Intensity non-Uniformity*);
- sea modificable para otros tipos de imágenes;
- iguale o supere la calidad de la segmentación obtenida con otros métodos hallados en la bibliografía.

La realización del mismo se verá reflejada en los siguientes **aportes**:

- Un método de segmentación sencillo y eficiente computacionalmente, que iguale o supera en desempeño a los métodos presentados en los últimos años.

- La incorporación de los operadores de la LDC para implementar las operaciones difusas.
- La exploración de la inclusión de predicados difusos en la función de evaluación de un AG, con el fin de maximizar su valor de verdad.

## 1.5. Trabajos previos propios

---

En relación con las líneas de Inteligencia Computacional y segmentación de imágenes que dieron origen a experiencias anteriores a esta tesis pueden mencionarse los siguientes **trabajos propios**:

- Mediante características provenientes del análisis de texturas se segmentaron imágenes microscópicas de biopsias de médula ósea [Meschino and Moler, 2004a]. se segmentaron las mismas mediante herramientas de Morfología Matemática [Pastore et al., 2006].
- Se ha presentado un esquema de procesamiento de RM de cerebro con Redes Neuronales de Regresión generalizada [Ballarin et al., 2005]. Las mismas redes se han utilizado para la clasificación de características de texturas [Meschino et al., 2004b].
- Para la visualización de características de la imagen, se ha desarrollado una representación pseudocolor con Mapas Autoorganizados de Kohonen [Meschino et al., 2006].
- Se ha empleado un Sistema de Inferencia Difuso en la obtención de marcadores para la transformada Watershed [Gonzalez et al., 2007].
- Se utilizó un modelo difuso optimizado con AG con el fin de estimar la “edad arterial” según un supuesto de coherencia entre la morfología de la señal de pulso y la edad cronológica del individuo [Scandurra et al., 2007].

El fundamento de esta tesis ha sido presentado en el Congreso de Matemática Aplicada, Computacional e Industrial (MACI) 2007, publicado posteriormente en la Serie Mecánica Computacional [Meschino et al., 2007].

## 1.6. Estructura de la tesis

---

A continuación y a modo de guía de la lectura del documento, se describe la estructura de la tesis y los contenidos esenciales de cada capítulo.

En el Capítulo 2 (“Imágenes de Resonancia Magnética”) se presentan los fundamentos físicos y tecnológicos de la Resonancia Magnética, con el objetivo de comprender el origen de los diferentes tipos de imágenes que se utilizarán en esta tesis. También se presentan aplicaciones de la cuantificación de las imágenes de Resonancia Magnética de cerebro para la detección de ciertas enfermedades.

En el Capítulo 3 (“Lógica Difusa”) se introduce la Lógica Difusa en forma general y se definen los conjuntos difusos y sus operaciones como generalización de la Lógica Booleana. Se presenta a la Lógica Difusa como herramienta adecuada para problemas de asistencia en la toma de decisiones, para lo que se definen las expresiones difusas y las operaciones que las relacionan. Se presenta la propuesta de la Lógica Difusa Compensatoria, con sus propiedades y ventajas. Finalmente se explica la manera en la que puede trabajarse con modelos de expresiones lingüísticas a través de predicados difusos y se da un ejemplo de su utilización.

En el Capítulo 4 (“Procesamiento de Imágenes de Resonancia Magnética: método propuesto”) se detalla la propuesta central de esta tesis. Se explica minuciosamente el método propuesto como alternativa a los ya existentes. Se describe también la manera en que se evalúa el método a través de un coeficiente para probar su adecuado funcionamiento y robustez, tanto en imágenes reales como en imágenes provenientes de simulaciones.

En el Capítulo 5 (“Resultados”) se muestran las imágenes segmentadas obtenidas con imágenes de prueba simuladas y con imágenes reales obtenidas de dos equipos diferentes. Se grafican los resultados de las medidas de calidad de la segmentación obtenidos cuando se utiliza el sistema en diferentes condiciones de ruido y distorsión. Se analizan los resultados obtenidos y se comparan las características del método propuesto con otras técnicas y con los resultados reportados por otros autores.

En el Capítulo 6 (“Conclusiones”) se presentan interesantes líneas de trabajo futuro y se resume todo lo expuesto.

---

## 2. Imágenes de Resonancia Magnética

## 2. Imágenes de Resonancia Magnética

---

### 2.1. Introducción

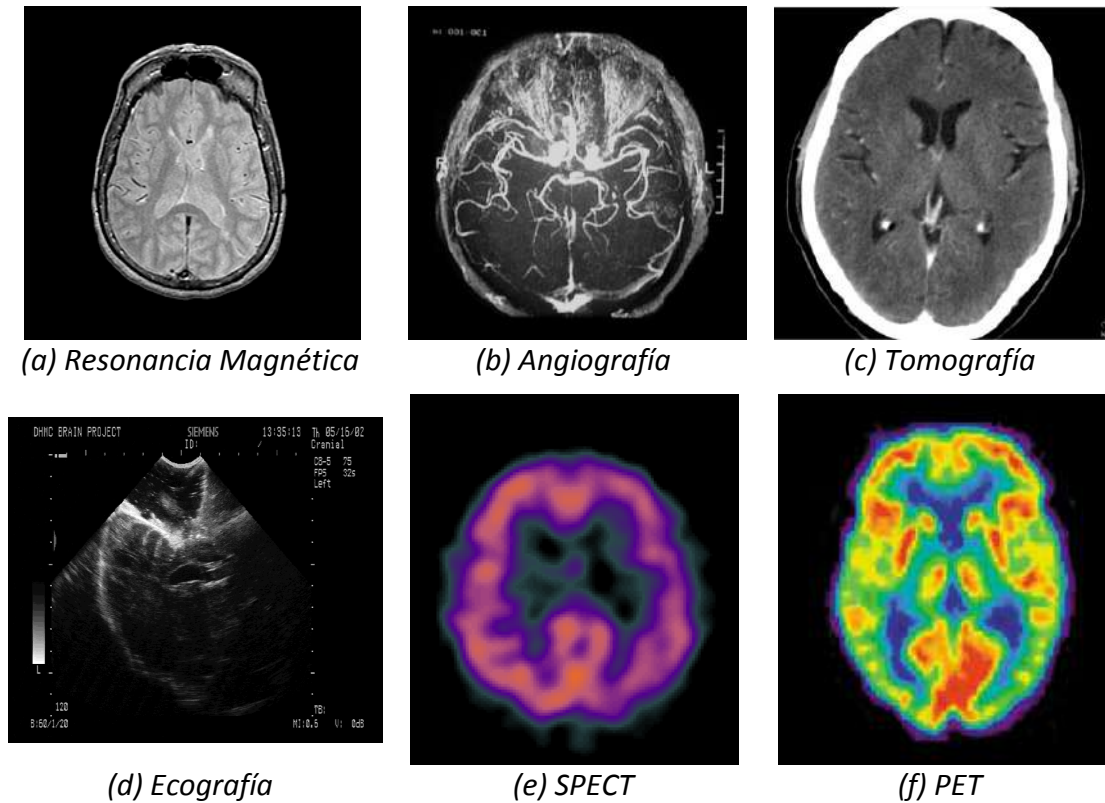
---

Desde su inclusión en 1971, los sistemas de Imágenes por RM ocupan un lugar muy importante en la larga lista de procedimientos que conforman la especialidad del Diagnóstico por Imágenes, como los que se muestran en la Figura 2. A diferencia de muchos de los otros métodos, la RM requiere del especialista un conocimiento relativamente completo de los fenómenos que suceden cuando diferentes tejidos se ubican en poderosos campos magnéticos y al mismo tiempo se someten a la acción de ondas de radiofrecuencia.

La física de la RM es muy compleja y existen numerosos textos y obras que ayudan a la comprensión de los fenómenos en ella involucrados [Blink, 2004, Coussement, 2000, Edelman et al., 1996, Hornak, 2007]. En este capítulo se describen las bases del tema con el fin de entender el porqué de los procesos y orientar a un mejor aprovechamiento de la técnica en los diferentes tejidos.

Sin embargo, dado que no constituye el objetivo de esta tesis, no es posible abarcar en este documento la enorme cantidad de avances desarrollados desde el nacimiento de la RM. Todo esto no hace más que confirmar la vertiginosa rapidez de los progresos tecnológicos. Para profundizar en estos deberá recurrirse a trabajos especializados [Storey, 2006].

En radiología convencional y en **tomografía axial computarizada** (TAC), el hecho de que las estructuras aparezcan blancas, negras o grises es de fácil comprensión y se origina en una ley muy simple: "la intensidad de la imagen es proporcional a la intensidad de rayos X que los tejidos atravesados han dejado pasar". En RM, un tejido cualquiera, por ejemplo un líquido, puede aparecer blanco o negro según los parámetros escogidos. Esto complica particularmente la comprensión y análisis de las imágenes, frente a lo cual no hay otra solución que tratar de entender lo que sucede [Coussement, 2000].



**Figura 2: Diferentes tecnologías para Diagnóstico por Imágenes.**

Se muestran a modo de ejemplo las imágenes obtenidas por diferentes equipos, cuyos fundamentos físicos son diferentes. SPECT: *Single Photon Emission Computed Tomography*, Tomografía computarizada por emisión de fotón único; PET: *Positron Emission Tomography*, tomografía de emisión de positrones.

La RM está basada en principios físicos totalmente diferentes de los de la TAC; la RM coincide con la TAC en que una energía es radiada dentro del paciente, pero presentan una diferencia fundamental a la hora de detectar la energía que produce la imagen: en la TAC la energía detectada es la remanente que el paciente no ha absorbido; mientras que en la RM esa energía se detecta al emerger del propio paciente [Liang et al., 2000].

En el caso de la RM la energía radiada consiste en ondas de radiofrecuencia en lugar de rayos X como ocurre en la TAC. Al igual que en la TAC, la energía detectada a la salida se correlaciona con varios parámetros característicos del tejido.

La gran ventaja de la RM es que es capaz de diferenciar tejidos con densidades muy próximas entre sí o incluso similares (densidades radiológicas) si su composición es distinta. Ello la hace muy superior a otros métodos de imagen, sobre todo en el estudio del sistema nervioso central y del sistema musculoesquelético, donde prácticamente todas las densidades -salvo el hueso- están comprendidas entre la grasa

y el agua. Además, la capacidad para diferenciar tantos tejidos permite obtener unas imágenes con un detalle anatómico excepcional y en cualquier plano del espacio.

## 2.2. Fundamentos de la Resonancia Magnética

---

### 2.2.1. Bases físicas

La materia está compuesta de moléculas y átomos; los átomos a su vez contienen núcleos con diferente número de protones y neutrones. Cuando las partículas cargadas giran o se mueven, generan alrededor de ellas un campo o momento magnético asociado, como resultado de las propiedades de giro (spin) de sus nucleones. Incluso los neutrones, que no tienen carga neta, contribuyen al spin.

La intensidad del campo magnético asociado de determinado núcleo depende del grado en que se suman o anulan entre sí los campos magnéticos de los nucleones individuales. La RM depende absolutamente de la existencia de un momento magnético neto del núcleo. De los núcleos de importancia biológica, el de hidrógeno proporciona la mayor sensibilidad de todos. Esto es debido tanto a su momento magnético elevado como a su gran abundancia en el cuerpo; ya sea en forma de agua como de otras moléculas de importancia biológica. Se entiende así por qué la RM es más sensible que la TAC [Bradley et al., 1984]: las diferencias en el contenido de agua, que a su vez implican diferencias en el contenido de hidrógeno en los tejidos biológicos, son mucho mayores que las diferencias correspondientes a la atenuación de los rayos X.

En un núcleo que gira, existe un vector magnético nuclear alineado en forma perpendicular al plano de giro. En un medio donde todos los núcleos tienen sus vectores alineados al azar en todas las direcciones, debido a las interacciones aleatorias con los núcleos vecinos, estos vectores se suman y anulan entre sí (Figura 3). Esto implica que no existe un campo electromagnético externo del medio considerado globalmente: los vectores nucleares se cancelan entre sí. Sin embargo presentan una susceptibilidad magnética potencial, lo que ocurre también en los tejidos biológicos.

Si el medio es colocado en un campo magnético externo fuerte  $B_0$ , los vectores magnéticos nucleares tienden a alinearse con este campo magnético externo [Blink,

2004]. El hidrógeno tiene dos formas de alineación posible: en la misma dirección del campo magnético externo, alineamiento paralelo, y en la dirección opuesta, alineamiento anti-paralelo (Figura 4).

El estado de alineamiento paralelo es el de menor energía; esto implica que un pequeño exceso de protones se alinea en esta dirección. A causa del pequeño exceso de alineamiento paralelo, permanece cierta magnetización neta en esa dirección, indicada como  $M_z$ . La fuerza de esa magnetización varía con la intensidad del campo  $B_0$ . En realidad los vectores no se alinean exactamente de modo paralelo o anti-paralelo, sino que efectúan un movimiento de precesión alrededor de la dirección del campo externo (Figura 4b).

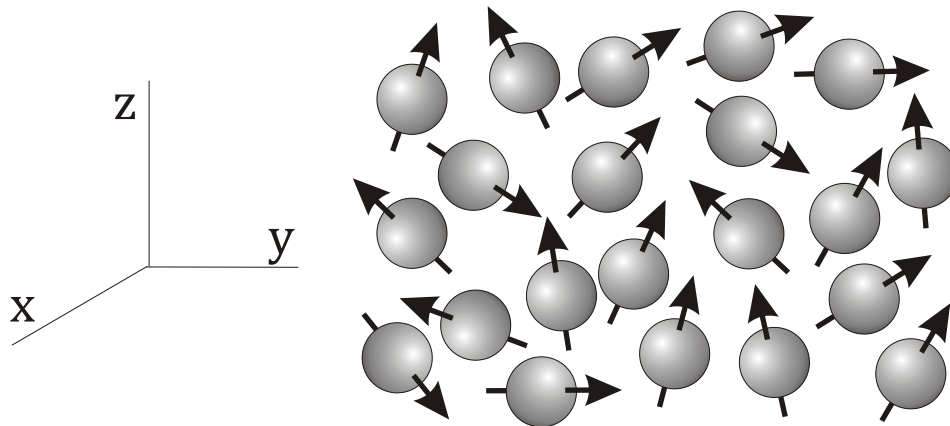
Este movimiento de precesión tiene una frecuencia que depende de  $B_0$  regida por la ecuación de Larmor:

$$f = k B_0 ,$$

donde  $k$  es una constante para cada núcleo llamada constante giromagnética. Esta constante es para el hidrógeno mayor que para cualquier otro núcleo de importancia biológica.

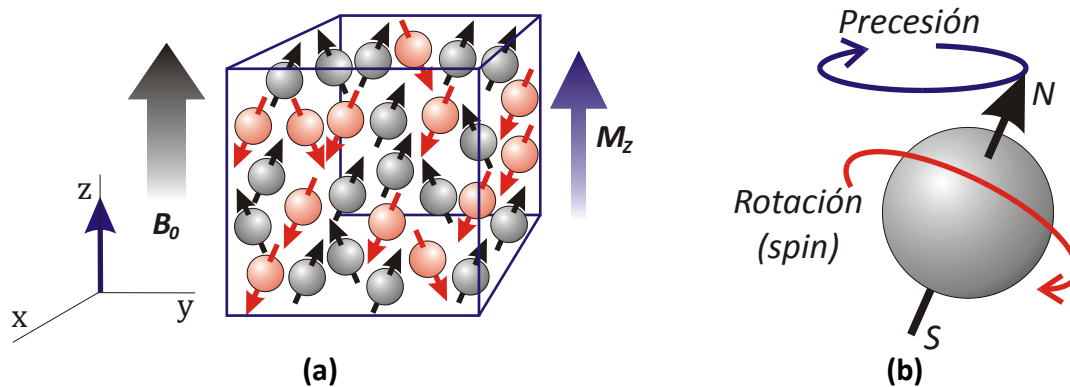
El conjunto de los vectores magnéticos participantes de todos los núcleos de una región forman un cono de vectores al girar en precesión alrededor de un campo magnético externo estático, cada uno con fase aleatoria. Observando la proyección en un instante dado, y al no haber direcciones preferidas en el espacio, las proyecciones transversales se anulan entre sí. Esto implica que no hay magnetización neta en el campo transversal. Sin embargo las proyecciones longitudinales se suman resultando un vector magnetización apreciable. Este vector está completamente alineado en la dirección longitudinal. Se dice que en esta situación el sistema está relajado y en equilibrio dinámico.

No hay señal detectable en el estado de equilibrio, ya que no hay componente magnética transversal. Sin embargo, si se altera el estado de equilibrio se podrá producir una señal.



**Figura 3: Núcleos de Hidrógeno con sus vectores alineados aleatoriamente en todas las direcciones.**

Debido a las interacciones aleatorias con los núcleos vecinos, estos vectores se suman y anulan entre sí. No existe un campo electromagnético externo.



**Figura 4: Efecto apreciable en los núcleos de Hidrógeno ante la presencia de un fuerte campo magnético externo.**

Los momentos nucleares se alinean en forma paralela o anti-paralela (a) y adquieren un movimiento de precesión alrededor del campo externo (b).

### 2.2.2. Proceso de excitación

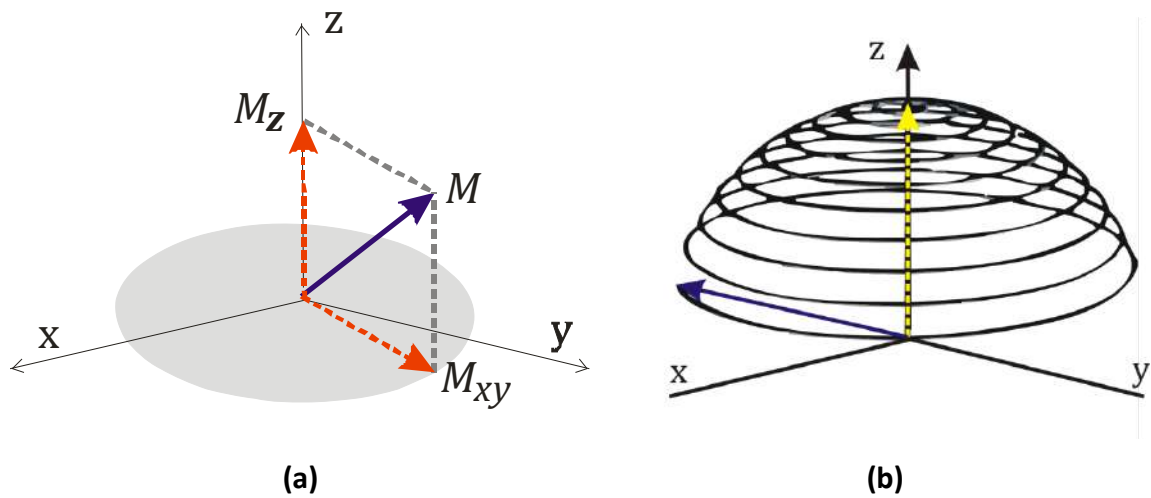
Para poder detectar una señal debe haber algún grado de magnetización sobre el eje transversal. Lo que se hace es excitar intencionalmente los núcleos para sacarlos de su alineamiento longitudinal hacia el plano transversal y de esa manera obtener una señal. Mientras el núcleo se relaja al equilibrio esta señal decrece con un patrón temporal característico que depende del contenido de hidrógeno del tejido (o sea de su química particular). Esta señal proporcionará la información de las propiedades en cada punto del objeto inmerso en el campo magnético.

Las ondas de radiofrecuencia (RF) son las más apropiadas para excitar los núcleos, dado que la constante giromagnética para los tejidos humanos es tal que la frecuencia obtenida por la ecuación de Larmor cae en el rango de las ondas de radiofrecuencia. Si se aplica una señal con la frecuencia de Larmor ésta será transferida al núcleo, ya que se trata de su frecuencia de resonancia, y éste se excitará. El efecto obtenido es que el vector magnetización se inclina sobre el plano transversal y los vectores que antes precesaban de manera aleatoria alrededor de  $B_0$  ahora lo hacen coherentemente girando en fase. Como consecuencia de esta coherencia de fase ya no se anulan entre sí en el plano transversal, quedando una magnetización neta transversal apreciable.

El vector magnetización se inclinará con un ángulo que dependerá del tiempo de aplicación de la RF. El vector de magnetización  $M$  comienza a precesar en un movimiento en espiral alejándose del alineamiento longitudinal original (Figura 5). A medida que el vector magnetización desciende en espiral, la magnetización longitudinal se reduce y la transversal aumenta. Con un ángulo de  $90^\circ$  respecto al eje  $x$  la magnetización longitudinal es cero y la transversal es máxima.

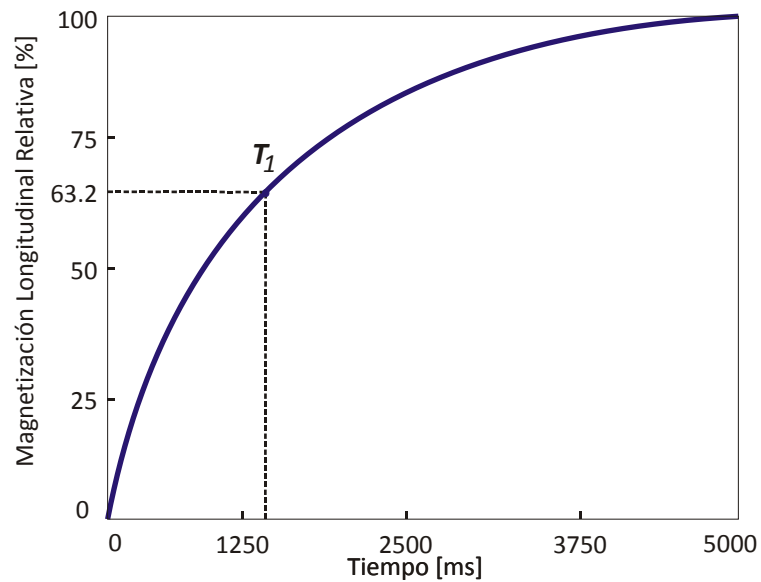
Resumiendo, la RF excita a los núcleos y tal excitación perturba el alineamiento longitudinal existente en la condición de equilibrio para crear una componente transversal. Esta componente es la que genera una señal detectable. La RF incide sobre los tejidos a una frecuencia que no distorsiona los enlaces electrostáticos de los átomos y moléculas constituyentes de los tejidos: no hay efectos indeseables.

Cuando cesa la aplicación de RF, los protones de hidrógeno estarán en un determinado estado de excitación manifestado por el desplazamiento del vector a cierto ángulo en relación al eje longitudinal y generando determinada señal debida a la componente de magnetización transversal inducida. Los núcleos revierten gradualmente su alineamiento longitudinal por la pérdida de la energía de excitación al medio. **Estas interacciones son llamadas spin-lattice.** Los núcleos también intercambian energía magnética entre ellos. Estas interacciones causan el desfase de los vectores de giro individual, es decir al interactuar entre sí pierden su coherencia de fase aunque el medio permanezca en excitación. Estas interacciones son llamadas spin-spin.



**Figura 5: Vector de magnetización y precesión en espiral.**

Diagrama de las componentes del vector magnetización  $M$  (a). Cuando se aplica una señal de radiofrecuencia, el vector de magnetización comienza a precesar en un movimiento en espiral alejándose del alineamiento longitudinal original (b).



**Figura 6: Tiempo de relajación  $T_1$ .**

Se define como el tiempo transcurrido hasta la recuperación del 63.2% del valor de la magnetización longitudinal original.

El tiempo asociado con la recuperación de la magnetización longitudinal, interacciones spin-lattice, se conoce como tiempo de relajación  $T_1$  mientras que el tiempo asociado con la pérdida de fase, interacciones spin-spin, se conoce como tiempo de relajación  $T_2$ . Ambos procesos conducen a la pérdida de la magnetización transversal y en consecuencia la señal detectada disminuye.

La señal se detecta mediante una bobina conductora que permite determinar la cantidad de magnetización transversal. Cuando el vector magnetización precesa, induce en la bobina una señal eléctrica y la intensidad de dicha señal eléctrica depende del ángulo que exista en cada momento entre el vector y la posición de la bobina. El tiempo de relajación  $T_1$  se define en términos de recuperación del alineamiento longitudinal, concretamente  **$T_1$  se define como el tiempo de recuperación del 63.2% del valor original.**

En los sólidos (materia gris, materia blanca) las moléculas están más cerca y más entrelazadas que en los líquidos (líquido ceforraquídeo). Como consecuencia la pérdida de energía por interacciones spin-lattice es más rápida, lo que implica que los tiempos  $T_1$  son menores como se ve en la Figura 6. Según los tejidos y sus enlaces intermoleculares resultarán diferentes tiempos de relajación  $T_1$ .

Las **interacciones spin-spin** se dan entre átomos vecinos que individualmente intercambian energía. En este proceso la energía es transferida entre los núcleos provocando la aceleración de algunos y el retardo de otros. Como consecuencia la coherencia de fase inicial que existe después de la excitación desaparece, conduciendo a cero la magnetización transversal rápidamente ya que las componentes transversales de los diferentes vectores vuelven a cancelarse entre si y por ende la señal detectada disminuye a cero con rapidez. Este proceso es independiente de la recuperación de la magnetización longitudinal y se produce antes de haber alcanzado el equilibrio. El tiempo de desfasaje conduce a cero la señal, ya que las componentes transversales de los diferentes vectores vuelven a cancelarse entre sí y la magnetización transversal neta tiende a cero por un proceso que resulta independiente de la recuperación de la magnetización longitudinal y que sucede antes de haber alcanzado el equilibrio.

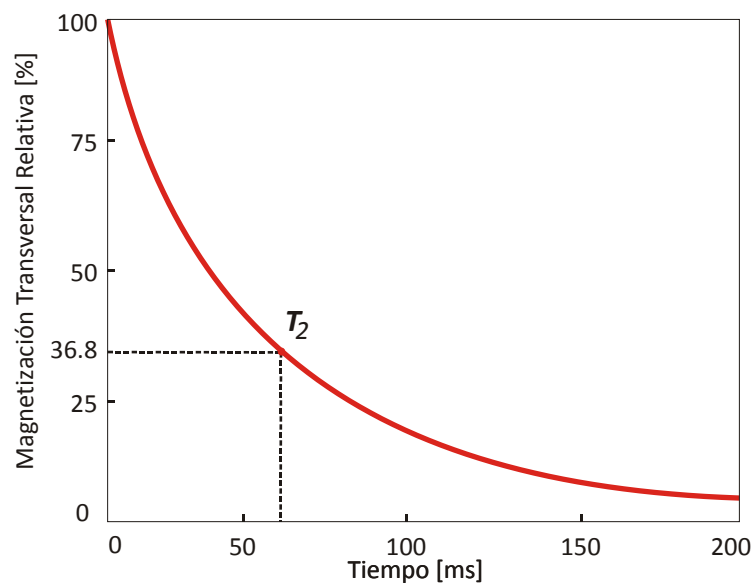
El tiempo de desfasaje de los vectores y de la caída de la señal está dado por  $T_2$  (Figura 7).  **$T_2$  se define como el tiempo necesario para que la señal se reduzca a un**

**36.8% de su valor original** (tanto el valor 36.8% como el 63.2% simplifican las ecuaciones matemáticas relacionadas con los fenómenos exponenciales y el número  $e$  y definen las llamadas “constantes de tiempo”). En general los tiempos  $T_2$  son del orden de la décima parte de  $T_1$  para los tejidos biológicos ( $T_2$  decenas de milisegundos y  $T_1$  cientos de milisegundos).

Un tercer mecanismo de caída de señal que contribuye a las interacciones spin-spin es debido al hecho de que un campo magnético no es totalmente uniforme debido a las heterogeneidades locales. Éstos aumentan las interacciones spin-spin, acelerando los procesos de desfasaje y la caída de la señal correspondiente. Este desfasaje es tan rápido que enmascara toda la información  $T_1$  y  $T_2$  de la señal. Este tiempo de relajación es denominado  $T_2^*$  y es indeseado, ya que los valores que interesan son los de  $T_1$  y  $T_2$  (Figura 8). Por otro lado  $T_2$  a su vez enmascara la relajación  $T_1$  que es relativamente larga. Por lo tanto, se plantea el problema de cómo separar  $T_1$  de  $T_2$  y ambos de  $T_2^*$ .

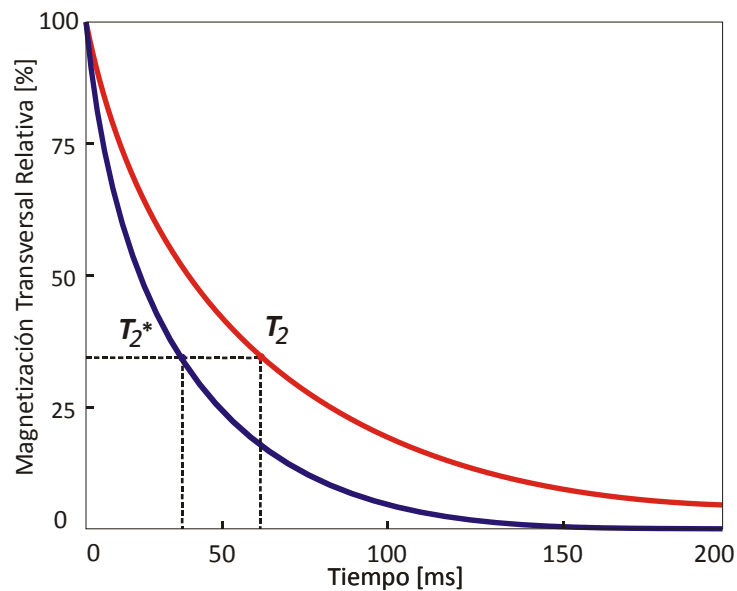
### 2.2.3. Secuencias de lectura de la señal

Para captar adecuadamente la señal es necesaria la aplicación de pulsos de excitación de RF durante el proceso de relajación. Inmediatamente después se mide la señal obtenida, generalmente en forma de eco. Para obtener estas señales puede ser necesaria la aplicación de uno o más pulsos en secuencia. Estas secuencias también son necesarias para eliminar los efectos de  $T_2^*$ . A continuación se citan algunas de las posibles.



**Figura 7: Tiempo de relajación  $T_2$ .**

Se define como el tiempo transcurrido hasta que la magnetización transversal se reduce al 36.8 % de su valor original.



**Figura 8: Tiempo de relajación  $T_2^*$ .**

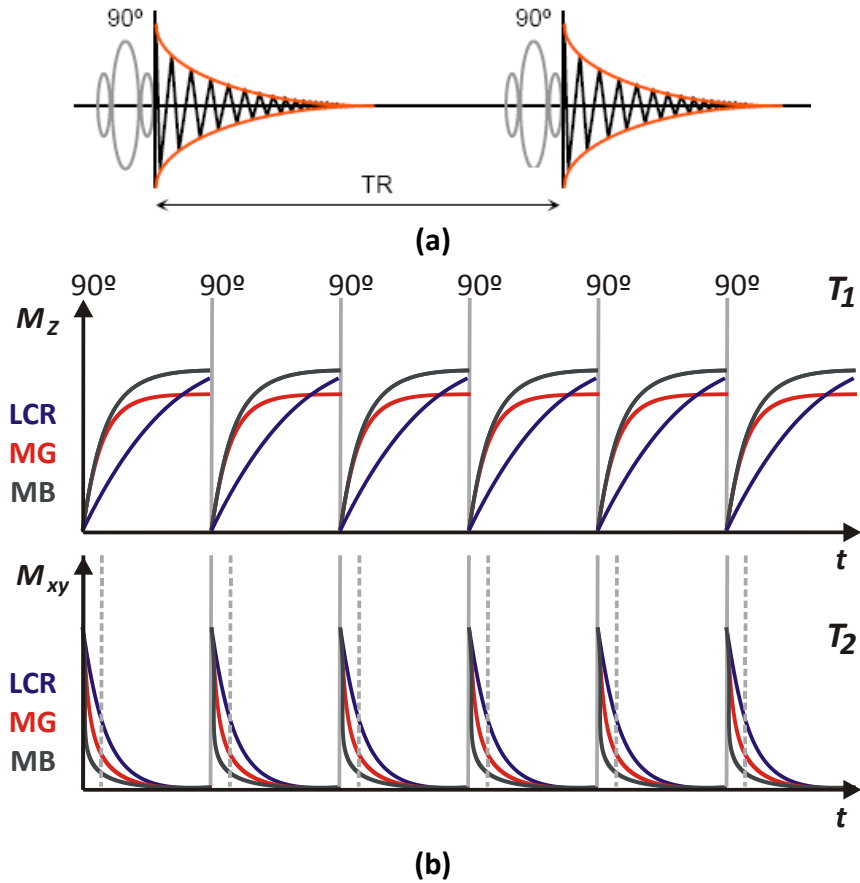
Las interacciones spin-spin en un campo magnético que no es totalmente uniforme hacen decaer la señal rápidamente. Este tiempo es tan pequeño que enmascara al tiempo  $T_2$ , que a su vez enmascara a  $T_1$ .

### 2.2.3.1. Saturación - Recuperación

Si se repite el pulso de RF se obtendrá una nueva señal. Un ejemplo de ellas es la *saturación - recuperación*. Aunque no es de uso habitual ilustra bien los procesos de RM (Figura 9).

La secuencia saturación-recuperación es una secuencia de pulsos de caída libre de la inducción, más conocidos como **Free Induction Decay** (FID), repetida; es decir, la aplicación de varios pulsos de  $90^\circ$  separados por un tiempo de repetición  $T_R$ .

Después de cada pulso de  $90^\circ$ , todos los vectores magnéticos quedan transversales y en fase. Comienzan inmediatamente a desfasarse y a volver a la posición longitudinal recorriendo una espiral. Cuando se repiten los pulsos de  $90^\circ$ , el segundo pulso capturará vectores en posiciones intermedias, es decir, saturados, antes de alcanzar la recuperación longitudinal completa y los proyectará  $90^\circ$  en el cuadrante siguiente. De esta manera, permanece una proyección neta productora de señal. La intensidad de esta señal dependerá del grado de recuperación longitudinal, o sea del tiempo de relajación  $T_1$  que haya tenido lugar en cada punto. Se dispone así de una forma de distinguir entre los diferentes tiempos de relajación  $T_1$  y de los diferentes tejidos que representan. Al aplicar un nuevo pulso los vectores vuelven a ponerse en fase. En todos los casos la señal es detectada independientemente de  $T_2^*$ .



**Figura 9: Secuencia saturación-recuperación.**

Se aplican pulsos de radiofrecuencia de 90° repetidos, separados por el tiempo de repetición TR (a), produciendo el efecto en la medición de los tiempos T<sub>1</sub> y T<sub>2</sub> que se observa en la figura (b), diferentes para distintas sustancias. LCR: Líquido Ceforraquídeo, MG: Materia Gris, MB: Materia Blanca

La señal detectada es igual a:

$$S_N = N \left( 1 - e^{-\frac{T_R}{T_1}} \right),$$

donde  $N$  es la densidad de protones presentes. Si el tiempo de repetición  $T_R$  es corto en relación a  $T_1$ , la magnetización longitudinal se recupera parcialmente y los siguientes pulsos de 90° proporcionan señales menos intensas. Por el contrario si la relación entre  $T_R$  y  $T_1$  es grande entonces las señales serán relativamente altas y mejora el contraste. En otras palabras, el contraste entre tejidos de una imagen es simplemente la diferencia entre las curvas de recuperación de cada tejido. Se debe buscar el  $T_R$  óptimo para visualizar los diferentes tejidos.

La secuencia saturación-recuperación es sensible a los errores de determinación del ángulo y con ella se obtiene una escala de contrastes pobres en la imagen obtenida comparada con otros tipos de secuencias.

### 2.2.3.2. Inversión - Recuperación

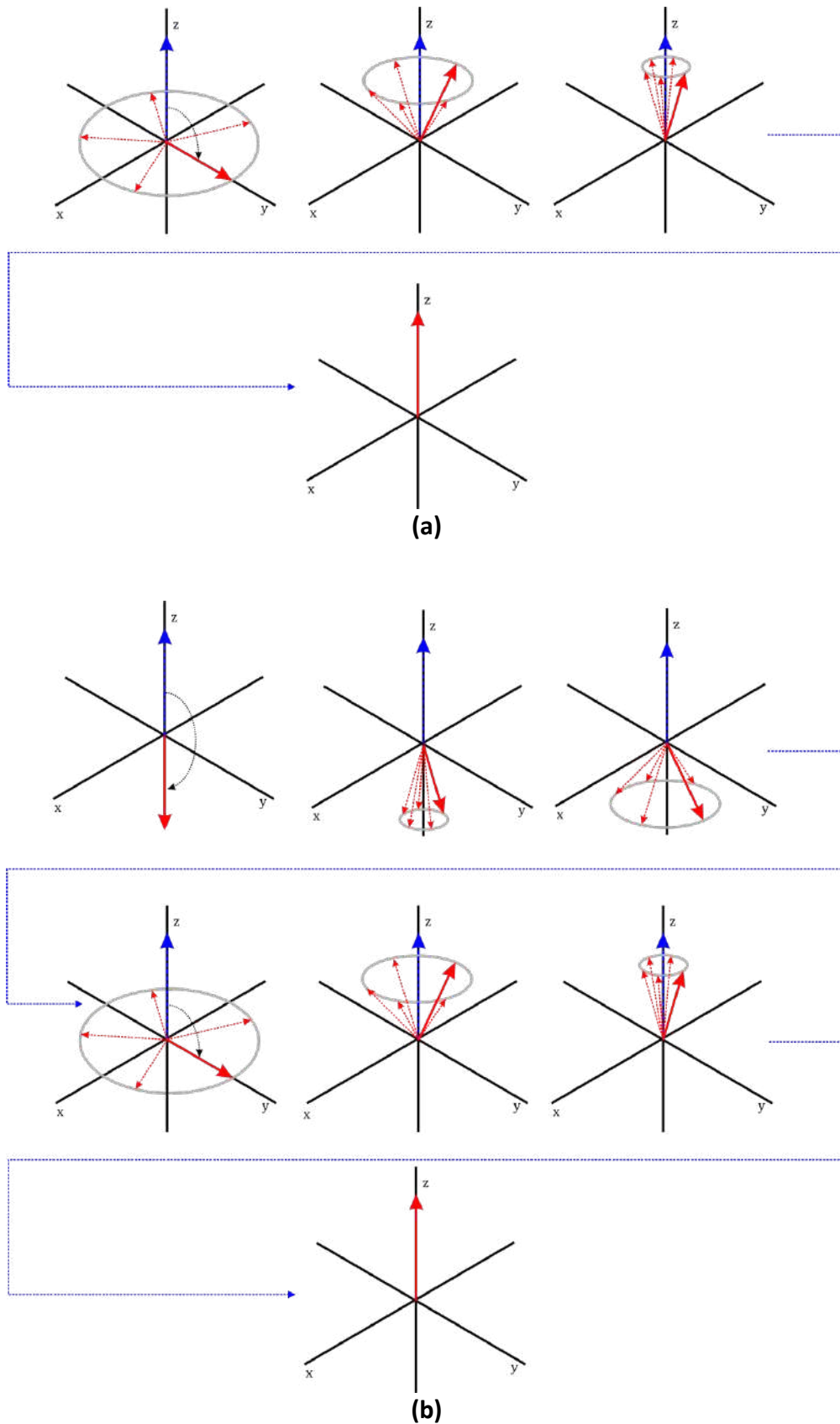
La inversión-recuperación es una secuencia de pulsos utilizada para obtener imágenes ponderadas en  $T_1$ . Se aplica inicialmente un pulso de  $180^\circ$  que invierte la magnetización hacia la dirección antiparalela. Inmediatamente después que cesa el pulso, aunque todos los vectores están en fase, no hay señal, ya que no hay componente de magnetización transversal.

A medida que pasa el tiempo y mientras los vectores se recuperan tampoco se produce señal dado el desfase que tiene lugar, como se ve en la Figura 10. Después de un intervalo de tiempo denominado **tiempo de inversión**  $T_{INV}$ , los vectores habrán recuperado la magnetización en distinto grado de acuerdo a los valores de  $T_1$  de cada uno de los tejidos. En ese momento se aplica un segundo pulso de  $90^\circ$ . Este pulso tiene el efecto de trasladar el vector magnetización al siguiente cuadrante y poner de nuevo en fase los spines proporcionando una señal medible.

La secuencia de inversión-recuperación es similar a la de saturación-recuperación pero los vectores tienen más espacio para recuperarse,  $180^\circ$  en vez de  $90^\circ$  (Figura 11). De esta manera se pueden observar mayores diferencias entre ellos.

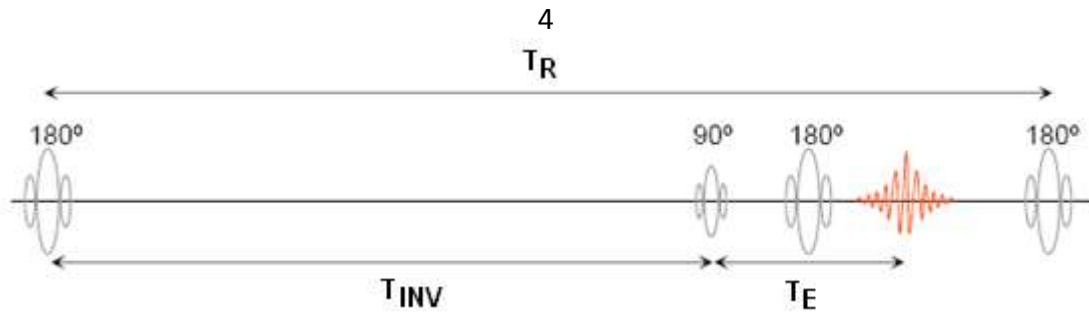
Las señales dependen de  $T_{INV}$  y de las características de los tejidos. Si  $T_{INV}$  es demasiado corto se dará poco tiempo para que se produzcan diferencias apreciables; si  $T_{INV}$  es demasiado largo se efectuará toda la recuperación longitudinal y habrá poca variación entre los tejidos. Si el vector magnetización se encontrara justo a  $0^\circ$  al aplicar el pulso de  $180^\circ$ , en el instante inmediatamente posterior irá a  $-90^\circ$  implicando una magnetización transversal nula y por lo tanto no se producirá señal.

Un aspecto interesante de la imagen es el hecho de que el contraste relativo puede ser invertido, es decir lo brillante verse oscuro y viceversa. Si el tiempo  $T_{INV}$  es corto puede darse esta inversión. Puede existir un punto donde dos tejidos proporcionen la misma señal; es decir, que no haya contraste y por lo tanto no se puedan distinguir (Figura 12).



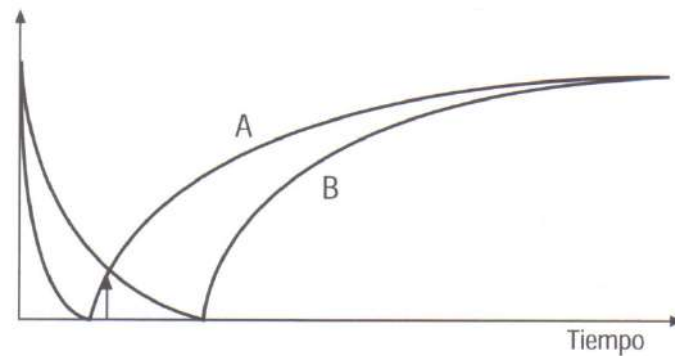
**Figura 10: Secuencia inversión-recuperación.**

Se observa el efecto de esta secuencia sobre las componentes longitudinal y transversal del vector magnetización con una excitación con un pulso de  $90^\circ$  (a) y de  $180^\circ$  (b).



**Figura 11: Secuencia inversión-recuperación. (Reproducida de [Blink, 2004])**

Se efectúan pulsos de 180° intercalados con pulsos de 90°.



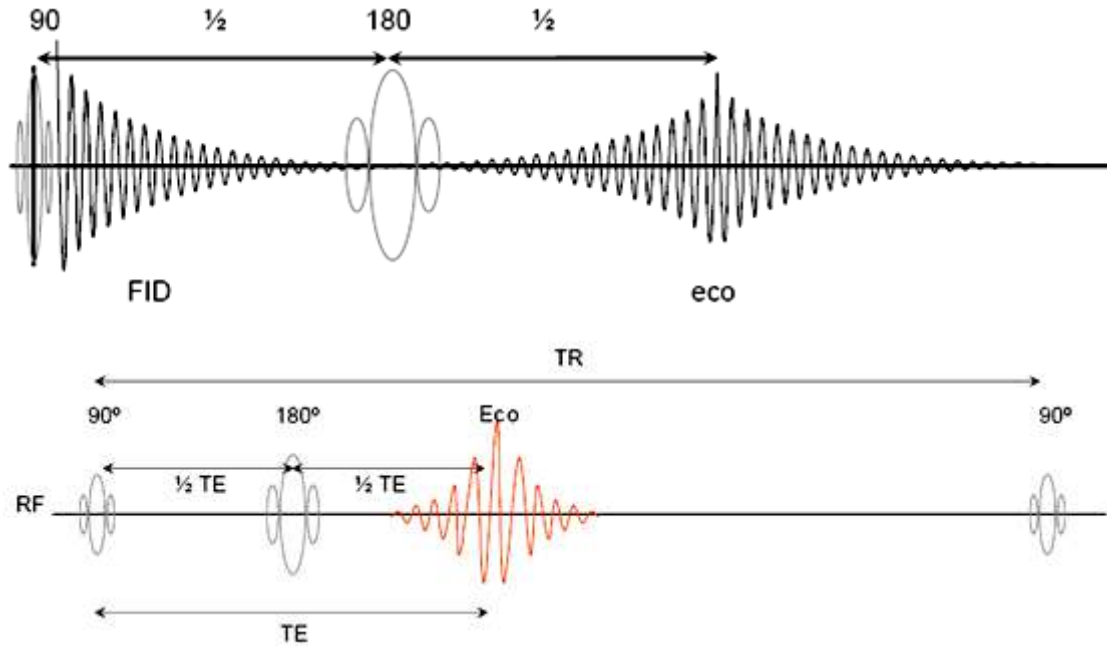
**Figura 12: Variación de la magnitud del vector de magnetización longitudinal (Reproducida de [Coussement, 2000]).**

No se observa la diferencia en la medición de los tiempos  $T_1$  para distintos tejidos A y B, pues la señal medida es la misma en el tiempo indicado con la flecha.

### 2.2.3.3. Secuencia de pulso spin-eco

Para obtener una imagen ponderada en  $T_2$  se debe eliminar la influencia del tiempo de relajación  $T_2^*$ . Primero se aplica un pulso de 90° que coloca el vector magnetización en el plano transversal. Inmediatamente después comienza a tener lugar un desfase. Este desfase depende tanto de las interacciones spin-spin como de las heterogeneidades del campo magnético externo. Con el tiempo, los vectores de giro más lento se separan de los de precesión de giro más rápido provocando que la magnetización transversal, como la señal, tienda a cero. Después se proporciona un nuevo pulso de 180° con un intervalo de tiempo aplicado  $\frac{T_E}{2}$ . Este pulso está diseñado para trasladar a los vectores al lado opuesto del plano transversal, tal que los vectores que precesan se muevan ahora hacia el punto de partida tendiendo a volver estar en fase. Es decir, los más rápidos alcanzan a los más lentos, volviendo ambos, rápidos y lentos, al mismo tiempo al punto de partida, volviendo a estar en fase en el tiempo  $T_E$ . Tan pronto como los vectores vuelven a estar en fase se produce una señal detectable

o señal de eco. Este eco aumenta alcanzando su valor máximo en  $T_E$ , cuando los vectores están totalmente en fase. En la Figura 13 se observa la secuencia de pulso spin-eco.



**Figura 13: Secuencia de pulso spin-eco (Reproducida de [Blink, 2004]).**

Primero se aplica un pulso de  $90^\circ$  que coloca el vector magnetización en el plano transversal. Después se proporciona un nuevo pulso de  $180^\circ$ .

Esta vuelta a estar en fase tiene sólo lugar para los desfases debido a las heterogeneidades locales, ya que están fijas en el espacio y el desfase que causan es reversible. Por otro lado el desfase debido al verdadero  $T_2$  es aleatorio e irreversible y no es afectado en la vuelta de fase. La señal entonces corresponde al desfase  $T_2$  que hubiera tenido si  $T_2^*$  no hubiera estado presente.

Si  $T_E$  es largo se deja más tiempo a los vectores para separar a sus diferentes  $T_2$  y la imagen final es ponderada en  $T_2$ . Si  $T_E$  es muy corto se produce poco desfase y la secuencia se aproxima a la saturación parcial. Esto implica que las imágenes proporcionadas son ponderadas en  $T_1$ . También este caso se puede invertir el contraste relativo. Esta relación entre la señal detectada y los tiempos  $T_E$ ,  $T_1$  y  $T_2$  queda plasmada en la ecuación siguiente:

$$S = Ne^{-\frac{T_E}{T_2}} \left( 1 - e^{-\frac{T_R}{T_1}} \right).$$

En la práctica clínica, ya que la señal de eco tiene un tratamiento matemático más favorable, tanto la secuencia de saturación-recuperación como la de inversión-recuperación son complementadas con un pulso final de  $180^\circ$ . Los  $T_E$  empleados son tan cortos que las imágenes resultantes son independientes de  $T_2$ .

Resumiendo, el contraste entre dos tipos de tejidos es la diferencia de intensidad entre las señales que emanan de los diferentes tejidos. La intensidad depende de las combinaciones específicas de los parámetros de tiempo  $T_{INV}$ ,  $T_R$  y  $T_E$  involucrados en la secuencia de pulso.

El tiempo  $T_R$  determina el grado en que la magnetización longitudinal se puede recuperar. Afecta la ponderación en  $T_1$ .

El tiempo  $T_E$  afecta la ponderación en  $T_2$ . Con un tiempo  $T_R$  largo las diferencias casi no dependen de  $T_1$ .

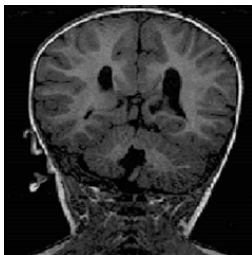
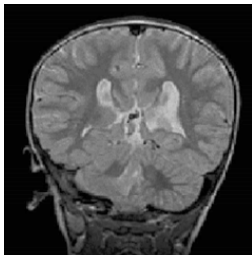
Estas combinaciones de tiempos se resumen en la Tabla I. Con una selección adecuada de  $T_R$  y  $T_E$  se tiene forma de cambiar la ponderación relativa en  $T_1$  y  $T_2$  de la imagen.

La intensidad o brillo en una posición dada de la IRM refleja la amplitud de la señal de RM en esa posición, la cual es proporcional a la concentración de protones móviles de hidrógeno en esa región.

Si los parámetros de la imagen son seleccionados de manera tal de reflejar sólo las variaciones de densidad de protones la imagen se conoce como imagen de densidad de protones PD. Un ejemplo de imagen PD se muestra en la Figura 14.

En las imágenes PD, la intensidad no varía mucho de estructura a estructura. La razón para esta apariencia es que la densidad de protones no varía mucho para los diferentes tejidos. Si  $T_R$  es muy largo ( $T_R \gg T_1$ ) y  $T_E$  es muy corto ( $T_E \ll T_2$ ) entonces ambos tejidos recuperan totalmente su magnetización longitudinal entre cada secuencia spin-echo de excitación-detección, aun si los tejidos poseen diferente valores de  $T_1$ . Por lo tanto las diferencias de  $T_1$  no afectarán la cantidad de magnetización longitudinal disponible para la próxima secuencia spin-echo del tren de pulsos.

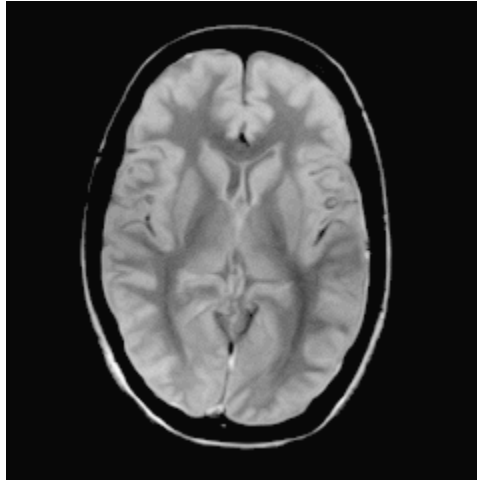
**Tabla I:** Combinaciones específicas de los parámetros de tiempo  $T_R$  y  $T_E$  y su relación con las diferentes ponderaciones  $T_1$  y  $T_2$ .

| TEJIDO      | $T_R$ corto                                                                                | $T_R$ largo                                                                                 |
|-------------|--------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------|
| $T_E$ corto | <br>$T_1$ | <br>PD    |
| $T_E$ largo |                                                                                            | <br>$T_2$ |

Teniendo en cuenta que el tiempo de relajación  $T_2$  describe la relación de disminución de la señal en plano transversal, un  $T_E$  largo pondrá en evidencia los diferentes valores de las señales procedentes de tejidos que poseen diferentes  $T_2$ . Sin embargo para  $T_E$  muy cortos la señal de ningún tejido tendrá tiempo de decaer, por lo que ni las diferencias en  $T_1$  ni en  $T_2$  afectarán el nivel de la intensidad de la imagen [Fletcher and Hornak, 1994].

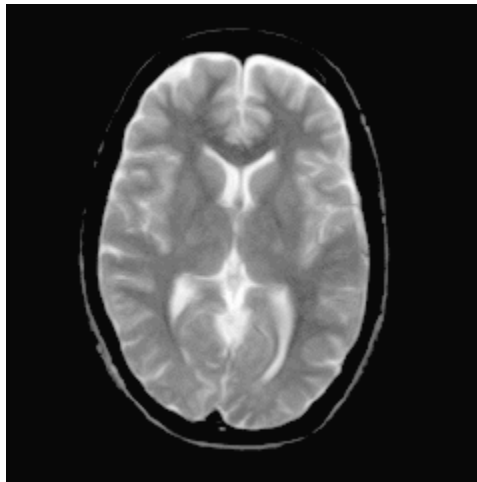
En las imágenes  $T_1$  y  $T_2$  la intensidad varía significativamente de tejido a tejido logrando muy buen contraste. Aun los tejidos en donde la densidad protónica es similar pueden poseer tiempos de relajación bien diferenciados.

Las IRM de cerebro pesadas en  $T_2$  presentan un buen contraste entre la **materia gris (MG)** y el **líquido cefalorraquídeo (LCR)** y poco contraste entre la MG y la **materia blanca (MB)**; esto se puede atribuir a que el tiempo  $T_2$  está directamente relacionado con las interacciones spin-spin. Como en el caso anterior, si  $T_R$  es largo permite una completa recuperación de la magnitud longitudinal en ambos tejidos entre la secuencia spin-echo, más allá de las variaciones de los valores de  $T_1$  de los tejidos.



**Figura 14: Imagen de densidad de protones (PD).**

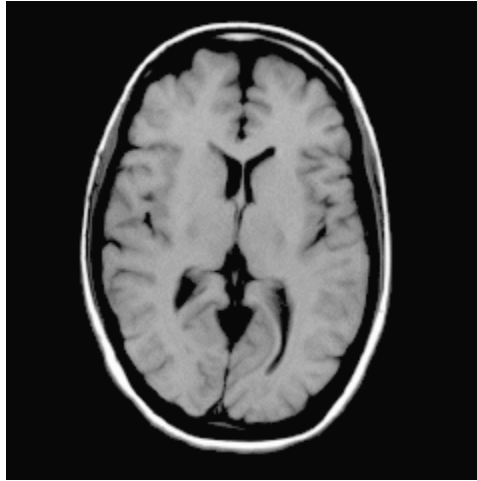
Esta imagen se obtuvo utilizando un tiempo de repetición  $T_R$  largo de 6000 ms y un tiempo de eco  $T_E$  corto de 20 ms.



**Figura 15: Imagen pesada en  $T_2$ .**

Esta imagen se obtuvo utilizando un tiempo de repetición  $T_R$  largo de 3300 ms y un tiempo de eco  $T_E$  largo de 112 ms.

A diferencia del caso anterior, las diferencias de  $T_2$  serán reflejadas en las amplitudes de las señales. Dado que la señal no es detectada hasta un tiempo  $T_E$  después del pulso de  $90^\circ$ , la magnitud transversal de un tejido con  $T_2$  corto decaerá más rápidamente que la de un tejido con un  $T_2$  largo. El tejido de  $T_2$  largo contribuye con más señal, causando en la IRM, una intensidad mayor de nivel de gris. Esta imagen se conoce como imagen pesada en  $T_2$  (Figura 15). En estas imágenes se ve con mayor intensidad el LCR, luego la MG, la MB y por último con el menor nivel de gris, la grasa subcutánea.



**Figura 16: Imagen pesada en  $T_1$ .**

Esta imagen se obtuvo utilizando un tiempo de repetición  $T_R$  intermedio, de 400 ms, y un tiempo de eco  $T_E$  corto, de 20 ms. Las estructuras de larga relajación son de baja intensidad en esta imagen.

Por último, si se acorta  $T_R$  al rango de valores de  $T_1$  de los tejidos ( $T_E \gg T_2$ ), como en el primer caso, un  $T_E$  muy corto no provee tiempo suficiente para que decaiga la magnetización transversal significativamente, de manera que las variaciones de  $T_2$  no afectan la intensidad de la señal. Sin embargo dado que  $T_R$  es corto, la cantidad de magnitud longitudinal presente inmediatamente antes de cada secuencia de excitación spin-echo dependerá de cuán rápido se recupere del pulso de excitación previo. Por lo tanto si el tejido tiene un  $T_1$  menor, producirá una señal mayor y por ende una intensidad en la imagen también mayor. Esta imagen se conoce como imagen pesada en  $T_1$  (Figura 16).

En resumen,  $T_R$  controla  $T_1$  y  $T_E$  controla  $T_2$ . Los tejidos con tiempos de relajación  $T_2$  cortos son oscuros en las imágenes pesadas en  $T_2$  mientras que los tejidos con tiempos de relajación  $T_1$  cortos son brillantes en las imágenes pesadas en  $T_1$ . La Tabla II da una idea de los valores de los tiempos de relajación  $T_1$  y  $T_2$  para los diferentes tejidos.

**Tabla II: Representación de los tiempos de relajación para diferentes tejidos (Valores en milisegundos medidos a 1 Tesla).**

| TEJIDO                  | $T_1$ [ms] | $T_2$ [ms] |
|-------------------------|------------|------------|
| Materia blanca          | 390        | 90         |
| Materia gris            | 520        | 100        |
| Líquido cefalorraquídeo | 2000       | 300        |
| Músculo esquelético     | 600        | 40         |
| Grasa                   | 180        | 90         |
| Sangre                  | 800        | 180        |

Para finalizar, se da en la Tabla III una descripción de los principales tejidos presentes en una RM de cerebro y las características en la intensidad de la imagen que presentan en las imágenes pesadas en  $T_1$ ,  $T_2$  y  $PD$  [Edelman et al., 1996]. Esta tabla es de gran utilidad para la automatización de la segmentación de este tipo de imágenes, presentada en esta tesis.

**Tabla III: Intensidades en las distintas imágenes para diferentes tejidos.**

| TEJIDO                                     | $T_1$          | $T_2$                  | $PD$          |
|--------------------------------------------|----------------|------------------------|---------------|
| Materia blanca                             | Brillante      | Oscura o Gris          | Brillante     |
| Materia gris                               | Oscura o Gris  | Brillante              | Muy Brillante |
| Líquido cefalorraquídeo en los ventrículos | Muy Oscura     | Muy Brillante o blanco | Gris          |
| Líquido cefalorraquídeo en flujo rápido    | Oscuro o Negro | Oscuro o Negro         |               |
| Agua                                       | Oscura         | Clara o blanca         |               |
| Grasa                                      | Brillante      | Clara                  | Brillante     |
| Músculo                                    | Gris           | Gris                   |               |
| Aire                                       | Negro          | Negro                  |               |
| Hueso cortical                             | Negro          | Negro                  |               |
| Médula ósea con grasa                      | Brillante      | Clara                  |               |
| Hígado                                     | Gris           | Gris                   |               |
| Riñón                                      | Gris           | Gris                   |               |

## 2.3. Distorsiones en las imágenes

---

La intensidad de señal obtenida puede variar espacialmente. Esta variación espacial se denomina **no-uniformidad de la intensidad** (INU, *Intensity Non-uniformity*) y se debe principalmente a la no-uniformidad de la bobina que actúa como receptor o emisor, pero también la intensidad de la señal es alterada por distorsiones geométricas [Wang and Doddrell, 2005].

La INU constituye un fenómeno adverso que se manifiesta como variaciones lentas de intensidad del mismo tejido sobre todo el dominio de la imagen. Este fenómeno puede tener serias implicancias en el procesamiento de IRM. Por ejemplo, la INU incrementa el solapamiento entre distribuciones de intensidad de los diferentes tejidos, lo que hace que la segmentación sea mucho más dificultosa y menos precisa [Likar et al., 2002].

La INU hace que diferentes regiones de la imagen se vean como si estuvieran “iluminadas” por una fuente heterogénea. Se han propuesto numerosos métodos para intentar minimizar este efecto [Vovk et al., 2007].

Otro fenómeno de relevancia lo constituye el **ruido**. Éste se presenta en la imagen como un patrón granular irregular, lo que degrada la información útil presente en la imagen. Sus causas son variadas: el movimiento térmico, que provoca una emisión de RF; las bobinas; los amplificadores, etc. [Macovski, 1996]

La imagen adquirida se encuentra contaminada por ruido blanco aditivo, generalmente descrito por una función densidad de probabilidad gaussiana. La varianza del ruido es un importante factor a tener en cuenta cuando se procesan las IRM, dado que lo obtenido con los métodos aplicados puede estar afectado por este valor. También se han propuesto métodos para reducir su influencia [Sijbers et al., 2007].

## 2.4. El rol de la RM en la detección de enfermedades cerebrales

---

La neuroimagen provee un importante apoyo para el diagnóstico clínico de algunas enfermedades, como por ejemplo la enfermedad de Alzheimer, la demencia vascular y la demencia asociada con atrofia cortical.

La detección de patrones de normalidad mediante el análisis de IRM de cerebro es constante motivo de investigación para el diagnóstico y la caracterización de patologías.

La IRM provee una representación anatómica del cerebro que permite la diferenciación de los distintos tejidos y la identificación de estructuras neuroanatómicas. Puede utilizarse para descartar lesiones e identificar patrones de atrofia en distintas enfermedades neurodegenerativas a diferencia de la Resonancia Magnética Funcional (RMF), que puede diferenciar patrones de hipometabolismo, indicando la presencia de un proceso degenerativo [Manes, 2000].

Como ejemplo, en la enfermedad de Alzheimer se observa un patrón de hipometabolismo bitemporoparietal, detectable mediante RMF, pero este estudio es realizado en muy pocas instituciones en Argentina. En cambio la RM anatómica está relativamente mucho más disponible y debido a su alta resolución y mejor contraste de los tejidos blandos (respecto a la TAC) posee una gran sensibilidad para los cambios que se producen en esta enfermedad.

En 1994 el Subcomité de Estándares de Calidad de la Academia Americana de Neurología (*Quality Standards Subcommittee of the American Academy of Neurology*) publicó guías de "práctica y parámetros" para la evaluación diagnóstica de personas sospechadas de padecer demencias [American Academy of Neurology, 1994]. Estas guías establecían el examen neurológico estándar y ciertos análisis de laboratorio como los componentes esenciales del procedimiento diagnóstico. Sin embargo, los estudios con neuroimágenes fueron sugeridos como opcionales. Pero en el año 2001 en una actualización de las guías se incluyeron los estudios con IRM como apropiados para la evaluación inicial de pacientes con demencia [Knopman et al., 2001].

Uno de los objetivos primordiales de estos procedimientos diagnósticos es la detección de los primeros cambios en el cerebro que producen las diferentes enfermedades. Esto permitiría que alguna terapia pudiera detener el avance de la atrofia del cerebro, siendo administrada tempranamente.

El *Multicentre Consortium to Establish a Registry for Alzheimer's Disease* (CERAD) demostró que la evaluación visual de la atrofia generalizada era insatisfactoria y que la interpretación de las imágenes era altamente subjetiva. El CERAD determinó en 1992 que se necesitaban técnicas de mayor precisión para medir la atrofia global [Davis et al., 1992]. Desde entonces se ha presentado una gran cantidad de métodos automáticos y semiautomáticos para solucionar este problema.

Sin embargo, el método que se practica en la mayoría de las instituciones de Diagnóstico por Imágenes en Argentina en la clínica de rutina sigue siendo la evaluación visual de los especialistas en imágenes y su opinión sobre si los surcos o el volumen ventricular están fuera del rango de normalidad. Este hecho puede deberse a que el software asociado a los equipos no suele incluir herramientas de segmentación de tejidos y otras disponibles son de gran complejidad para los operadores habituales.

Con el advenimiento de equipos informáticos potentes y mejores métodos de segmentación automática y semiautomática, se ha efectuado gran cantidad de estudios para establecer patrones de normalidad de los volúmenes absolutos y relativos de MG, MB y LCR, lo que sigue siendo un tema de constante actualización.

Diversos estudios han establecido que las IRM ayudan a establecer un patrón normal de los cambios en el espacio intracraneal, del cerebro completo, la MG, la MB y el LCR en las distintas etapas de la vida. La medición de estos volúmenes es necesaria en la investigación de la detección de anormalidades [Courchesne et al., 2000]. Se ha determinado que la pérdida de volumen de la MG es una función lineal de la edad en los adultos, mientras que el déficit de volumen de la MB parece retrasarse hasta la mitad de la edad adulta, independientemente del sexo. El análisis cuantitativo de los porcentajes de MG y MB ayuda a la determinación de patologías cuando se tiene un patrón de cambio normal debido a la edad [Ge et al., 2002].

Un caso de sumo interés que no ha cesado de ser investigado desde los años 90 es el del autismo: se ha detectado una tasa anormal del crecimiento del cerebro, que se traduce en hiperplasia en la MG cerebral y la MB del cerebelo en los primeros años de vida, lo que fue cuantificado con estudios de RM [Courchesne et al., 2001, Courchesne et al., 2007].

Se ha demostrado correlación entre atrofia cerebral y fracciones de MG y MB con la valoración dada en la escala de demencia "*Mini-Mental Status Examination and Blessed Dementia Scale*". Se ha hallado disminución en las fracciones de MG y MB y aumento en la de LCR, lo que resulta en alertas para el diagnóstico prematuro de la enfermedad de Alzheimer [Brunetti et al., 2000].

Se ha asociado a la esquizofrenia con déficits sutiles en los diferentes tejidos cerebrales. Se ha reportado también que mayores volúmenes en la MB en pacientes con esquizofrenia y descendientes de los mismos sugiere una alteración en la maduración genética de la mielinización que se espera durante la adolescencia y la adultez. También en estos pacientes se observa déficit de la MG. El incremento del LCR subcalloso y subaracnoideo ocurre en la edad adulta, pero aparece prematuramente en individuos masculinos con esquizofrenia. La reducción de la MG es un efecto natural de la edad pero se encuentra debajo de un umbral en estos pacientes [Narr et al., 2003, Hulshoff Pol et al., 2002, Hulshoff Pol and Kahn, 2008]. Por todo esto, un estudio adecuado de los volúmenes de los diferentes tejidos tiene gran valor predictivo en el diagnóstico temprano de la esquizofrenia [Ho, 2007].

Los volúmenes de MG, MB y LCR relativos al volumen total del cerebro se vieron significativamente reducidos en pacientes que manifestaron un primer episodio de esquizofrenia y aún no habían sido medicados. Se encontró menos MG en el núcleo caudado, el giro cingulado, giro hipocampal, giro temporal, cerebelo y tálamo. También se ha observado significativa reducción de la MB en el sistema límbico anterior derecho y significativamente aumentado el LCR especialmente en el ventrículo lateral derecho [Chua et al., 2007].

Mediante la segmentación de IRM en las diferentes sustancias, seguido por una edición manual para diferenciar la MG cortical y no cortical, pudo efectuarse un estudio sobre anomalías de los volúmenes de LCR y en la cantidad de MG cortical

en individuos con desorden de personalidad esquizotípica, patología que tiene la misma predisposición genética que la esquizofrenia [Dickey et al., 2002].

La caracterización del crecimiento de la MG y de la MB en recién nacidos mediante IRM ha permitido interesantes conclusiones tales como que desde el nacimiento ya se encuentran algunos patrones adultos de dimorfismo sexual (diferencias sistemáticas en individuos de diferente sexo) y asimetrías cerebrales pero otros se desarrollan posteriormente. Se ha medido la MG y la WN diferenciando la velocidad de crecimiento de las mismas en las diferentes regiones cerebrales [Gilmore et al., 2007].

El estudio de las RM permite cuantificar los daños cerebrales que se han producido en individuos alcohólicos. Se ha demostrado gran déficit en las densidades de MG y de MB en determinadas regiones del cerebro [Mechtcheriakov et al., 2007].

Se conoce que también se produce atrofia en el curso de la esclerosis múltiple pero es motivo de constante investigación el determinar cuándo, cuáles tejidos son afectados y con qué ritmo ocurre esta atrofia [Dalton et al., 2004]. El análisis de IRM puede conducir a futuras guías para el diagnóstico prematuro de esclerosis múltiple y para el monitoreo del tratamiento. Algunas técnicas para medir la atrofia cerebral requieren la estimación precisa de los volúmenes de MG, MB y LCR.

Estudios preliminares sobre epilepsia han utilizado IRM para analizar características morfométricas especiales del cerebro en niños. Se ha demostrado una fuerte asociación entre el desarrollo cognitivo y el volumen cerebral, especialmente de la MB, y que esta asociación está ausente en niños epilépticos. Los niños con problemas de aprendizaje mostraron reducción en la MG occipital y parietal [Hermann et al., 2006].

El procesamiento de imágenes por computadora es una herramienta poderosa. La alta calidad de las IRM hace que sean apropiadas para gran cantidad de estudios diagnósticos. Los esfuerzos apuntan entonces a cuantificar la cantidad (en proporción o en valores de volumen) de los diferentes tejidos cerebrales de una manera más automática, más sensible y sobre todo menos subjetiva. La objetividad y repetitividad del análisis permiten determinar patrones de evolución o cambios en los pacientes y

dar así un paso hacia la elaboración de guías para la asistencia al diagnóstico temprano de las enfermedades, su tratamiento y su evolución.

---

## 3. Lógica Difusa

## 3. Lógica Difusa

---

### 3.1. Introducción

---

La Lógica Difusa (LD) es una rama de la Inteligencia Computacional que se funda en la incertidumbre, lo cual permite manejar información vaga o de difícil especificación si se quiere utilizar objetivamente esta información con un fin específico [Kecman, 2001]. Constituye una herramienta adecuada para modelar el conocimiento o la comprensión de conceptos.

La LD permite implementar procesos de razonamiento por medio de reglas y predicados que generalmente se refieren a cantidades indefinidas o inciertas. Estos predicados pueden obtenerse con sistemas que “aprenden” al “procesar” datos reales o pueden también ser formulados por un experto humano o, mejor aún, por el consenso entre varios de ellos.

La flexibilidad de la LD la hace apropiada para los sistemas de asistencia en la toma de decisiones [Sousa and Kaymak, 2003, Li and Yen, 1995b]. Su capacidad para elaborar modelos lingüísticos es útil para incorporar el conocimiento de especialistas en forma de predicados y transformar ese conocimiento en algoritmos de cálculo. En este sentido, se han realizado estudios basados en Lógica Multivaluada en los que se han presentado diversos operadores para las operaciones entre valores de verdad de sentencias (o entre conjuntos difusos, que se definirán más adelante).

Los modelos basados en LD pueden verse como “cajas blancas”, dado que permiten ver su estructura explícitamente. En contraposición a los modelos basados en datos exclusivamente, como las RN, que corresponderían a “cajas negras”. Los modelos difusos que combinan experiencia y conocimiento con datos numéricos, como el presentado en esta tesis, son vistos como “cajas grises” [Kecman, 2001].

## 3.2. Historia

---

La LD puede ser vista como un lenguaje que permite trasladar sentencias en lenguaje natural a un lenguaje matemático formal. Mientras la motivación original fue ayudar a manejar aspectos imprecisos del mundo real, la práctica temprana de la LD permitió el desarrollo de aplicaciones prácticas. Aparecieron numerosas publicaciones que presentaban los fundamentos básicos con aplicaciones potenciales. Surgió una fuerte necesidad de distinguir la LD de la Teoría de Probabilidad. Tal como se aborda actualmente, la teoría de conjuntos difusos y la Teoría de Probabilidad tienen diferentes tipos de incertidumbre.

En 1994, la teoría de la LD se encontraba en la cumbre. Hubo una gran explosión de aplicaciones exitosas, aun cuando su origen era motivo de ataque. Varios autores afirman que esta disciplina estuvo bajo el nombre de “Lógica Difusa” durante 25 años, pero sus orígenes se remontan varios siglos hacia atrás [Elkan et al., 1994].

En el siglo XVIII el filósofo y obispo anglicano Irlandés, George Berkeley y David Hume describieron que el núcleo de un concepto atrae conceptos similares [Berkeley, 1710]. Hume en particular creía en la lógica del sentido común, el razonamiento basado en el conocimiento que la gente adquiere en forma ordinaria mediante vivencias en el mundo [Ayer, 1940]. En Alemania, Immanuel Kant, consideraba que sólo los matemáticos podían proveer definiciones claras, y muchos principios contradictorios no tenían solución. Por ejemplo la materia podía ser dividida infinitamente y al mismo tiempo no podía ser dividida infinitamente [Kuehn, 2001].

Particularmente la escuela americana de la filosofía llamada pragmatismo fundada a principios de siglo por Charles Sanders Peirce, cuyas ideas se fundamentaron en estos conceptos, fue el primero en considerar “vaguedades”, más que falso o verdadero, como forma de acercamiento al mundo y a la forma en que la gente funciona [Peirce, 1883].

La idea de que la lógica produce contradicciones fue popularizada por el filósofo y matemático británico Bertrand Russell, a principios del siglo XX. Estudió las vaguedades del lenguaje, concluyendo con precisión que la vaguedad es un grado [Garciadiego, 1992]. El filósofo austríaco Ludwig Wittgenstein estudió las formas en

las que una palabra puede ser empleada para muchas cosas que tienen algo en común [Wittgenstein, 1958].

La primera lógica de vaguedades fue desarrollada en 1920 por el filósofo Jan Lukasiewicz, quien visualizó los conjuntos con un posible grado de pertenencia con valores de 0 y 1 y después los extendió a un número infinito de valores entre 0 y 1. Fue un pionero en la investigación de las Lógicas Multivaluadas. Los historiadores de la lógica sostienen que la raíz de la Lógica Multivaluada es el Indeterminismo, cuya relación fue establecida en 1910. Lukasiewicz Introdujo el cálculo axiomático con Lógica Tri-valuada en 1917 [Lukasiewicz, 1970, Moler, 1998].

En los años sesenta, Lofti Zadeh propuso la LD como disciplina, que combina los conceptos de la lógica y de los conjuntos de Lukasiewicz mediante la definición de grados de pertenencia [Zadeh, 1965].

### 3.3. Conceptos de Lógica Difusa

---

#### 3.3.1. Imprecisión

La mayoría de los fenómenos son imprecisos; es decir, tienen implícito un cierto grado de vaguedad en la descripción de su naturaleza, que puede estar asociada con su forma, posición, momento, color, textura, o incluso en la semántica que describe lo que son.

En muchos casos el mismo concepto puede tener diferentes grados de imprecisión en diferentes contextos o tiempos. Un día “cálido” en invierno no es exactamente lo mismo que un día “cálido” en primavera. La definición de cuándo la temperatura va de “templada” a “caliente” no es exacta: no es natural identificar un punto simple en el que variando un grado se cambie de “templado” a “caliente”. Esta situación, asociada continuamente a los fenómenos, es común en todos los campos de estudio: sociología, física, biología, finanzas, ingeniería, oceanografía, psicología, etc.

Se acepta a la imprecisión como una consecuencia natural de “la forma de las cosas en el mundo”. La dicotomía entre el rigor y la precisión del modelado matemático en todos los campos y la intrínseca incertidumbre de “el mundo real” no

es generalmente aceptada por los científicos, filósofos y analistas de negocios. Generalmente se aproximan estos eventos a través de funciones o procesamiento numéricos y se escoge un resultado, a veces sin hacer un análisis del conocimiento empírico. Sin embargo la mente humana procesa y entiende fácilmente, de manera implícita, la imprecisión de la información. Está capacitada para formular planes, tomar decisiones y reconocer conceptos compatibles con altos niveles de vaguedad y ambigüedad.

Algunas sentencias:

“La cámara de combustión está demasiado caliente.”

“La inflación actual aumenta rápidamente.”

“Los precios están por debajo de los precios de la competencia.”

“Esta región de la imagen es hiperintensa.”

Estas proposiciones responden al concepto de “cómo es”. Sin embargo, son incompatibles con el modelado tradicional y el diseño de sistemas de información. Incorporando estos conceptos se logra que los sistemas sean potentes y se aproximen directamente a la realidad.

Algunas proposiciones pueden expresar su veracidad como un nivel de pertenencia a un conjunto denominado “conjunto difuso”. En estas proposiciones la LD permite una doble interpretación: puede pensarse en el grado de verdad de que la expresión “ $x$  pertenece a  $A$ ” sea cierta, o bien puede pensarse en grados de pertenencia del elemento  $x$  al conjunto  $A$ .

Aunque la pertenencia a un conjunto o el grado de verdad parecieran asociarse a probabilidades, no se está trabajando en este contexto, sino en el de *posibilidades*. Hay esenciales diferencias entre ambos espacios, que quedan claras al presentar las operaciones entre grados de pertenencia o grados de verdad [Zadeh, 1965].

A continuación se presenta una formalización del concepto de Conjuntos Difusos.

### 3.3.2. Conjuntos Difusos

La teoría de conjuntos difusos constituye una generalización de la teoría de conjuntos clásica, introduciendo la noción de que la pertenencia de un elemento a un conjunto ya no es solamente verdadera o falsa. Se expondrá brevemente un enfoque práctico que permita utilizar los conjuntos difusos como herramienta para la cuantificación de valores de verdad de predicados.

Sea  $X$  un conjunto de objetos, llamado **universo**, cuyos elementos son denotados por  $x$ , La pertenencia de esos objetos a un conjunto  $A$  se puede representar por medio de una función  $\mu_A(x)$ , de la siguiente manera:

$$\mu_A(x) = \begin{cases} 1 & \text{si y sólo si } x \in A \\ 0 & \text{si y sólo si } x \notin A \end{cases} \quad (1)$$

Donde la imagen de  $\mu_A(x)$ , el conjunto  $\{0, 1\}$ , es el conjunto de evaluación.

A partir de ésta se llega a la definición de **Conjunto Difuso**, cuando el conjunto de evaluación de  $\mu_A(x)$  es el intervalo  $[0, 1]$  [Zadeh, 1965], mediante propiedades de los espacios de posibilidad, que debe ser diferenciado de los espacios de probabilidad [Dubois and Prade, 2000].

DEFINICIÓN 1: Sea  $g$  una función de  $\mathcal{P}(X)$  en  $[0, 1]$ , donde  $\mathcal{P}(X)$  es el conjunto de partes de  $X$ . Se dice que  $g$  es una **medida difusa** si y sólo si:

1.  $g(\emptyset) = 0$ ;  $g(X) = 1$ ;
2.  $\forall A, B \in \mathcal{P}(X)$ , si  $A \subset B$  entonces  $g(A) \leq g(B)$ ;
3. Si  $\forall i \in \mathbb{N}$ ,  $A_i \in \mathcal{P}(X)$  y  $(A_i)_{i \in \mathbb{N}}$  es una sucesión de conjuntos monótona creciente o decreciente (esto es  $A_1 \subseteq A_2 \subseteq \dots \subseteq A_n \subseteq \dots$  o bien  $A_1 \supseteq A_2 \supseteq \dots \supseteq A_n \supseteq \dots$ ), entonces

$$\lim_{n \rightarrow \infty} g(A_i) = g\left(\lim_{n \rightarrow \infty} A_i\right),$$

Donde  $\lim_{n \rightarrow \infty} A_i = \bigcup_{i \in \mathbb{N}} A_i$  si  $(A_i)_{i \in \mathbb{N}}$  es creciente y  $\lim_{n \rightarrow \infty} A_i = \bigcap_{i \in \mathbb{N}} A_i$  si  $(A_i)_{i \in \mathbb{N}}$  es decreciente.

La propiedad 3 se conoce como continuidad.

En el contexto anterior debe diferenciarse la medida de probabilidad de la medida de posibilidad.

DEFINICIÓN 2: Un **espacio de probabilidad** es una terna  $(X, \mathcal{F}, P)$ , donde  $X$  es el espacio muestral,  $\mathcal{F}$  es una  $\sigma$ -álgebra de los subconjuntos de  $X$  y  $P$  es una medida de probabilidad en  $X$  que cumple:

1.  $\forall A, P(A) \in [0,1]; P(X) = 1$ ;
2. Si  $\forall i \in \mathbb{N}, A_i \in \mathcal{P}(X)$  y  $\forall i \neq j A_i \cap A_j = \emptyset$ , entonces

$$P\left(\bigcup_{i \in \mathbb{N}} A_i\right) = \sum_{i=1}^{\infty} P(A_i).$$

Propiedad: Una medida de probabilidad  $P$  es una medida difusa.

Una medida de posibilidad puede verse como una función de pertenencia  $\mu_A(x)$  de los elementos de  $X$  al *conjunto*  $A$ , **lo que ubica a la LD en el contexto de la teoría de posibilidades.**

DEFINICIÓN 3: Un **espacio de posibilidad** es una terna  $(X, \mathcal{P}(X), \Pi)$ , donde  $X$  es el espacio muestral,  $\mathcal{P}(X)$  es el conjunto de partes de  $X$  y  $\Pi$  es una medida de posibilidad es una función de  $\mathcal{P}(X)$  en  $[0, 1]$  tal que:

1.  $\Pi(\emptyset) = 0; \Pi(X) = 1$ ;
2. Para cualquier colección  $\{M_i\}$  de subconjuntos de  $X$ ,

$$\Pi\left(\bigcup_i M_i\right) = \sup_i \Pi(M_i).$$

Propiedad: Una medida de posibilidad  $\Pi$  es una medida difusa.

Cuanto más cercano a 1 es el valor de  $\mu_A(x)$ , más pertenece el elemento al *conjunto*. Por lo tanto puede verse a  $A$  como un *subconjunto* de  $X$  que no tiene límites bien definidos.

Un conjunto difuso queda completamente caracterizado por los pares formados entre elementos y sus grados de pertenencia.

DEFINICIÓN 4:  $A$  es un conjunto difuso si:

$$A = \{(x, \mu_A(x)) \text{ con } x \in X\}, \quad (2)$$

siendo  $\mu_A(x)$  una medida de posibilidad  $\mu: X \rightarrow [0, 1]$ .

**Ejemplo:** Sea  $X = \{1, 2, 3, 4, 5, 6, 7, 8, 9\}$  y los conjuntos difusos definidos a continuación:

|        |   |                                                                  |
|--------|---|------------------------------------------------------------------|
| POCOS  | = | $\{(1, 0.4), (2, 0.8), (2, 1.0), (4, 0.4)\}$                     |
| VARIOS | = | $\{(3, 0.5), (4, 0.8), (5, 1.0), (6, 1.0), (7, 0.8), (8, 0.5)\}$ |
| MUCHOS | = | $\{(6, 0.4), (7, 0.6), (8, 0.8), (9, 0.9), (10, 1.0)\}$ .        |

En este ejemplo, el elemento “4” pertenece en grado 0.4 al conjunto POCOS, en grado 0.8 al conjunto VARIOS y en grado 0 a MUCHOS.

En esta tesis se utilizarán los conjuntos difusos para la cuantificación del grado de verdad de predicados lógicos. Siguiendo el ejemplo anterior, se define el siguiente predicado en un determinado contexto:

$p$  = “En la sección B hay pocos sectores”.

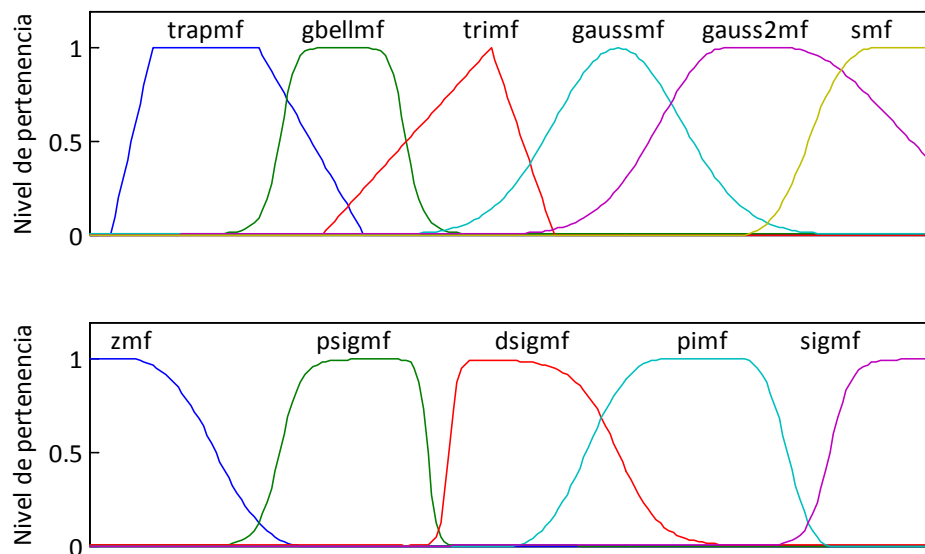
El valor de verdad difuso del predicado  $p$  podrá ser determinado utilizando la definición del conjunto difuso “POCOS” y sabiendo cuántos sectores hay “en la sección B”. El grado de pertenencia al conjunto difuso “POCOS” será el grado de verdad del predicado  $p$ , que se escribirá como  $\mu_p(x)$  siendo  $x$  la “cantidad de sectores en la sección B”.

El razonamiento humano es de tipo aproximado. Las proposiciones dependen del contexto en que se declaran y el ser humano describe su mundo físico y espiritual en términos vagos. La definición de conceptos imprecisos forma parte del pensamiento humano [Kecman, 2001].

Por ejemplo, puede utilizarse un conjunto difuso en la modelización de una expresión que será parte fundamental del trabajo realizado en esta tesis: “gris oscuro”, refiriéndose a una región de una imagen. Este concepto es impreciso, relativo y subjetivo; depende de la “luminosidad” de la imagen en general, de la percepción, del experto en imágenes. Este concepto es un firme candidato a ser modelado con conjuntos difusos.

Usualmente las funciones de pertenencia que definen los conjuntos difusos pueden tener distintas formas y su elección depende de cada caso particular y es subjetiva y dependiente del problema concreto.

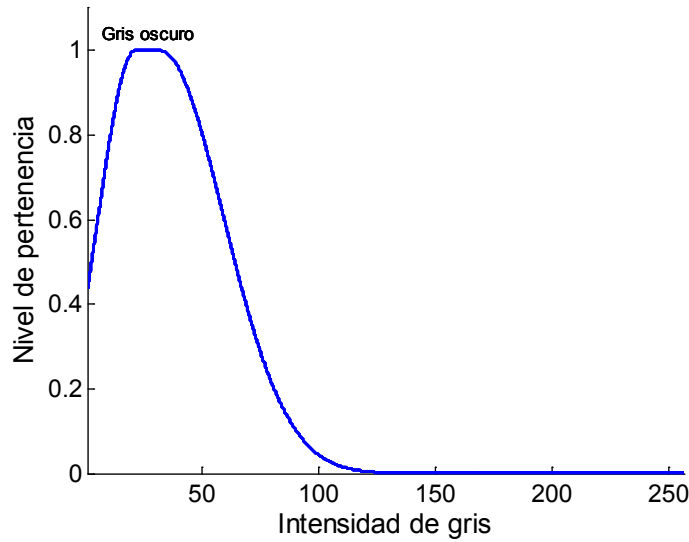
En general, cualquier función  $\mu(x) \rightarrow [0,1]$  puede ser utilizada para definir una función de pertenencia asociada a un conjunto difuso. La Figura 17 muestra algunas formas habituales, implementadas en MatLab®, donde pueden verse con los nombres con que se denominan las funciones que permiten obtenerlas.



**Figura 17: Diferentes formas de funciones de pertenencia.**

Se muestra la forma que adoptan las diferentes opciones de funciones de pertenencia que se encuentran en el software MatLab®. Se incluyen los nombres de las funciones que las generan.

Siguiendo con el ejemplo “gris oscuro”, considerando como variable a la intensidad de gris que presenta una región de una imagen en el rango que va de 1 a 256, un posible conjunto difuso continuo que podría representar este concepto puede verse en la Figura 18.



**Figura 18: Función de pertenencia para definir el conjunto difuso “Gris oscuro” aplicado a la variable “Intensidad de Gris”**

### 3.3.3. Operaciones entre Conjuntos Difusos

Las principales operaciones entre conjuntos son el complemento, la intersección y la unión. La intersección se asocia con el operador lógico AND, la unión con el operador lógico OR y el complemento con el operador lógico NOT.

Hay una gran variedad de operadores que pueden ser utilizados para modelar la intersección y la unión. En las aplicaciones de ingeniería tradicionalmente se utilizan las funciones mínimo (*min*) y máximo (*max*) para modelar la intersección y la unión respectivamente. Si se utilizaran estos operadores, se tendría la siguiente definición:

DEFINICIÓN: sean  $A$  y  $B$  dos conjuntos difusos,  $A = \{(x, \mu_A(x)) \text{ con } x \in X\}$  y  $B = \{(x, \mu_B(x)) \text{ con } x \in X\}$ . Se definen nuevos conjuntos difusos  $A \wedge B$  (intersección) y  $A \vee B$  (unión) según las siguientes expresiones:

$$\mu_{A \wedge B}(x) = \min[\mu_A(x), \mu_B(x)] , \quad (3)$$

$$\mu_{A \vee B}(x) = \max[\mu_A(x), \mu_B(x)] . \quad (4)$$

Pueden utilizarse otras definiciones para las operaciones intersección y unión. En la teoría de LD, los operadores intersección se denominan **T-Normas** y los operadores para la unión se denominan **S-Normas**. La Tabla IV muestra algunos tipos de T-Normas y S-Normas que han sido definidos.

En secciones posteriores se presentarán algunas ventajas y desventajas que presentan las diferentes T y S-Normas y se propondrán nuevos operadores.

**Tabla IV: Diferentes operadores para unión e intersección.**

| <b>Intersección (AND)</b><br><b>T-Norma</b><br>$\mu_{A \wedge B}(x) = T[\mu_A(x), \mu_B(x)]$                       | <b>Unión (OR)</b><br><b>S-Norma</b><br>$\mu_{A \vee B}(x) = S[\mu_A(x), \mu_B(x)]$                                     |
|--------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------|
| <b>Mínimo</b><br>$\min[\mu_A(x), \mu_B(x)]$                                                                        | <b>Máximo</b><br>$\max[\mu_A(x), \mu_B(x)]$                                                                            |
| <b>Producto Algebraico</b><br>$\mu_A(x) \cdot \mu_B(x)$                                                            | <b>Suma Algebraica</b><br>$\mu_A(x) + \mu_B(x) - \mu_A(x) \cdot \mu_B(x)$                                              |
| <b>Producto drástico</b><br>$\min[\mu_A(x), \mu_B(x)]$ si<br>$\max[\mu_A(x), \mu_B(x)] = 1$<br>0 en otro caso      | <b>Suma drástica</b><br>$\max[\mu_A(x), \mu_B(x)]$ si<br>$\min[\mu_A(x), \mu_B(x)] = 0$<br>1 en otro caso              |
| <b>AND de Lukasiewicz</b><br>$\max[0, \mu_A(x) + \mu_B(x) - 1]$                                                    | <b>OR de Lukasiewicz</b><br>$\min[1, \mu_A(x) + \mu_B(x)]$                                                             |
| <b>Producto de Einstein</b><br>$\frac{\mu_A(x) \cdot \mu_B(x)}{2 - \mu_A(x) + \mu_B(x) - \mu_A(x) \cdot \mu_B(x)}$ | <b>Suma de Einstein</b><br>$\frac{\mu_A(x) + \mu_B(x)}{1 + \mu_A(x) \cdot \mu_B(x)}$                                   |
| <b>Producto de Hamacher</b><br>$\frac{\mu_A(x) \cdot \mu_B(x)}{\mu_A(x) + \mu_B(x) - \mu_A(x) \cdot \mu_B(x)}$     | <b>Suma de Hamacher</b><br>$\frac{\mu_A(x) + \mu_B(x) - 2 \cdot \mu_A(x) \cdot \mu_B(x)}{1 - \mu_A(x) \cdot \mu_B(x)}$ |
| <b>Operador de Yager</b><br>$1 - \min[1, [1 - \mu_A(x)]^b + [1 - \mu_B(x)]^b]^{\frac{1}{b}}$                       | <b>Operador de Yager</b><br>$\min[1, [\mu_A(x)^b + \mu_B(x)^b]^{\frac{1}{b}}]$                                         |

### 3.4. Toma de decisiones con Lógica Difusa

Los modelos matemáticos de la racionalidad son frecuentemente la base de los métodos y sistemas de ayuda a la decisión [French, 1986]. Al margen de las aplicaciones exitosas, que demuestran sus cualidades, estos modelos tienen, sin embargo, dificultades para lidiar con la subjetividad humana que caracteriza la amplia variedad de situaciones de decisión.

El uso de la Teoría de Probabilidades, esencialmente a partir del concepto de Probabilidad Subjetiva (entendiendo como tal a la probabilidad de que suceda un

evento específico que asigna una persona con base en cualquier información disponible), para incorporar la incertidumbre asociada con la representación de escenarios posibles, tropieza con dificultades prácticas [Kahneman and Tversky, 2005, Tversky and Kahneman, 2005, Hogarth, 1991]. Las exigencias axiomáticas del concepto de probabilidad dificultan el manejo de valores aportados subjetivamente. Asimismo, la construcción de funciones de preferencia de valor cardinal o funciones utilidad, tropieza con problemas prácticos asociados a sus comportamientos axiomáticos [Hogarth, 1991].

Por otra parte, la Psicología y la Economía Experimental han estudiado durante las últimas décadas la manera en que el hombre juzga y elige, aportando evidencias del incumplimiento de los modelos racionales con la conducta humana real en situaciones de decisión [Fox, 1994, Smith, 2000, Tversky and Kahneman, 2005]. Además aportan un estudio fundamental sobre sesgos asociados con la valoración de probabilidades y el comportamiento decisorio en condiciones de riesgo. Entre otros resultados, encuentran que los decisores actúan de acuerdo con una valoración de pérdidas y ganancias, y que esta idea describe mejor el comportamiento de los decisores humanos.

Otro avance en el ámbito de la Toma de Decisiones es el desarrollo de sistemas expertos, derivado particularmente de la Inteligencia Artificial [Krause and Clark, 1994, Fox, 1994]. Aunque la Lógica Matemática, y en particular los sistemas multivalentes han estado relacionados con la Inteligencia Artificial, el énfasis de estos sistemas está generalmente orientado a reglas no formalizadas matemáticamente y a modelos matemáticos de propagación de la incertidumbre. Sin embargo, los sistemas expertos son pioneros en la idea de obtener modelos partiendo de expresiones verbales, de manera que los decisores pueden aplicar su experiencia esencial en problemas concretos.

En los últimos tiempos se ha desarrollado una disciplina matemático-informática llamada *Soft-Computing* o Inteligencia Computacional [Bonissone, 1997, Verdegay, 2005]. Entre los fundamentos de esta disciplina está la LD [Zadeh, 1965, Dubois and Prade, 1980, Lindley, 1994]. Una parte considerable de los modelos inteligentes utiliza modelos difusos de manera pura o en combinación con elementos

como las RN y los Algoritmos Evolutivos (Modelos Híbridos) combinando poco a poco la Investigación Operativa y la Inteligencia Artificial. La “vaguedad” es, junto a la incertidumbre, objeto de modelado de la LD, y esta perspectiva ha permitido la modelado del conocimiento y la toma de decisiones basados en expresiones verbales como información de entrada [Fox, 1994].

### 3.4.1. La Lógica Difusa y la modelado de la decisión

Una manera de poner en práctica el “Principio de gradualidad”, propiedad esencial de la LD, es la definición de lógicas donde los predicados son funciones del universo  $X$  en el intervalo  $[0,1]$  y las operaciones de conjunción, disyunción, negación, e implicación, se definen de modo que al ser restringidas al dominio  $\{0,1\}$  se transforman en la Lógica Booleana. Las distintas formas de definir las operaciones y sus propiedades determinan diferentes lógicas multivalentes que son parte del paradigma de la LD [Dubois and Prade, 1980]. Las lógicas multivalentes se definen en general como aquéllas que permiten valores intermedios entre la verdad absoluta y la falsedad total de una expresión.

Las aplicaciones de esta disciplina en el campo de la Toma de Decisiones han sido hechas básicamente a partir del concepto de operador, más que en el de Lógica Multivalente [Dubois and Prade, 1985]. Los operadores son clasificados en conjuntivos, disyuntivos e interactivos y utilizados por analistas de la decisión de acuerdo con su experiencia y su intuición para lograr a través de la selección de alguno de ellos una “confluencia” de objetivos y restricciones. Sin embargo, esta manera de abordar las decisiones no proporciona la mejor base para utilizar la capacidad de la LD para la transformación del conocimiento y las preferencias de los expertos humanos en fórmulas lógicas; en otras palabras, no permite usar esta lógica a la manera de la Ingeniería del Conocimiento. El uso del lenguaje como elemento de comunicación entre un Ingeniero del Conocimiento y un Experto apunta más al uso de una combinación armónica de operadores, que hacia el uso de sólo uno de ellos.

La LD ha sido aplicada con buenos resultados al Control Automático para decidir a través de reglas los valores aconsejables de las variables de control. El Control Difuso suele escoger reglas partiendo de una formulación verbal de las mismas,

utilizando las llamadas “variables lingüísticas”. Se seleccionan de manera pragmática operadores que definen conectivas lógicas, incluidas en las reglas, y un procedimiento llamado *defuzzificación*. Éste es un recurso extra-lógico asociado con la incapacidad de los modelos lógicos multivalentes para obtener resultados favorables en esta esfera. La *defuzzificación* actúa como un grado de libertad en un modelo basado en la combinación pragmática de operadores, pero sin un enlace axiomático armónico que justifique la denominación de “Lógica”. En este enfoque se recomienda el uso de reglas poco complejas; la selección pragmática de operadores, en combinación con la *defuzzificación*, sólo permite buenos resultados con reglas simples [Passino and Yurkovich, 1998, Zimmermann, 2001].

En el campo de la Toma de Decisiones se necesitan desarrollos teóricos para el logro de mejores resultados. Para su aplicación a la toma de decisiones es deseable, por ejemplo, que los valores de verdad de las Lógicas Multivalentes posean sensibilidad a los cambios de los valores de verdad de los predicados básicos y conserven el “significado verbal” de los valores de verdad calculados. Esto es difícil de lograr con los requerimientos axiomáticos tradicionales. En esta tesis se hace especial hincapié en la LDC [Espín Andrade et al., 2007], que constituye una Lógica Multivalente que supera las dificultades señaladas. Ésta se propone como una alternativa desde posiciones de la lógica al enfoque normativo de la decisión, uniendo la modelación de la decisión y el razonamiento sobre bases afines al paradigma racional que aquél sustenta.

### 3.4.2. Expresiones difusas y diferentes lógicas

Así como es posible operar con los valores de pertenencia a conjuntos difusos, se definen operaciones entre valores de verdad de expresiones o predicados difusos como una extensión natural de la lógica tradicional o bivaluada [Mendel, 2001]. Algunas de las equivalencias matemáticas más importantes entre lógica y teoría de conjuntos son:

| Lógica           | Teoría de Conjuntos                  |
|------------------|--------------------------------------|
| AND [ $\wedge$ ] | Intersección [ $\cap$ ]              |
| OR [ $\vee$ ]    | Unión [ $\cup$ ]                     |
| NOT [ $-$ ]      | Complemento [ $\overline{(\cdot)}$ ] |

Adicionalmente, hay una correspondencia entre lógica elemental y álgebra booleana. Cualquier predicado o sentencia que sea verdadera en un sistema lo es en el otro, simplemente a través de los siguientes cambios en la notación:

| Lógica                             | Álgebra de Boole    |
|------------------------------------|---------------------|
| Verdadero                          | 1                   |
| Falso                              | 0                   |
| AND [ $\wedge$ ]                   | $\times$            |
| OR [ $\vee$ ]                      | $+$                 |
| NOT [ $-$ ]                        | Complemento [ $'$ ] |
| Equivalencia [ $\leftrightarrow$ ] | $=$                 |

Sean  $x_1, x_2, \dots$  variables que toman valores en  $[0,1]$ . Se definen las operaciones negación, conjunción y disyunción entre estas variables [Dubois and Prade, 1980].

Una **expresión difusa** es una función de  $[0,1]^n$  en  $[0,1]$  definida por las siguientes reglas:

- 0, 1 y las variables  $x_i$ ,  $i = 1, \dots, n$  son expresiones difusas.
- Si  $f$  es una expresión difusa, entonces  $-f$  también lo es.
- Si  $f$  y  $g$  es una expresión difusa, entonces  $f \wedge g$  y  $f \vee g$  también lo son.

Todas las expresiones booleanas pueden ser expresiones difusas (lógicas multivalentes) si se extiende su dominio a  $[0,1]^n$ . Esto se debe a que todas las expresiones booleanas pueden expresarse sólo en términos de los operadores “-” (*not*), “ $\wedge$ ” (*and*) y “ $\vee$ ” (*or*). Análogamente al caso de las operaciones entre conjuntos difusos, como se vio en la Tabla IV, existen operadores diferentes para el cálculo de estas operaciones.

La evaluación de expresiones difusas puede escribirse de la siguiente manera: se denota como  $v(P)$  al valor de verdad de una proposición  $P$ , siendo  $v(P) \in [0,1]$ , así como el grado de pertenencia a un conjunto difuso se escribe como  $\mu_A(x) \in [0,1]$ .

Cada conjunto de operadores difusos propuesto tiene sus propias propiedades y características. La evaluación de la negación suele ser común a todos ellos y se da como  $v(-P) = 1 - v(P)$ , y por lo tanto  $v(- - P) = v(P)$ .

Todas las lógicas multivalentes constituyen extensiones de la lógica bivalente clásica. Para cada una de ellas puede analizarse si se cumplen las propiedades que se exponen en la Tabla V.

Como ejemplo, tomando las operaciones  $v(P \wedge Q) = \min[v(P), v(Q)]$  y  $v(P \vee Q) = \max[v(P), v(Q)]$  para la conjunción y disyunción respectivamente, se cumplirán todas las propiedades enunciadas, a excepción de la ley del tercero excluido, que no se satisfará.

Tomando las operaciones:

$$v(P \wedge Q) = \max[0, v(P) + v(Q) - 1],$$

$$v(P \vee Q) = \min[1, v(P) + v(Q)]$$

para la conjunción y disyunción respectivamente (AND y OR de Lukasiewicz), se cumplirán las propiedades conmutativa, asociativa, De Morgan y la ley del tercero excluido pero no se cumplirán la idempotencia y la propiedad distributiva.

**Tabla V: Propiedades de la conjunción y la disyunción.**

|                          |                                                                                                                            |
|--------------------------|----------------------------------------------------------------------------------------------------------------------------|
| Conmutativa              | $v(P \wedge Q) = v(Q \wedge P)$<br>$v(P \vee Q) = v(Q \vee P)$                                                             |
| Asociativa               | $v[(P \wedge Q) \wedge R] = v[P \wedge (Q \wedge R)]$<br>$v[(P \vee Q) \vee R] = v[P \vee (Q \vee R)]$                     |
| Idempotente              | $v(P \wedge P \wedge P \dots \wedge P) = v(P)$<br>$v(P \vee P \vee P \dots \vee P) = v(P)$                                 |
| Distributivas            | $v[(P \wedge Q) \vee R] = v[(P \vee R) \wedge (Q \vee R)]$<br>$v[(P \vee Q) \wedge R] = v[(P \wedge R) \vee (Q \wedge R)]$ |
| Ley del tercero excluido | $v(P \vee \neg P) \neq 1$<br>$v(P \wedge \neg P) \neq 0$                                                                   |
| Absorción                | $v[P \vee (P \wedge Q)] = v(P)$<br>$v[P \wedge (P \vee Q)] = v(P)$                                                         |
| De Morgan                | $v[\neg(P \wedge Q)] = v(\neg P \vee \neg Q)$<br>$v[\neg(P \vee Q)] = v(\neg P \wedge \neg Q)$                             |
| Equivalencia             | $v[(\neg P \vee Q) \wedge (P \vee \neg Q)] = v[(P \wedge Q) \vee (\neg P \wedge \neg Q)]$                                  |
| Disyunción exclusiva     | $v[(\neg P \wedge Q) \wedge (P \wedge \neg Q)] = v[(P \vee Q) \wedge (\neg P \vee \neg Q)]$                                |

Tomando las operaciones  $v(P \wedge Q) = v(P) \cdot v(Q)$  y  $v(P \vee Q) = v(P) + v(Q) - v(P) \cdot v(Q)$  para la conjunción y disyunción respectivamente (AND y OR probabilísticos), se cumplirán las propiedades conmutativa, asociativa y De Morgan pero no se cumplirán la idempotencia y la propiedad distributiva.

Cada una de estas lógicas propuestas considera a su vez diferentes definiciones para la implicación, basadas en las operaciones de conjunción y disyunción. La implicación se define comúnmente como  $v(P \Rightarrow Q) = v(\neg P \vee Q)$  pero se han propuesto otras definiciones que se dan en la Tabla VI.

**Tabla VI: Diferentes definiciones para la implicación.**

|                                                                                                   |
|---------------------------------------------------------------------------------------------------|
| Implicación de Zadeh                                                                              |
| $v(P \Rightarrow Q) = v[\neg P \vee (P \wedge Q)]$                                                |
| Implicación probabilística                                                                        |
| $v(P \Rightarrow Q) = 1 - v(P) + v(P)v(Q)$                                                        |
| Implicación de Lukasiewicz                                                                        |
| $v(P \Rightarrow Q) = \min[1, 1 - v(P) + v(P)v(Q)]$                                               |
| Implicación Brouweriana                                                                           |
| $v(P \Rightarrow Q) = \begin{cases} 1 & v(P) \leq v(Q) \\ v(Q) & \text{en otro caso} \end{cases}$ |
| Otra implicación considerada por Zadeh                                                            |
| $v(P \Rightarrow Q) = \max[1 - v(P), \min(v(P), v(Q))]$                                           |

La equivalencia o doble implicación se define como:

$$v(P \Leftrightarrow Q) = v(P \Rightarrow Q) \wedge v(Q \Rightarrow P).$$

### 3.4.3. Modificadores

Un modificador es una operación que hace variar el significado de un término en una expresión lógica o también el nombre de un conjunto difuso. Por ejemplo, si se considera la expresión “La presión es baja”, entonces “La presión es muy baja”, “La presión es más o menos baja”, “La presión es extremadamente baja” hacen uso de modificadores que se han aplicado al concepto “baja”, que a su vez puede ser cuantificado a través de un conjunto difuso.

Se dan a continuación algunos ejemplos pero pueden encontrarse muchas otras definiciones en la literatura [Cox, 1994].

#### 3.4.3.1. Concentración

Este operador se define mediante la operación:

$$con[v(P)] = v(P)^N. \quad (5)$$

Como ejemplos,

$$v(\text{"La presión es muy baja"}) = v(\text{"La presión es baja"})^2,$$

$$v(\text{"La presión es extremadamente baja"}) = v(\text{"La presión es baja"})^4.$$

### 3.4.3.2. Dilatación

Este operador se define mediante la operación:

$$dil[v(P)] = v(P)^{\frac{1}{N}}. \quad (6)$$

Como ejemplos,

$$v(\text{"La presión es más o menos baja"}) = v(\text{"La presión es baja"})^{\frac{1}{2}},$$

$$v(\text{"La presión es algo baja"}) = v(\text{"La presión es baja"})^{\frac{1}{4}}.$$

Como alternativa a las definiciones de diferentes lógicas presentadas, se presenta a continuación la LDC [Espín Andrade et al., 2007, Espin Andrade et al., 2003, Espin Andrade et al., 2004], adecuada a la resolución del problema que se plantea en la presente tesis.

### 3.4.4. Lógica Difusa Compensatoria (LDC)

En los procesos que requieren toma de decisiones, el intercambio con los expertos lleva a obtener formulaciones complejas y sutiles que requieren de predicados compuestos. Los valores de verdad obtenidos sobre estos predicados compuestos deben poseer sensibilidad a los cambios de los valores de verdad de los predicados básicos. Esta necesidad se satisface con el uso de la LDC, que renuncia al cumplimiento de las propiedades clásicas de la conjunción y la disyunción, contraponiendo a éstas la idea de que el aumento o disminución del valor de verdad de la conjunción o la disyunción provocadas por el cambio del valor de verdad de una de sus componentes, puede ser “compensado” con la correspondiente disminución o aumento de la otra.

La LDC es un modelo lógico multivalente que renuncia a varios axiomas clásicos para lograr un sistema multivalente idempotente y “sensible” que asimila virtudes de la escuela descriptiva y la escuela normativa de la Toma de Decisiones, pues permite la compensación de los atributos (en este caso los valores de verdad de predicados simples), pero si son violados ciertos umbrales hay un veto que impide la compensación.

Al mismo tiempo, las propiedades que satisface hacen posible de manera natural el trabajo de traducción del lenguaje natural al de la Lógica, incluidos los predicados extensos si éstos surgen del proceso de modelación.

Se denotará, para simplificar la notación, como  $\mu_{p1} = v(P_1)$ ,  $\mu_{p2} = v(P_2)$ , ...,  $\mu_{pn} = v(P_n)$  a los valores de verdad de los predicados  $P_1, P_2, \dots, P_n$ .

En esta propuesta, el **operador conjunción**, expresado como **c** es la media geométrica:

$$c(p_1, p_2, \dots, p_n) = (\mu_{p1} \times \mu_{p2} \dots \times \mu_{pn})^{1/n}. \quad (7)$$

La **negación n** es la función:

$$n(p) = 1 - \mu_p. \quad (8)$$

La **disyunción d** es el operador dual de la media geométrica, que garantiza el cumplimiento de las reglas de De Morgan [Espin Andrade et al., 2004]:

$$d(p_1, p_2 \dots, p_n) = 1 - [(1 - \mu_{p1})(1 - \mu_{p2}) \dots (1 - \mu_{pn})]^{1/n}. \quad (9)$$

La **implicación** con propiedades más deseables es la implicación de Zadeh generalizada:

$$p \Rightarrow q = -p \vee (p \wedge q),$$

$$i(p, q) = d[n(p), c(p, q)];$$

aunque ha sido estudiada también la implicación siguiente:

$$p \Rightarrow q = -p \vee q,$$

$$i(p, q) = d[n(p), q].$$

La **equivalencia** ("si y sólo si") se define a partir del operador implicación de la siguiente manera:

$$p \Leftrightarrow q = (p \Rightarrow q) \wedge (q \Rightarrow p),$$

$$e(p, q) = c[i(p, q), i(q, p)].$$

Los **cuantificadores universal y existencial** son introducidos a través de las siguientes fórmulas:

$$\begin{aligned}
\forall_{x \in U} p(x) &= \bigwedge_{x \in U} p(x) = \sqrt[n]{\prod_{x \in U} p(x)} = \\
&= \begin{cases} \exp \left[ \frac{1}{n} \sum_{x \in U} \ln[p(x)] \right], & \text{si } p(x) \neq 0 \\ 0, & \text{en cualquier otro caso} \end{cases}, \\
\exists_{x \in U} p(x) &= \bigvee_{x \in U} p(x) = 1 - \sqrt[n]{\prod_{x \in U} [1 - p(x)]} = \\
&= \begin{cases} 1 - \exp \left[ \frac{1}{n} \sum_{x \in U} \ln[1 - p(x)] \right], & \text{si } p(x) \neq 0 \\ 0, & \text{en cualquier otro caso} \end{cases}.
\end{aligned}$$

### 3.4.5. Modelización de expresiones

Como ejemplo de construcción de predicados a partir de una expresión verbal, se describirá brevemente a continuación un modelo capaz de ordenar un conjunto de empresas con respecto a una línea de productos en un mercado competitivo. En su construcción participaron como expertos especialistas de la empresa consultora BIOMUNDI, la cual ha alcanzado un gran desarrollo en la oferta de servicios de inteligencia competitiva en Cuba.

A continuación aparecen las formulaciones verbales y su traducción al lenguaje del Cálculo de Predicados:

Una empresa es competitiva si cumple los siguientes requisitos:

- La economía de la empresa es sólida
  - Una empresa es económicamente sólida si tiene un buen estado financiero y buenas ventas. Si el estado financiero fuera algo malo debe ser compensado con muy buenas ventas.
- Su posición tecnológica es de avanzada.
  - Una empresa tiene una posición tecnológica de avanzada si su tecnología actual es buena y además es dueña de patentes, o tiene productos en investigación desarrollo, o dedica cantidades importantes de dinero a esta

actividad. Si su tecnología es algo atrasada, entonces debe tener muchas patentes, o muchos productos en investigación desarrollo, o dedicar cantidades muy importantes de recursos a esta actividad.

- Es muy fuerte en una línea de productos en un mercado dado
  - Una empresa es fuerte en una línea de productos si tiene una posición fuerte en el mercado, la línea de productos es variada y tiene independencia del proveedor.

Se definen los siguientes predicados simples:

$C(x)$  = “La empresa x es competitiva”

$S(x)$  = “La empresa x tiene una economía sólida”

$T(x)$  = “La empresa x tiene una posición tecnológica de avanzada”

$L(x)$  = “La empresa x es fuerte en la línea de productos”

$EF(x)$  = “La empresa x tiene un buen estado financiero”

$BV(x)$  = “La empresa x tiene buenas ventas”

$TB(x)$  = “La empresa x tiene actualmente buena tecnología”

$DP(x)$  = “La empresa x es dueña de patentes”

$ID(x)$  = “La empresa x tiene productos en investigación y desarrollo”

$DD(x)$  = “La empresa x dedica gran cantidad de dinero en I+D”

$FM(x)$  = “La empresa x tiene fortaleza en el mercado”

$LV(x)$  = “La empresa x tiene una línea variada de productos”

$IP(x)$  = “La empresa x es independiente del proveedor”.

Entonces el modelo es el siguiente predicado compuesto:

$$C(x) = S(x) \wedge T(x) \wedge L^2(x),$$

donde:

$$S(x) = EF(x) \wedge BV(x) \wedge [-EF^{0.5}(x) \rightarrow BV^2(x)],$$

$$T(x) = TB(x) \wedge [DP(x) \vee ID(x) \vee DD(x)]$$

$$\wedge [-TB^{0.5}(x) \rightarrow [DP^2(x) \vee ID^2(x) \vee DD^2(x)]],$$

$$L(x) = FM(x) \wedge LV(x) \wedge IP(x).$$

En este modelo se utilizan modificadores lingüísticos basados en los operadores de Novak con el propósito de modelar los adverbios “muy” y “algo”.

Con los valores de verdad de los predicados simples se procede a calcular los valores de verdad de los predicados compuestos  $S(x)$ ,  $T(x)$  y  $L(x)$  para determinar finalmente el grado de verdad de  $C(x)$ . Así, un conjunto de empresas podría ordenarse según estos grados de verdad para determinar un ranking de “empresas competitivas” de mayor a menor.

Los modelos así obtenidos pueden ser optimizados en problemas para los que se dispone de datos reales numéricos. El método de optimización también puede provenir de la Inteligencia Computacional. Este enfoque constituye el fundamento de los sistemas híbridos. En este contexto, los AG presentan una alternativa interesante y utilizada con éxito en casos similares.

---

## 4. Procesamiento de Imágenes de Resonancia Magnética: método propuesto

## 4. Procesamiento de Imágenes de Resonancia Magnética: método propuesto

---

### 4.1. Introducción

---

Los médicos especialistas en diagnóstico por imágenes suelen explicar la interpretación de las IRM enunciando algunas consideraciones basadas en inspecciones visuales de esas imágenes que luego aplican de forma automática. Para explicar cómo reconocer regiones en las que se encuentra agua o algún líquido suelen enunciar sentencias del tipo *“el líquido se ve negro en T1, blanco en T2”*. Implícitamente lo que se está explicando es que, cuando se trata de agua, en la secuencia  $T_1$  se verán regiones “oscuras” y las mismas regiones se verán “claras” en la secuencia  $T_2$ .

Se tienen entonces conceptos sencillos de interpretar visualmente, pero no adecuados para la implementación computacional con herramientas tradicionales de segmentación.

Los primeros cuestionamientos que surgen son: ¿qué intensidades de gris presentan las regiones de la imagen consideradas “negras” u “oscuras”? ¿Cuáles intensidades representan la idea de “blanco” o “brillante”?

En esta etapa resulta de suma importancia contar con imágenes previamente clasificadas, en lo posible por los mismos expertos que enuncian los predicados utilizados para la discriminación de los tejidos presentes, con el fin de hacer un análisis cuantitativo de estas ideas.

La LD es un enfoque que suele resolver con éxito este tipo de situaciones, puesto que permite utilizar estas ideas expresadas lingüísticamente para efectuar cálculos matemáticos objetivos que reflejen el resultado razonado.

Para formalizar un proceso de cálculo se requiere expresar predicados diferentes para cada tejido que se desea detectar en un determinado tipo de imagen.

En especial las IRM son interesantes para este enfoque por la riqueza de la complementariedad de información que presentan las distintas secuencias de adquisición.

En las siguientes secciones se explica de manera detallada el proceso por el cual se propone el análisis de IRM píxel a píxel en el cual para cada tejido a reconocer se ha definido un predicado que involucra las intensidades presentes en las secuencias T1, T2, y eventualmente, de ser necesaria, la secuencia PD en imágenes de cerebros.

El objetivo es decidir a qué tejido, dentro de los posibles, corresponde cada píxel de la imagen. Los tejidos (o sustancias) a reconocer preliminarmente son: LCR, MB, MG. También se reconocen los píxeles de fondo.

## 4.2. Determinación de predicados

---

En la etapa inicial de este trabajo se establecieron extensas entrevistas con especialistas en diagnóstico por imágenes con el fin de determinar los criterios utilizados para la interpretación de las IRM.

Luego del análisis de los conceptos aportados por los especialistas, se elaboró un conjunto de predicados y se interactuó nuevamente con los expertos a fin de expresarlos según su parecer.

De esta manera, a través de estas entrevistas, investigando sobre distintos tipos de imágenes, reconociendo diferentes sustancias o tejidos en las imágenes, surgieron diversas descripciones subjetivas, formales e informales, como por ejemplo “gris claro”, “intensidad constante”, “zonas brillantes”, “gris”, “gris intenso”, “regiones hipointensas”, “prácticamente negro”, “blanco”, “gris oscuro”, y todas sus combinaciones a través de conectivos lógicos.

Mediante el consenso de varios expertos se determinó el siguiente conjunto preliminar de predicados que describen las sustancias:

---

$E_1 =$  “El líquido cefalorraquídeo es *Oscuro* en T1, *Muy Intenso* en T2 y *Muy Intenso* en PD o también puede ser *Blanco* en T2.”

$E_2 =$  “La materia gris es *Gris* en T1, *Intenso* en T2 y *Muy Intenso* pero *No-Blanco* en PD.”

$E_3 =$  “La materia blanca es *Intenso* en T1, *Gris* en T2 e *Intenso* en PD.”

$E_4 =$  “El fondo es *Negro* en T1 y *Negro* en T2.”

---

Acercando más estas expresiones a una formalización, utilizando siglas para los nombres de las sustancias para abreviar y el símbolo  $v(X)$  para denotar al grado de verdad del predicado  $X$ , se pueden escribir las siguientes expresiones lingüísticas para adecuarlas al análisis de cada píxel. Formalizando estos predicados se tienen las expresiones lógicas que se dan a continuación:

---

$E_1 =$  “El píxel corresponde a LCR” cuando “el píxel es *Oscuro* en T1”, “el píxel es *Muy Intenso* en T2” o “el píxel es *Blanco* en T2” y “el píxel es *Muy Intenso* en PD”

$$v(P_1) = v(P_{11}) \wedge [v(P_{12}) \vee v(P_{13})] \wedge v(P_{14})$$

$P_1 =$  “El píxel corresponde a LCR”

$P_{11} =$  “El píxel es *Oscuro* en T1”

$P_{12} =$  “El píxel es *Muy Intenso* en T2”

$P_{13} =$  “El píxel es *Blanco* en T2”

$P_{14} =$  “El píxel es *Muy Intenso* en PD”

---

$E_2 =$  “El píxel corresponde a MG” cuando “el píxel es *Gris* en T1”, “el píxel es *Intenso* en T2” y “el píxel es *Muy Intenso* en PD” pero “el píxel No es *Blanco* en PD”

$$v(P_2) = v(P_{21}) \wedge v(P_{22}) \wedge [v(P_{23}) \wedge (\neg v(P_{24}))]$$

$P_2 =$  “El píxel corresponde a MG”

$P_{21} =$  “El píxel es *Gris* en T1”

$P_{22} =$  “El píxel es *Intenso* en T2”

$P_{23} =$  “El píxel es *Muy Intenso* en PD”

$P_{24}$  = "El píxel es *Blanco* en PD"

---

$E_3$  = "*El píxel corresponde a MB*" cuando "*el píxel es Intenso en T1*", "*el píxel es Gris en T2*" y "*el píxel es Intenso en PD*"

$$v(P_3) = v(P_{31}) \wedge v(P_{32}) \wedge v(P_{33})$$

$P_3$  = "El píxel corresponde a MB"

$P_{31}$  = "El píxel es *Intenso* en T1"

$P_{32}$  = "El píxel es *Gris* en T2"

$P_{33}$  = "El píxel es *Intenso* en PD"

---

$E_4$  = "*El píxel corresponde a Fondo*" cuando "*el píxel es Negro en T1*" y "*el píxel es Negro en T2*"

$$v(P_4) = v(P_{41}) \wedge v(P_{42})$$

$P_4$  = "El píxel corresponde a Fondo"

$P_{41}$  = "El píxel es *Negro* en T1"

$P_{42}$  = "El píxel es *Negro* en T2"

---

El valor de verdad logrado en cada predicado no debe ser tomado como absoluto, sino que intervendrá como resultado parcial en la toma de decisión de la sustancia que se asignará a cada píxel. Precisamente, al píxel en análisis se le asignará la sustancia para la que corresponda el máximo grado de verdad obtenido entre todos los predicados [Espin Andrade et al., 2004].

Algunos tejidos pueden requerir predicados más complejos para su reconocimiento, lo que no implica un problema, dado que pueden incluirse todos los conectores lógicos o incluso implicaciones.

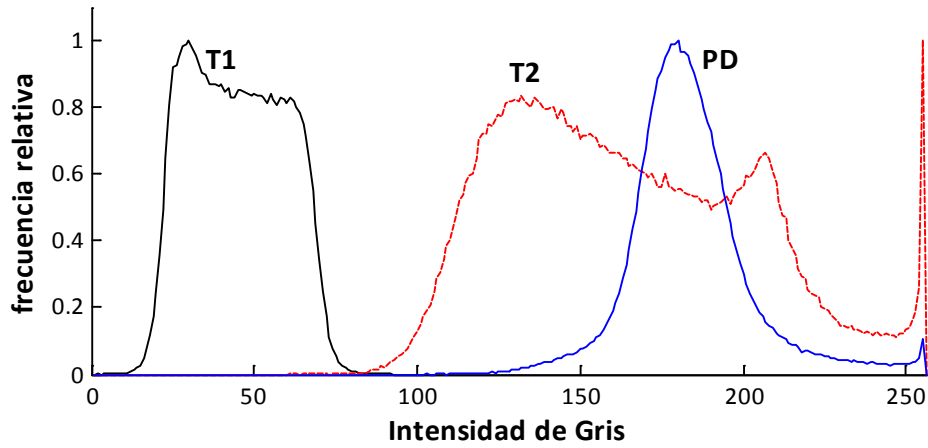
### 4.3. Cuantificación de los valores de verdad de los predicados difusos

---

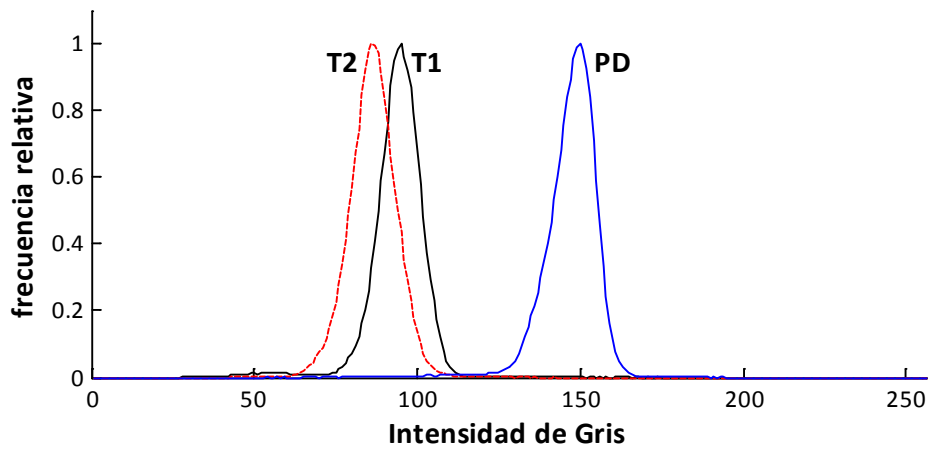
Con el fin de cuantificar el grado de verdad de los predicados definidos en la sección anterior, se procedió a definir conjuntos difusos para representar los conceptos (adjetivos) asociados a la intensidad de gris de las diferentes imágenes. El conjunto soporte para estos conjuntos difusos es el rango completo de valores de gris (en las imágenes procesadas, 8 bits, valores entre 0 y 255).

En una etapa inicial se procedió a la representación con histogramas de las intensidades de gris para las regiones pre-reconocidas de las diferentes sustancias, para lo que se requirió un conjunto de imágenes segmentadas previamente. Con estas representaciones puede comprenderse que las ideas expresadas, como por ejemplo “negro” o “blanco”, en realidad representan intervalos de intensidades de grises, a veces con gran diversidad de valores. Esto se debe a la poca sensibilidad del ojo humano para detectar variaciones de intensidad pequeñas. Así es como, en el intervalo de valores de 0 a 255, los valores entre 0 y 20 son prácticamente iguales y visiblemente “negros” y los valores entre 230 y 255 son prácticamente “blancos”.

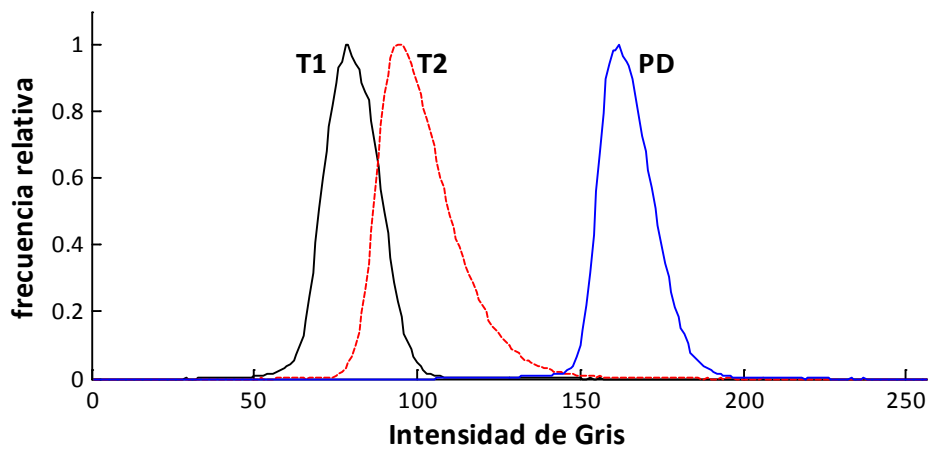
En la Figura 19 pueden verse los histogramas correspondientes a las secuencias T1, T2 y PD para el LCR (a), MB (b) y MG (c). En todos los casos el eje horizontal muestra las intensidades de gris y el eje vertical la cantidad de píxeles, normalizada. Puede concluirse, por ejemplo, que el LCR se encuentra principalmente en el rango de intensidades entre 20 y 60 para la secuencia T1, pero también hay unos pocos píxeles con intensidades entre 10 y 20 o entre 60 y 75.



(a) Líquido Ceforraquídeo



(b) Materia Blanca



(c) Materia Gris

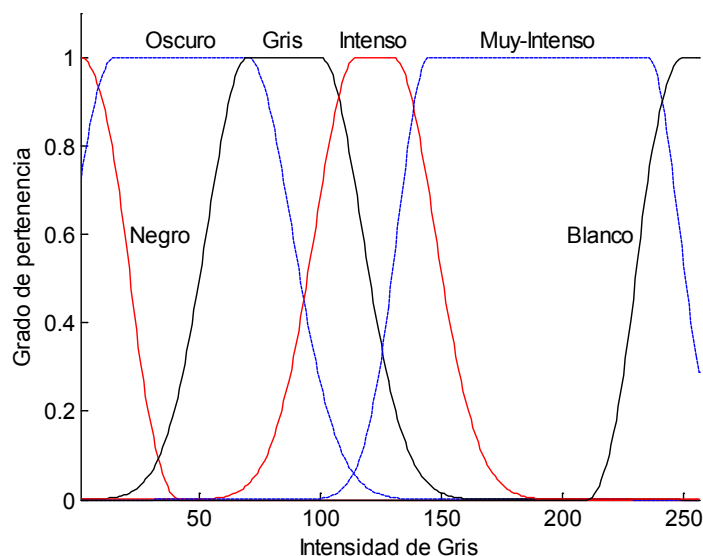
**Figura 19: Histogramas de frecuencias relativas de intensidades de gris.**

Se observan los histogramas de intensidades de las imágenes pesadas en T1, T2 y PD correspondientes al líquido ceforraquídeo (a), materia blanca (b) y materia gris (c).

De las entrevistas con los expertos se tenía el conocimiento de que el LCR se ve “oscuro en T1”. De este análisis debería diseñarse, en principio heurísticamente, una función de pertenencia que represente este concepto. Esta función debería devolver valores cercanos a 1 para los valores de gris en que existe gran cantidad de píxeles “oscuros” en el histograma, 0 para los valores de intensidad en que no aparecen píxeles y su transición debe ser gradual para reflejar el carácter difuso del concepto. Esto lleva a pensar en formas de campanas de gauss para las funciones de pertenencia. Se han utilizado funciones definidas con dos campanas, lo que permite que el crecimiento y decrecimiento de la función sea diferente.

En particular los conceptos “Blanco” y “Negro” abarcarían los rangos superior e inferior de intensidades de gris, por lo cual en estos dos casos es adecuado el empleo de funciones del tipo sigmoidea. Existen dos definiciones que cumplen esta característica: las funciones denominadas “S” y “Z”.

De este análisis surge que el concepto “oscuro” (un conjunto difuso) en este contexto podría ser representado por la función de pertenencia que se visualiza en la Figura 20, indicado con la etiqueta “Oscuro”. De esta manera se completa el análisis para disponer, de una forma preliminar, de todos los conjuntos difusos que aparecen en cada uno de los predicados enunciados.



**Figura 20: Funciones de pertenencia preliminares.**

Gráfico de las funciones de pertenencia, basadas en los histogramas, que definen los conjuntos difusos “Negro”, “Oscuro”, “Gris”, “Intenso”, “Muy Intenso” y “Blanco”. Estas funciones serán posteriormente optimizadas.

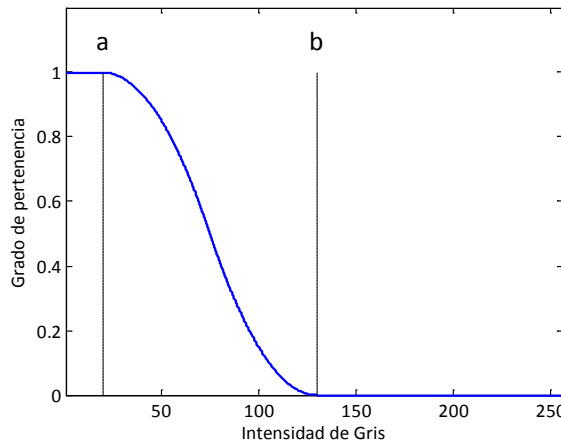
Las funciones de pertenencia de los conjuntos difusos definidos se simbolizan mediante:  $\mu_{NEGRO}(x)$ ,  $\mu_{OSCURO}(x)$ ,  $\mu_{GRIS}(x)$ ,  $\mu_{INTENSO}(x)$ ,  $\mu_{MUY-INTENSO}(x)$ ,  $\mu_{BLANCO}(x)$ , donde  $x \in [0,255]$  representa el valor de intensidad de gris.

La función de pertenencia que define el conjunto difuso “Negro” se eligió de tipo “Z”. Esta función requiere 2 parámetros ( $a$  y  $b$ ) para quedar determinadas su posición y forma (Figura 21):

$$\mu_{NEGRO}(x) = \begin{cases} 1 & , \quad x \leq a \\ 1 - 2 \left( \frac{x-a}{b-a} \right)^2 & , \quad a < x \leq \frac{a+b}{2} \\ 2 \left( \frac{x - \frac{a+b}{2}}{b-a} \right)^2 & , \quad \frac{a+b}{2} < x < b \\ 0 & , \quad x \geq b \end{cases}.$$

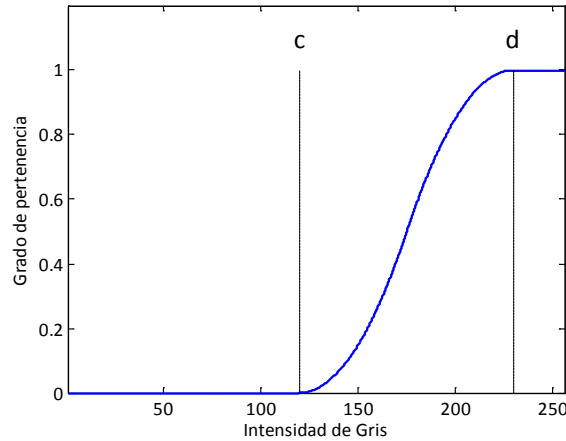
La función que define el conjunto difuso “Blanco” es de tipo “S”, también con 2 parámetros requeridos ( $c$  y  $d$ ) para su determinación. La forma de la misma puede observarse en la Figura 22. Su expresión es la siguiente:

$$\mu_{BLANCO}(x) = \begin{cases} 0 & , \quad x \leq c \\ 2 \left( d \frac{x}{d-c} \right)^2 & , \quad c \leq x \leq \frac{c+d}{2} \\ 1 - 2 \left( \frac{d-x}{d-c} \right)^2 & , \quad \frac{c+d}{2} \leq x \leq d \\ 1 & , \quad x \geq d \end{cases}.$$



**Figura 21: Función de pertenencia tipo “Z” y sus parámetros.**

Esta función queda determinada a través de dos parámetros,  $a$  y  $b$ , que definen el comienzo y el fin de la zona de transición entre 1 y 0.



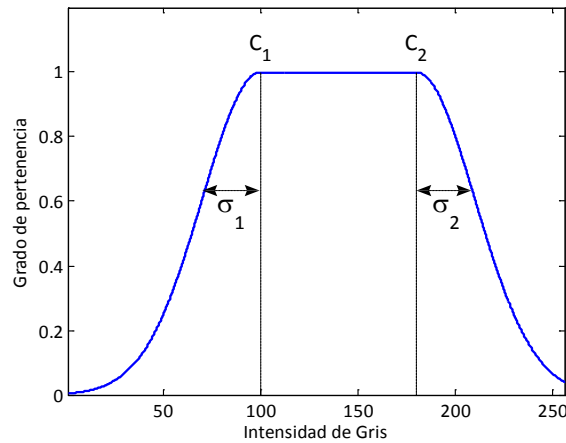
**Figura 22: Función de pertenencia tipo “S” y sus parámetros.**

Esta función queda determinada a través de dos parámetros,  $c$  y  $d$ , que definen el comienzo y el fin de la zona de transición entre 0 y 1.

Las otras funciones de pertenencia son de tipo gaussianas asimétricas. Éstas permiten que en un determinado rango de grises, los valores de pertenencia tomen el valor máximo y luego desciendan en forma suave. Estas funciones requieren 4 parámetros cada una para quedar determinadas, e indican los puntos de comienzo de descenso y la velocidad del mismo (2 para el lado izquierdo  $-C_1$  y  $\sigma_1-$  y 2 para el lado derecho  $-C_2$  y  $\sigma_2-$ ), lo que puede apreciarse en la Figura 23. La expresión es la siguiente:

$$\mu(x) = \begin{cases} \exp\left[-\frac{(x - C_1)^2}{2\sigma_1^2}\right] & , \quad x < C_1 \\ 1 & , \quad C_1 \leq x \leq C_2 \\ \exp\left[-\frac{(x - C_2)^2}{2\sigma_2^2}\right] & , \quad x > C_2 \end{cases}$$

La determinación de las ubicaciones de las funciones de pertenencia mediante la determinación de sus parámetros es, en esta etapa preliminar, basada en el análisis de los histogramas. Posteriormente se dará un criterio para su optimización.



**Figura 23: Función de pertenencia tipo gaussiana asimétrica y sus parámetros.**

Esta función queda determinada a través de cuatro parámetros. Los valores de  $c_1$  y  $c_2$ , definen la región que tomará el valor 1. Los otros dos parámetros definen la forma de la campana de gauss para las transiciones entre 0 y 1 en los dos laterales.

Como se mencionó, el procesamiento de la imagen consiste en un análisis de la misma píxel a píxel donde para cada uno se toma la decisión de la asignación a alguna de estas sustancias. Para esto se cuantifican los valores de verdad de cada uno de los predicados principales definidos ( $P_1, P_2, P_3, P_4$ ), escogiendo el que presenta el mayor de ellos. Esta cuantificación se hará a través de los conjuntos difusos definidos.

Dado un píxel con una intensidad  $x$ , la primera operación a efectuar es la obtención de un valor difuso entre 0 y 1 para cada uno de los conjuntos difusos que se aplicarán (“Gris”, “Intenso”, “Negro”, “Blanco”, etc.) mediante las ecuaciones anteriores. Luego se debe calcular el valor también difuso de las operaciones lógicas realizadas entre los predicados básicos que forman el predicado compuesto que describe a cada sustancia.

Por ejemplo, el cálculo del valor de verdad de la expresión:

$$P_3 = \text{"El píxel corresponde a MB"}$$

requiere, conocer el valor de verdad de las expresiones:

$$P_{31} = \text{"El píxel es Intenso en T1"}$$

$$P_{32} = \text{"El píxel es Gris en T2"}$$

$$P_{33} = \text{"El píxel es Intenso en PD"}$$

para efectuar la operación lógica:

$$v(P_3) = v(P_{31}) \wedge v(P_{32}) \wedge v(P_{33}) = \mu_{INTENSO}(x_{T1}) \wedge \mu_{GRIS}(x_{T2}) \wedge \mu_{INTENSO}(x_{PD}),$$

donde se ha definido:

$x_{T1}$  = intensidad de gris del píxel en análisis en la imagen T1.

$x_{T2}$  = intensidad de gris del píxel en análisis en la imagen T2.

$x_{PD}$  = intensidad de gris del píxel en análisis en la imagen PD.

Como se emplea el operador conjunción correspondiente a la LDC, definido en la ecuación (7), el valor de verdad se calcula mediante la siguiente expresión:

$$v(P_3) = [\mu_{INTENSO}(x_{T1}) \times \mu_{GRIS}(x_{T2}) \times \mu_{INTENSO}(x_{PD})]^{\frac{1}{3}}.$$

El procedimiento de cálculo se realiza para cada uno de los predicados correspondientes a las distintas sustancias. Se utilizarán siempre las operaciones de la LDC. Finalmente el píxel se asigna a la sustancia con mayor valor de verdad.

Si un píxel correspondiera en realidad a una sustancia no descrita en los predicados anteriores, igualmente sería clasificado o reconocido como de una de ellas, pero en estos casos puede ser factible determinar que la clasificación no es demasiado confiable, dado que los valores de verdad hallados para todos los predicados serán pequeños, menores a un cierto umbral de tolerancia determinado. Debe tenerse presente que se está haciendo una asignación de las sustancias “explicadas” en los predicados y que no se busca encontrar nuevo conocimiento o agrupar los píxeles en diferentes clases como suele hacerse con la aplicación de otras técnicas.

## 4.4. Implementación y formalización del método

Se presenta en esta sección la forma de operar con imágenes en forma matricial para reducir la cantidad de operaciones al mínimo.

- Primeramente se construyen vectores con los valores de intensidades de grises de las imágenes a analizar. Cada elemento de estos vectores representa un píxel de la imagen. Entonces, por ejemplo, si se trata de imágenes de 256 x 256 píxeles, se lograrían vectores de 65536 elementos. La cantidad de vectores depende de la cantidad de imágenes intervienen en las sentencias utilizándose imágenes pesadas en T<sub>1</sub>, T<sub>2</sub> y PD se requerirán 3 vectores.

- Se calculan los valores de pertenencia de todos los píxeles a los diferentes conjuntos difusos que intervienen en los predicados, cálculo que se hace en forma vectorial y por tanto con un bajo costo computacional.
- Se efectúan las operaciones lógicas que indican los predicados para cada una de las sustancias. Se obtendrán así cuatro vectores (o tantos como sustancias se desea reconocer) con los valores de verdad obtenidos para cada uno de estos predicados.
- Se analizan estos vectores elemento a elemento para determinar cuál es el mayor de ellos, tomando así la decisión de la sustancia que se asignará al píxel. Así se arma un nuevo vector con un valor diferente que identifique a cada sustancia.
- Se reconstruye una imagen final que muestre diferentes colores según la sustancia indicada. Esta imagen es la que se mostrará a los expertos para la evaluación de los resultados.

Un post-procesamiento podría considerar aisladamente los píxeles que tienen como resultado del grado de verdad de todos los predicados valores inferiores a 0.5. En este caso podría tratarse de sustancias no indicadas en los predicados y estos píxeles deberían rotularse como tales.

A continuación se presenta una formalización del método:

Sea  $X$  un vector de dimensión  $N$  cuyos elementos son las intensidades de gris de la imagen  $Y$ :

$$X_Y = [x_1 \ x_2 \ \dots \ x_N]^T ,$$

donde  $N$  representa la cantidad de píxeles a analizar. Por ejemplo, si se trabaja con imágenes de 256 x 256 píxeles,  $N = 65536$ .

Sean tres vectores:  $X_{T1}$ ,  $X_{T2}$  y  $X_{PD}$ , definidos para representar las imágenes T1, T2 y PD. La pertenencia a cada conjunto difuso se calcula como:

$$\mu_{fuzzySet}(X_Y) = mf_{fuzzySet}(X_Y) ,$$

donde  $mf_{fuzzySet}$  es la expresión de la función de pertenencia que define el conjunto difuso “fuzzySet” y  $X_Y$  es el vector de cualquiera de las imágenes consideradas. Se

calculan tantos vectores  $\mu_{fuzzySet}$  (también de dimensión  $N$ ) como conjuntos difusos se utilicen en los predicados.

Los valores de verdad de los  $K$  predicados se calculan operando con las pertenencias a los conjuntos difusos recién calculadas. Las operaciones serán diferentes, dependiendo de las definiciones de los predicados.

**Ejemplo #1:** el predicado P1 declara que “El Tejido TEJ1 es *Intenso* en T1, es *Oscuro* en T2 y es *Muy Intenso* en PD”. Usando LDC, se calcula la siguiente operación de conjunción (c):

$$\begin{aligned}\mu_{P1} &= c(\mu_{INTENSO}(X_{T1}), \mu_{OSCURO}(X_{T2}), \mu_{MUY-INTENSO}(X_{PD})) \\ &= [\mu_{INTENSO}(X_{T1}) \times \mu_{OSCURO}(X_{T2}) \times \mu_{MUY-INTENSO}(X_{PD})]^{\frac{1}{3}}.\end{aligned}$$

**Ejemplo #2:** el predicado P2 declara que “El Tejido TEJ2 es *Gris* en T1, es *Blanco* o *Muy Intenso* en T2 y es *Intenso* en PD”. En este caso se calcula la siguiente operación de conjunción (c) y disyunción (d):

$$\begin{aligned}\mu_{P2} &= c[\mu_{GRIS}(X_{T1}), d(\mu_{BLANCO}(X_{T2}), \mu_{MUY-INTENSO}(X_{T2})), \mu_{INTENSO}(X_{PD})] \\ &= \left[ \mu_{GRIS}(X_{T1}) \right. \\ &\quad \times \left( 1 - \sqrt{(1 - \mu_{BLANCO}(X_{T2}))(1 - \mu_{MUY-INTENSO}(X_{T2}))} \right) \\ &\quad \left. \times \mu_{INTENSO}(X_{PD}) \right]^{\frac{1}{3}}.\end{aligned}$$

Luego se construye una nueva matriz  $A$ , de dimensión  $N \times K$ , con esos vectores. Cada columna será un vector de grados de verdad, correspondiendo a los  $K$  predicados.

$$A = [\mu_{P1} \ \mu_{P2} \ \dots \ \mu_{PK}].$$

El objetivo de la matriz  $A$  es hallar la columna que contiene el máximo valor, operando fila por fila. El índice de la columna que contiene el máximo identifica al tejido que será asignado. Los resultados se almacenan en un nuevo vector  $D$ , también de  $N$  elementos.

$$D = indiceColMax(A),$$

donde la operación indicada como *indiceColMax* indica hallar el índice de la columna que contiene el valor máximo de cada fila.

Finalmente, el vector *D* se reacomoda para construir una matriz del mismo tamaño de la imagen original que contendrá la clasificación obtenida. Los valores de esta matriz son los índices que identifican los tejidos, siendo éstos enteros entre 1 y *K*.

Todo el proceso fue implementado en MatLab® 2006b. Los operadores de LDC fueron programados para esta aplicación.

## 4.5. Optimización con Algoritmos Genéticos

---

Las imágenes segmentadas con el método explicado en las secciones anteriores son satisfactorias y presentan gran similitud con las imágenes de referencia, resultado que también puede ser cuantificado como se explicará más adelante.

Sin embargo, las funciones de pertenencia que definen a los conjuntos difusos que permiten cuantificar los valores de verdad de los predicados son diseñadas en esta etapa de forma heurística, por lo tanto subjetiva, basándose únicamente en la inspección visual de los histogramas. Esto hace que el sistema es aplicable una vez definido, pero su utilización en otro conjunto de imágenes, por ejemplo provenientes de otro equipo, requiere un nuevo análisis de los parámetros de las funciones de pertenencia.

El sistema puede ser críticamente mejorado en su exactitud y en su adaptabilidad si las funciones de pertenencia se optimizan automáticamente para reflejar numéricamente la información contenida en las imágenes previamente segmentadas (las mismas que se utilizaron para construir los histogramas). En la siguiente sección se presenta un esquema de trabajo con ésta y otras opciones para la optimización.

### 4.5.1. Sistemas híbridos

Los sistemas de Inteligencia Computacional denominados “híbridos” son aquéllos que combinan dos o más técnicas, ofreciendo un mejor desempeño que las técnicas aisladas. En los últimos años se ha desarrollado el uso de diversas combinaciones entre sistemas que utilizan RN, LD, Algoritmos Evolutivos (particularmente AG) y aprendizaje de máquina (*machine learning*) [Dounias and Linkens, 2004].

Una de las propuestas más utilizadas consiste en la combinación de LD con AG, lo que logra un sistema Difuso-Genético (*Genetic Fuzzy System*) [Cordon, 2004, Cordon et al., 2004]. Este tipo de sistemas ha sido utilizado desde hace años en el contexto de la clasificación [Yuan and Zhuang, 1996] y es tema de constante investigación, principalmente en el área de optimización de sistemas de control [Bonissone et al., 2006, Pal and Pal, 2003]. La propuesta presentada en esta tesis consiste básicamente en un sistema difuso optimizado por un proceso de aprendizaje basado en un AG [Herrera, 2005].

Existen diversas opciones en el diseño de un sistemas híbridos difuso-genético [Alcalá et al., 2006]. Aplicadas en el caso del modelo presentado en esta tesis, las posibilidades podrían ser las siguientes:

- **Optimización de las funciones de pertenencia:** el algoritmo busca un adecuado conjunto de parámetros para determinar las características morfológicas de las funciones de pertenencia, a partir de la definición inicial, sin modificar su forma básica ni variar los predicados que intervienen en el sistema.

Ejemplo: si las funciones son gaussianas, se buscan los mejores valores para sus centros y dispersiones.

- **Selección de predicados:** si se da más de un predicado para cada tejido, el algoritmo puede seleccionar el mejor de ellos. Entonces, partiendo de un conjunto de predicados, queda el mejor subconjunto de ellos.

Ejemplo: seleccionar cuál de los siguientes predicados explica mejor las características en la MB.

$P_{3(\text{opción 1})} = \text{"El píxel corresponde a MB cuando es Intenso en T1, Gris en T2 e Intenso en PD"}$

$P_{3(\text{opción 2})} = \text{"El píxel corresponde a MB cuando es Intenso o Muy Intenso en T1, Oscuro en T2 e Intenso en PD"}$

- **Optimización de los predicados:** a partir de definiciones iniciales de los mismos, el algoritmo busca modificarlos para obtener un mejor resultado final en la segmentación. Este enfoque requiere estructuras definidas de los predicados que no podrán variarse, lo que quita cierta flexibilidad al sistema, pero brinda un mayor grado de automatismo sin perder el aprovechamiento que se hace del conocimiento inicial provisto por los especialistas.

Ejemplo: los predicados tienen la estructura fija:

$P_i = \text{"El píxel corresponde a } TEJIDO_i \text{ cuando es } mf1 \text{ en T1, } mf2 \text{ en T2 y } mf3 \text{ en PD"}$

No se utilizarían en esta propuesta otros operadores entre los predicados simples que forman el predicado compuesto.

El algoritmo debe determinar cuáles son los mejores conjuntos difusos representados por  $mf1$ ,  $mf2$  y  $mf3$ . Debe ampliar su espacio de búsqueda para hallar, además de los parámetros de las funciones de pertenencia, los mejores conjuntos difusos. Si se asigna un número a cada posible función de pertenencia:

1 = "Negro"  
2 = "Oscuro"  
3 = "Gris"  
4 = "Intenso"  
5 = "Muy-intenso"  
6 = "Blanco"

entonces puede formarse una secuencia de los conjuntos difusos que se utilizan en los sucesivos predicados y esta secuencia puede ser variada por el AG y por lo tanto deben ser incluidas en el individuo.

- **Incorporación de modificadores en los predicados:** en este enfoque se altera el efecto de las funciones de pertenencia, asociando estas alteraciones a etiquetas lingüísticas (*hedges*).

Ejemplo: los predicados tienen también una estructura fija:

$P_3$  = "El píxel corresponde a MB cuando es *hedge1* Intenso en T1, *hedge2* Gris en T2 y *hedge3* Intenso en PD"

El algoritmo debe optimizar cuáles serían *hedge1*, *hedge1* y *hedge3* eligiendo modificadores entre un conjunto de ellos (por ejemplo "marcadamente", "algo" o sin modificador). Así el mismo predicado queda expresado lingüísticamente de la siguiente manera:

$P_3$  = "El píxel corresponde a MB cuando es *algo* Intenso en T1, Gris en T2 y *marcadamente* Intenso en PD"

Los modificadores podrían ser codificados de la siguiente manera:

- 1 = "Algo"
- 2 = Ningún modificador
- 3 = "Muy"

Estos enfoques no son exhaustivos, pero constituyen los fundamentales para este sistema con predicados.

Los predicados así obtenidos pueden ser posteriormente analizados por los especialistas para comprobar si tienen una real validez lingüística o si son coherentes con su conocimiento.

Los posibles esquemas de optimización podrían abordarse de a uno (por ejemplo, primero modificar los predicados, luego optimizar las funciones de pertenencia y luego agregar modificadores) o más de uno simultáneamente (por ejemplo, buscar al mismo tiempo las mejores definiciones para los predicados y un conjunto de parámetros de las funciones de pertenencia adecuado).

Estas propuestas han sido implementadas en su totalidad y pueden verse en la interfaz gráfica de desarrollo que se presentó en la Figura 59. En la región inferior izquierda puede verse la codificación de un posible conjunto de predicados y su expresión lingüística. En esta interfaz, si los valores de los predicados codificados se

modifican por el usuario, cambia automáticamente su expresión lingüística y el siguiente procesamiento se efectuaría con esta nueva definición. Esto brinda un adecuado método para probar diversas configuraciones del sistema alternativas a las entregadas por el AG.

#### 4.5.2. Sistema de optimización propuesto

El problema de optimización elegido para esta tesis consiste en la búsqueda de un adecuado conjunto de parámetros de las funciones de pertenencia que intervienen en el método descrito. Estos parámetros definen ciertos aspectos de la forma y ubicación de las funciones de pertenencia dentro del conjunto soporte. Sin embargo la morfología general de las funciones se elige inicialmente para no variar durante la optimización. Las funciones tipo “Z”, tipo “S” y gaussianas seguirán siéndolo, pero el AG será capaz de ir variando sus centros, anchos y velocidades de caída.

Se ha definido una función de evaluación basada en la cuantificación de verdad de una expresión difusa, lo que hace el problema complejo o directamente intratable para otros métodos de optimización. Si bien la función es sumamente compleja para operar matemáticamente, esto no es una limitación para los AG ni para su implementación en MatLab®.

Se aprovechó la LDC como una herramienta de Ingeniería del Conocimiento que, por su mejor comportamiento axiomático, permite “atrapar” fácilmente el conocimiento a partir de su expresión verbal.

El individuo para el AG es un vector conteniendo los parámetros indicados. Los conjuntos difusos  $\mu_{NEGRO}(x)$  y  $\mu_{BLANCO}(x)$  requieren dos parámetros cada uno para identificar la ubicación y pendiente de caída. Los conjuntos difusos  $\mu_{OSCURO}(x)$ ,  $\mu_{GRIS}(x)$ ,  $\mu_{INTENSO}(x)$  y  $\mu_{MUY-INTENSO}(x)$  requieren cuatro parámetros cada uno de ellos (los centros y dispersiones de las dos funciones gaussianas que intervienen en cada uno). Así, se han ordenado los parámetros de las distintas funciones para definir el individuo:

$$INDIVIDUO = [\text{Par}(\mu_{NEGRO}) \text{Par}(\mu_{OSCURO}) \text{Par}(\mu_{GRIS}) \dots \\ \dots \text{Par}(\mu_{INTENSO}) \text{Par}(\mu_{MUY-INTENSO}) \text{Par}(\mu_{BLANCO})],$$

donde  $\text{Par}(\mu_A)$  indica la secuencia de parámetros de la función de pertenencia correspondiente al conjunto difuso  $A$ .

Según los parámetros requeridos para cada función de pertenencia, este vector será, finalmente:

$INDIVIDUO =$

$[a_{NEGRO} \ b_{NEGRO} \ \sigma_{1OSCURO} \ C_{1OSCURO} \ \sigma_{2OSCURO} \ C_{2OSCURO} \ \sigma_{1GRIS} \ C_{1GRIS} \ \sigma_{2GRIS} \ C_{2GRIS} \ \dots \ \sigma_{1INTENSO} \ C_{1INTENSO} \ \sigma_{2INTENSO} \ C_{2INTENSO} \ \sigma_{1MUY} \ C_{1MUY} \ \sigma_{2MUY} \ C_{2MUY} \ a_{BLANCO} \ b_{BLANCO}]$ .

Para crear una población inicial adecuada de la cual comenzar a ejecutar el AG se creó un individuo con los parámetros obtenidos a partir de los histogramas. El resto de los individuos iniciales se creó permitiendo variar aleatoriamente cada parámetro en un 40% en más o en menos.

Como es conocido, un AG busca iterativamente los parámetros tratando de minimizar una función de evaluación. La definición de la función de evaluación es un paso crítico, si no el más importante, para obtener un resultado adecuado.

En esta tesis el valor de la función de evaluación se define a partir de un conjunto de píxeles correctamente clasificados *a priori* con igual cantidad de ellos para cada sustancia a detectar. Se hicieron pruebas tomando de 50 a 100 píxeles que actúan a modo de un conjunto de datos de entrenamiento del algoritmo.

Como se explicó en párrafos anteriores, se recurrió nuevamente a la lógica de predicados difusos y su evaluación con LDC para definir la función de evaluación. Para cada píxel de entrenamiento, el objetivo es maximizar el grado de verdad obtenido para el predicado que le corresponde a la sustancia real a la que el píxel pertenece y minimizar los valores de verdad para los predicados que definen todas las sustancias restantes.

Por ejemplo, si un píxel de entrenamiento corresponde a MB, el caso ideal sería obtener 1 para el valor de verdad del predicado que describe a la MB ( $P_3$  en las definiciones dadas en la sección correspondiente) y 0 para el valor de verdad de todos los predicados restantes.

Entonces, para un único píxel, debe maximizarse el valor de verdad de la siguiente doble implicación, que se interpreta como “Si el píxel corresponde a MB entonces es detectado como MB y viceversa.”:

$$\begin{aligned} MB(pixel) &\Leftrightarrow MBreal(pixel) = \\ &= [MB(pixel) \Rightarrow MBreal(pixel)] \wedge [MBreal(pixel) \Rightarrow MB(pixel)] , \end{aligned}$$

donde:

$$MB(pixel) = \text{“El píxel es detectado como materia blanca”} = v(P_3)$$

$$MBreal(pixel) = \text{“El píxel corresponde a materia blanca”} = \begin{cases} 1 & \text{si es MB} \\ 0 & \text{si no es MB} \end{cases} .$$

Esta expresión lógica se repite para cada sustancia a detectar y todas las implicaciones dobles deben cumplirse al mismo tiempo, por lo que estarán conectadas con el conectivo AND, teniendo para cada píxel la expresión:

$$\begin{aligned} R(pixel) &= [LCR(pixel) \Leftrightarrow LCRreal(pixel)] \wedge [MB(pixel) \Leftrightarrow MBreal(pixel)] \wedge \\ &\quad \wedge [MG(pixel) \Leftrightarrow MGreal(pixel)] \\ &\quad \wedge [FONDO(pixel) \Leftrightarrow FONDOreal(pixel)] . \end{aligned}$$

Esta expresión debe cumplirse para todos los píxeles del conjunto de datos de entrenamiento. Para expresar esto, se consideró la proposición universal sobre todos los píxeles (la conjunción).

$$v(F) = \bigvee_{pixel \in T} v[R(pixel)] = \bigwedge_{pixel \in T} v[R(pixel)] ,$$

donde  $T$  es el conjunto de píxeles de entrenamiento.  $v(F)$  es el valor que se procura maximizar.

Para esta aplicación particular, pueden presentarse sólo dos posibles casos de dobles implicaciones. Pueden ser de tipo  $P \Leftrightarrow 0$  ó  $P \Leftrightarrow 1$ , dado que los píxeles pertenecen o no a un tejido específico.

Se eligió la implicación de Zadeh generalizada ( $P \Rightarrow Q = \neg P \vee (P \wedge Q)$ ) y se continuó con la aplicación de LDC. Así, las expresiones a calcular quedaron de la forma:

$$\begin{aligned}
P \Leftrightarrow 0 &= (P \Rightarrow 0) \wedge (0 \Rightarrow P) \\
&= [-P \vee (P \wedge 0)] \wedge [1 \vee (0 \wedge P)] \\
&= (-P \wedge 0) \wedge (1 \vee 0),
\end{aligned}$$

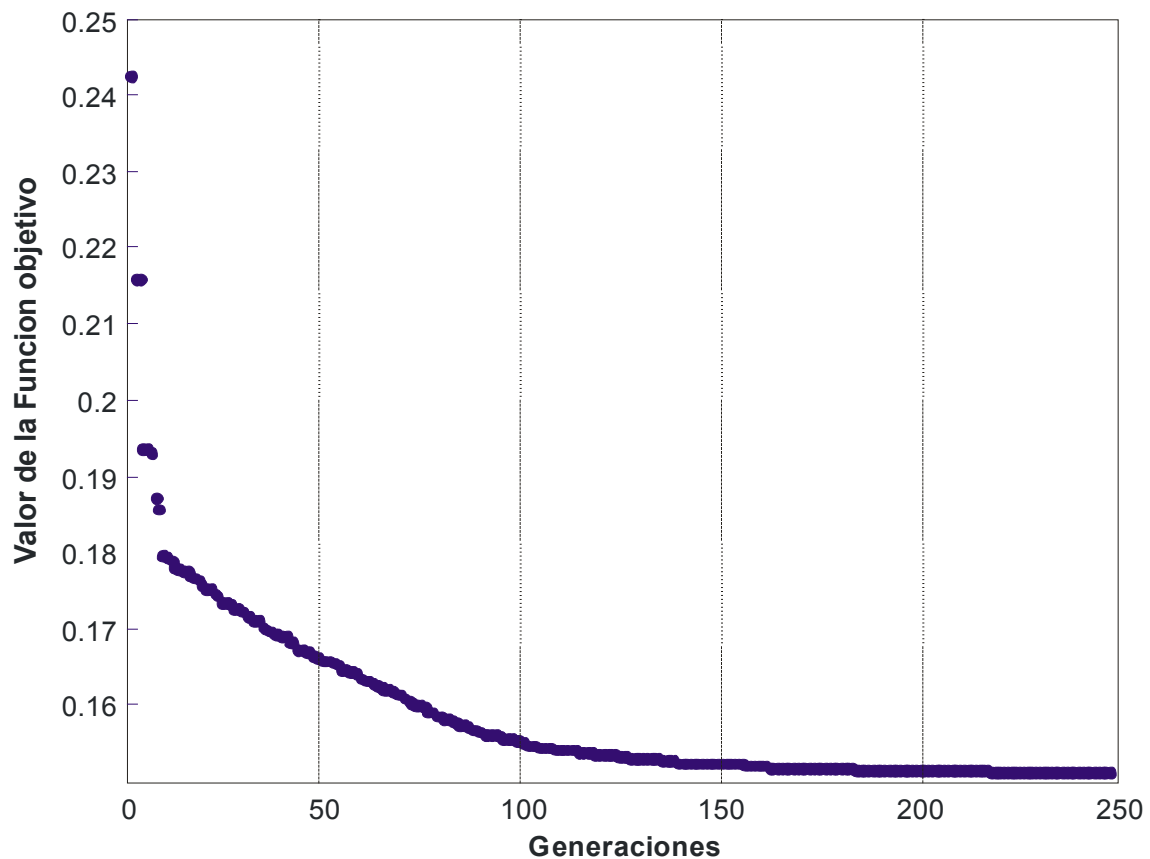
$$v(P \Leftrightarrow 0) = \sqrt{1 - \sqrt{v(P)}},$$

$$\begin{aligned}
P \Leftrightarrow 1 &= (P \Rightarrow 1) \wedge (1 \Rightarrow P) \\
&= [-P \vee (P \wedge 1)] \wedge [0 \vee (1 \wedge P)],
\end{aligned}$$

$$v(P \Leftrightarrow 1) = \sqrt{\left[1 - \sqrt{v(P) \left(1 - \sqrt{v(P)}\right)}\right] \left[1 - \sqrt{\left(1 - \sqrt{v(P)}\right)}\right]}.$$

El AG se ejecuta una única vez antes del procesamiento de un conjunto de imágenes o de una imagen 3D compuesta por una serie de cortes. Una vez que las funciones de pertenencia han sido optimizadas, el algoritmo queda preparado para procesar imágenes similares. Generalmente se tomará un único corte de un estudio 3D completo para tener algunos píxeles reconocidos *a priori* para cada tejido (conjunto de píxeles de entrenamiento) y luego se procesará el resto de los cortes con las funciones de pertenencia optimizadas.

El proceso iterativo del AG puede verse gráficamente representando el valor de la función de evaluación para el mejor individuo de las sucesivas poblaciones, y el valor medio de la población. Un gráfico típico de este proceso en MatLab® puede observarse en la Figura 24. En la figura puede verse que el valor de la función de evaluación que produce el mejor individuo (el mejor conjunto de parámetros) de cada población va sucesivamente disminuyendo. El valor medio de la población no es aplicable en esta aplicación, dado que a los individuos que no cumplen las restricciones que se darán más adelante, se les asigna valor infinito para su función de evaluación. Por ejemplo, en la figura puede observarse que el valor comienza siendo 0.228 y luego de 210 iteraciones el valor se estabiliza en 0.155.



**Figura 24: Progreso típico de entrenamiento de un Algoritmo Genético.**

El valor de la ordenada es el complemento a 1 del valor de verdad del predicado, tomado como función objetivo. Se observa que el valor que produce el mejor individuo disminuye en las sucesivas generaciones.

La configuración del AG fue elegida heurísticamente, pues no hay una forma sistemática de elegir las opciones que pueden variarse, que no son pocas. Luego de una gran cantidad de pruebas, se eligió la siguiente configuración:

- Fracción de crossover: 0.7. El 70% de los nuevos individuos se generarán por cruzamiento y el 30% por mutación.
- Cantidad de generaciones máxima: 250 generaciones.
- Tamaño de la población: 150 individuos.
- Elite: 2 individuos.
- Rango de la población inicial: comprendido por intervalos de los parámetros centrados en los valores obtenidos por el análisis visual de los histogramas, más / menos el 20%.

- Límite de tiempo: un valor alto de modo que no se detenga el algoritmo por tiempo de procesamiento.
- Se configuró para que el AG no se detenga aunque haya relativamente poca mejora entre una población y la siguiente.
- Mutación: uniforme, con una tasa de 0.9. Es un valor alto para permitir una búsqueda amplia en el espacio de soluciones.

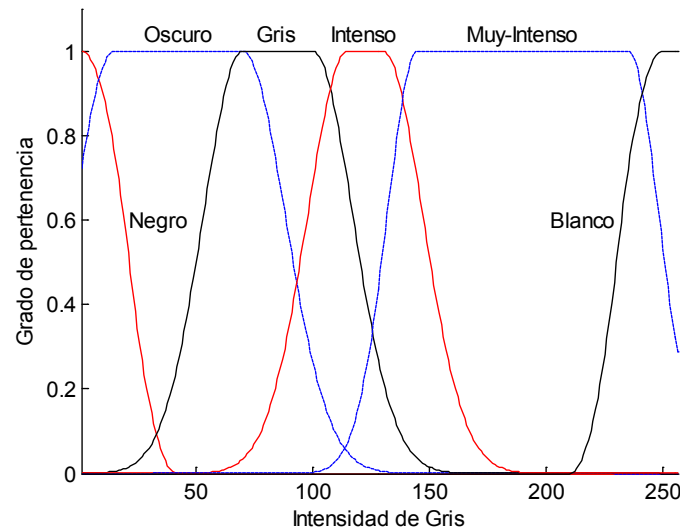
Además se incorporaron restricciones en los valores de los parámetros. Se configuró de manera que todas las funciones de pertenencia tengan sus centros en el intervalo [0, 255] y que sus pendientes no desciendan demasiado lentamente para que no se superpongan exageradamente entre ellas. A esos valores se impuso el valor máximo en 60.

Para seguir conservando la interpretabilidad del modelo empleado (por ejemplo, que la función de pertenencia correspondiente al conjunto difuso “Gris muy intenso” represente efectivamente valores de gris altos) se agregaron restricciones de modo que se mantuviera el orden en que las funciones gaussianas aparecen. Las restricciones exigen, por ejemplo, que el segundo centro de la función del conjunto “Oscuro” sea inferior al primer centro de la función del conjunto “Gris”. Formalmente, lo que se exige es:

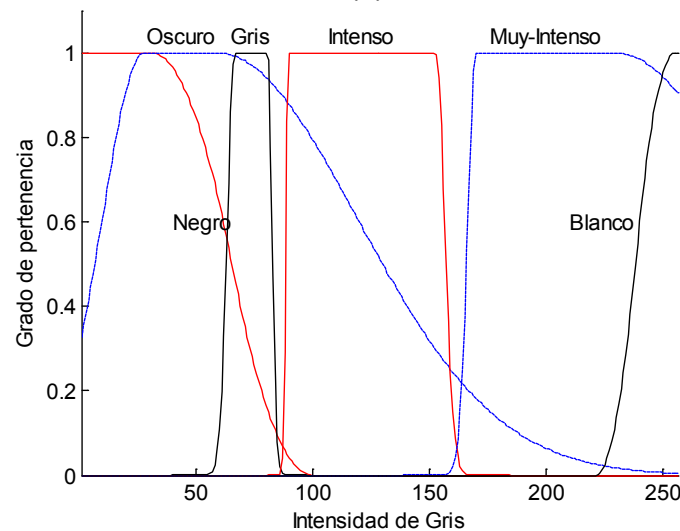
$$\begin{aligned}\frac{a_{NEGRO} + b_{NEGRO}}{2} &< C_{1OSCURO} , \\ C_{2OSCURO} &< C_{1GRIS} , \\ C_{2GRIS} &< C_{1INTENSO} , \\ C_{2INTENSO} &< C_{1MUY-INTENSO} , \\ C_{2MUY-INTENSO} &< \frac{a_{BLANCO} + b_{BLANCO}}{2} .\end{aligned}$$

Luego de ejecutar el AG, las funciones de pertenencia cambian sus posiciones y pendientes dentro del conjunto soporte tomando valores que optimizan los resultados del procesamiento.

En la Figura 25 pueden verse las formas de las funciones de pertenencia iniciales y las formas de las mismas luego de la optimización con AG para una imagen en particular.



(a)



(b)

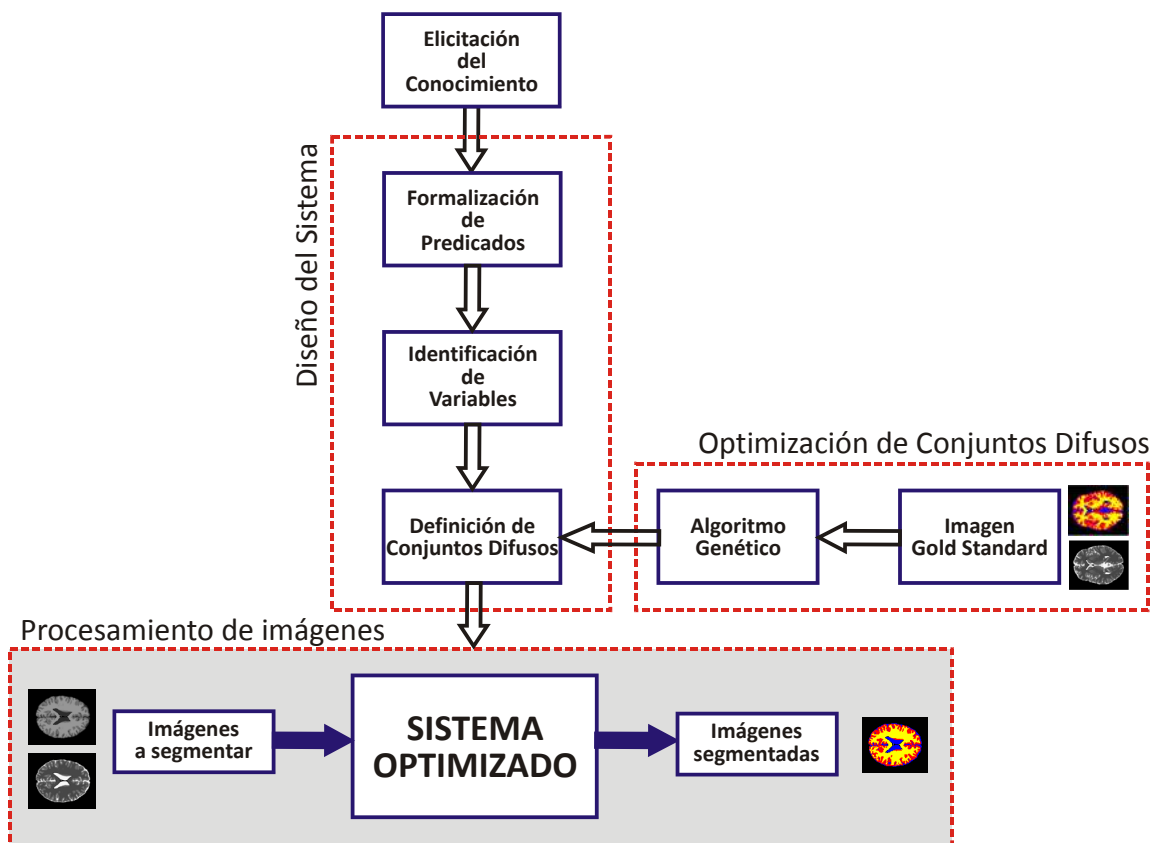
**Figura 25: Funciones de pertenencia antes y después de la optimización realizada por el Algoritmo Genético.**

Las funciones de pertenencia preliminares (a) son iterativamente modificadas en su morfología en las sucesivas generaciones del Algoritmo Genético, por medio de la variación de sus parámetros, hasta llegar a un estado final (b).

## 4.6. Esquema del método propuesto

A modo de resumen, pueden verse en la Figura 26 los pasos seguidos en la etapa de diseño y puesta a punto del sistema lógico de segmentación, en el que se explicita la optimización a través de AG.

También se muestra la secuencia de procesamiento de imágenes una vez que el sistema ha sido entrenado. Una vez efectuada la optimización, las funciones de pertenencias ya no se modifican durante la etapa de procesamiento, el que se realizaría corte a corte para un estudio 3D completo. Durante la validación del método, las imágenes involucradas deben ser distintas a la imagen utilizada para el entrenamiento del sistema.



**Figura 26: Esquema de los pasos de diseño y optimización del sistema.**

Una vez que se cuenta con el conocimiento de los expertos, se diseña y optimiza el sistema, para proceder al procesamiento de imágenes similares. Si estas imágenes provienen de otro protocolo, se debe ejecutar nuevamente la etapa de optimización.

## 4.7. Validación del método

---

Se han propuesto variadas maneras de medir cuantitativamente la calidad de la segmentación obtenida. Para el cálculo de las diferentes mediciones de calidad de la segmentación se requieren imágenes (o al menos una parte de las mismas) que puedan tomarse como referencia o Gold Standard. El problema de la evaluación de clasificadores que se utilizan en segmentación es constante tema de estudio [Barra and Lundervold, 2007, Bouix et al., 2007, Crum et al., 2006].

En la práctica, es difícil disponer de imágenes de referencia previamente segmentadas. El trazado manual de las diferentes sustancias que se desean detectar es un trabajo arduo que suele consumir varias horas a un especialista, aun si esta tarea se efectúa en una computadora. Además, los resultados pueden variar de un operador a otro, sin mencionar que la calidad de los resultados va empeorando debido al cansancio visual que produce este tipo de trabajo.

Como solución a este inconveniente se suele recurrir a imágenes provenientes de simulaciones computacionales que tienen las siguientes ventajas:

- presentan una segmentación de referencia objetiva,
- permiten comparaciones con otros autores y otros métodos de segmentación que hayan sido aplicadas a las mismas imágenes,
- ayudan a evaluar el desempeño de los diferentes métodos bajo condiciones de distorsiones conocidas.

Como desventaja, las simulaciones suelen entregar resultados mejores que los que se logran en una condición real, por lo que la valoración de un método con este tipo de imágenes como referencia no hace que pueda prescindirse de una evaluación con imágenes reales.

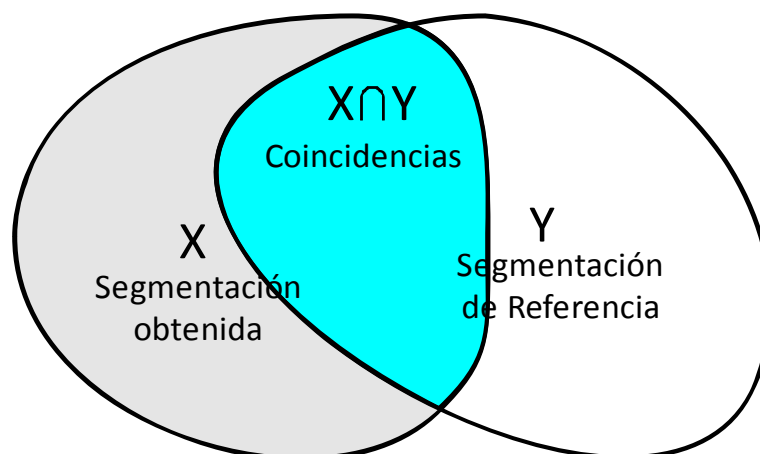
### 4.7.1. Medidas de similitud

Según lo expuesto anteriormente, para evaluar cuantitativamente los resultados obtenidos en algoritmos de segmentación se han definido diferentes

medidas de similitud. Para efectuar comparaciones de resultados obtenidos por diversos métodos, debe utilizarse siempre la misma medida.

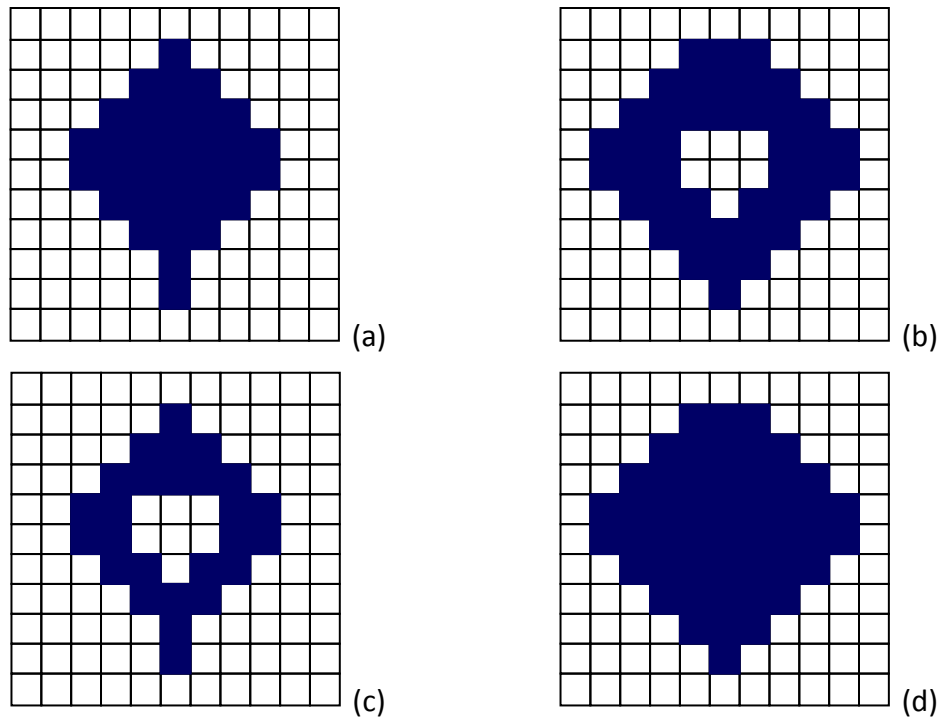
Aunque las segmentaciones halladas contienen diversas etiquetas (una por cada tejido), la segmentación de cada tejido puede representarse con una imagen binaria, donde los píxeles activos (valor 1, blanco) corresponden al tejido y los demás (valor 0, negro) al resto de la imagen, como se mostró en la Figura 62. Estas imágenes binarias deben ser comparadas con las obtenidas a partir de la segmentación de referencia.

En la Figura 27 se muestra esquemáticamente la región que debería haberse clasificado junto con la región obtenida y la región en la que ambas coinciden. Estas regiones permiten la definición de diferentes medidas de similitud, las que se dan a continuación, junto con un breve análisis de sus características.



**Figura 27: Esquema de las consideraciones necesarias para el cálculo de diferentes medidas de similitud.**

Partiendo de imágenes binarias de la segmentación obtenida y la segmentación de referencia de cada tejido, se establecen las regiones que se observan en la figura. Las cantidades de píxeles que contienen cada una se utilizan en las diferentes medidas de similitud.



**Figura 28: Ejemplo de imágenes binarias para el cálculo de diferentes medidas de calidad de una segmentación.**

(a) Imagen real X (33 píxeles activos), (b) imagen obtenida Y (42 píxeles activos), (c) intersección entre X e Y (26 píxeles activos), (d) unión de X e Y (49 píxeles activos).

La Figura 28 muestra un ejemplo de una imagen simulada de 11 x 11 píxeles. En (a) se observa la imagen real X, que presenta 33 píxeles activos; en (b) se muestra una posible imagen obtenida Y por algún método de segmentación (42 píxeles activos); en (c) se observa la imagen intersección entre X e Y (26 píxeles activos); en (d) se muestra la unión de X e Y (49 píxeles activos). Estas cantidades de píxeles activos serán utilizadas en las siguientes definiciones de medidas de calidad.

#### 4.7.1.1. Coeficiente de Tanimoto

El **Coeficiente de Tanimoto** (TC, *Tanimoto Coefficient*), también conocido como Tasa de Solapamiento Relativo (*Relative Overlap Ratio*) y también definido como coeficiente de *Jaccard*, ha sido ampliamente utilizado en aplicaciones similares a la presente [Bouix et al., 2007, Jimenez-Alaniz et al., 2006, Santalla et al., 2007, Song et al., 2004, Song et al., 2007].

La definición de este coeficiente, basándose en el esquema de la Figura 27 es [Duda et al., 2000]:

$$TC_{XY}(k) = \frac{n_{X \cap Y}(k)}{n_{X \cup Y}(k)},$$

$X$  = Imagen binaria que contiene los píxeles identificados como del tejido  $k$ ;

$Y$  = Imagen binaria de segmentación de referencia para el tejido  $k$ ;

$n_M(k)$  = Cantidad de píxeles activos de la imagen binaria genérica  $M$ ;

$TC_{XY}(k)$  = Coeficiente de Tanimoto para comparar la imagen  $X$  con la  $Y$ .

El TC se calcula para cada tejido. Cuando la segmentación es perfecta, la imagen binaria  $X$  coincide con la imagen  $Y$ , por lo que la intersección entre ellas ( $X \cap Y$ ) y la unión entre ellas ( $X \cup Y$ ) son la misma imagen resultado, y por tanto numerador y denominador son iguales: el coeficiente toma el valor 1. En el ejemplo de la Figura 28 se obtiene  $TC_{XY} = \frac{26}{49} = 0.53$ .

Cuando se detecta una cantidad mayor o menor de píxeles que los que realmente corresponden para un tejido, la intersección siempre contiene menos píxeles que la unión de la imagen resultado con la de referencia y entonces el coeficiente toma un valor menor a 1. Algunos ejemplos para diferentes situaciones simuladas pueden observarse en la Figura 29.

| Imagen X<br>Gold Standard | Imagen Y<br>Resultado | $X \cap Y$ | $X \cup Y$ | Coeficiente<br>de Tanimoto |
|---------------------------|-----------------------|------------|------------|----------------------------|
|                           |                       |            |            | 0.01                       |
|                           |                       |            |            | 0.29                       |
|                           |                       |            |            | 0.53                       |
|                           |                       |            |            | 1.00                       |

**Figura 29: Coeficientes de Tanimoto.**

El Coeficiente de Tanimoto toma un valor mínimo cuando las regiones obtenida y esperada no se solapan y valor máximo si la coincidencia entre ambas es total.

#### 4.7.1.2. Coeficiente de Exactitud

La medición de Exactitud (ACC, *Accuracy*) en la segmentación tiene en cuenta la cantidad de píxeles que han sido correctamente clasificados con respecto a la cantidad total de píxeles que se esperaba detectar. Esta medición se utiliza para la evaluación de algoritmos de clasificación y también se ha definido con el nombre de “coeficiente de coincidencia” [Sasikala et al., 2006, Bouix et al., 2007, Kang et al., 2008].

Basándose también en la Figura 27, se define de la siguiente manera:

$$ACC_{XY}(k) = \frac{n_{X \cap Y}(k)}{n_Y(k)},$$

$X$  = Imagen binaria que contiene los píxeles identificados como del tejido  $k$ ;

$Y$  = Imagen binaria de segmentación de referencia para el tejido  $k$ ;

$n_M(k)$  = Cantidad de píxeles activos de la imagen  $M$ .

$ACC_{XY}(k)$  = Coeficiente de Exactitud para comparar la imagen  $X$  con la  $Y$ ;

Esta medida puede darse para cada tejido clasificado o bien puede obtenerse de manera general, considerando la cantidad total de píxeles que fueron correctamente clasificados y la cantidad total de píxeles a clasificar.

En el ejemplo de la Figura 28 se obtiene  $ACC_{XY} = \frac{26}{33} = 0.79$ .

Los valores obtenidos para una misma segmentación resultante suelen ser más altos que los correspondientes al TC.

#### 4.7.1.3. Porcentaje de error en la clasificación

En algunos trabajos [Song et al., 2007] se dan los resultados en base a los valores del **Error de Clasificación** (*MCR*, *MisClassification Rate*), que indica la cantidad en porcentaje de píxeles incorrectamente clasificados.

Se encuentra directamente relacionado con la medición de exactitud (*ACC*), y puede calcularse de la siguiente manera:

$$MCR = (1 - ACC) \times 100\% .$$

En el ejemplo de la Figura 28 se obtiene  $MCR = (1 - 0.79) \times 100\% = 21\%$ .

#### 4.7.1.4. Coeficiente de Dice

El **Coeficiente de Dice (DICE)** se define:

$$DICE_{XY}(k) = \frac{2 n_{X \cap Y}(k)}{n_X(k) + n_Y(k)} ,$$

$X$  = Imagen binaria que contiene los píxeles identificados como del tejido  $k$ ;

$Y$  = Imagen binaria de segmentación de referencia para el tejido  $k$ ;

$n_M(k)$  = Cantidad de píxeles activos de la imagen  $M$ .

$DICE_{XY}(k)$  = Coeficiente de Dice para comparar la imagen  $X$  con la  $Y$ ;

Este coeficiente se encuentra relacionado algebraicamente con el TC según la siguiente ecuación:

$$DICE = \frac{2 TC}{1 + TC} .$$

En el ejemplo de la Figura 28 se obtiene  $DICE = \frac{2 \times 0.53}{1 + 0.53} = 0.69$ .

Considerando la misma segmentación, los valores de este coeficiente son siempre más altos, lo que puede llevar a confusión ante comparaciones.

#### 4.7.1.5. Matrices de Confusión

Las matrices de confusión se utilizan para mostrar el resultado de cualquier algoritmo de clasificación [Kohavi and Provost, 1998]. Aplicadas en la segmentación de imágenes, estas matrices permiten ver explícitamente cuántos píxeles han sido correctamente clasificados para cada tejido, pero además permiten evaluar cómo se han propagado los errores en los píxeles erróneamente clasificados. Se calculan los cocientes entre cantidad de píxeles detectados como un tejido y cantidad de píxeles reales que corresponden a cada tejido. Se espera que los mayores valores se encuentren en la diagonal principal, los cuales coinciden con la definición de Exactitud. La forma de estas matrices se muestra en la Tabla VII, donde se dan los valores de un caso de clasificación sin errores.

**Tabla VII: Matriz de confusión para un caso de clasificación sin errores.**

|                                |             | ... que han sido clasificados como... |             |             |
|--------------------------------|-------------|---------------------------------------|-------------|-------------|
|                                |             | ...tejido A                           | ...tejido B | ...tejido C |
| Píxeles correspondientes al... | ...tejido A | 1.00                                  | 0.00        | 0.00        |
|                                | ...tejido B | 0.00                                  | 1.00        | 0.00        |
|                                | ...tejido C | 0.00                                  | 0.00        | 1.00        |

Si bien las matrices de confusión dan más posibilidades de evaluar la segmentación obtenida, no resulta una manera resumida de expresar el resultado para efectuar comparaciones rápidas.

#### 4.7.2. Imágenes de prueba provenientes de simulaciones

Como ya se expuso, a medida que se van desarrollando cada vez más métodos y sistemas para el análisis cuantitativo y procesamiento de las imágenes, la dificultad para comparar tales métodos va creciendo. Desafortunadamente, no existe una “verdad absoluta” o un Gold Standard asociado a las imágenes adquiridas *in vivo* para evaluar cuantitativamente los resultados obtenidos y comparar diferentes técnicas.

Para solucionar este problema, como resultado del trabajo de científicos de la Universidad de McGill se ha construido una **base de datos de imágenes simuladas** (*Simulated Brain Database*) que contiene un conjunto de datos de RM producido por

un simulador. Estos datos se disponen libremente en la web y pueden ser utilizados para evaluar el desempeño de los métodos de análisis de imágenes, dado que se conoce la verdadera ubicación píxel a píxel de los tejidos [Montréal Neurological Institute, 2007].

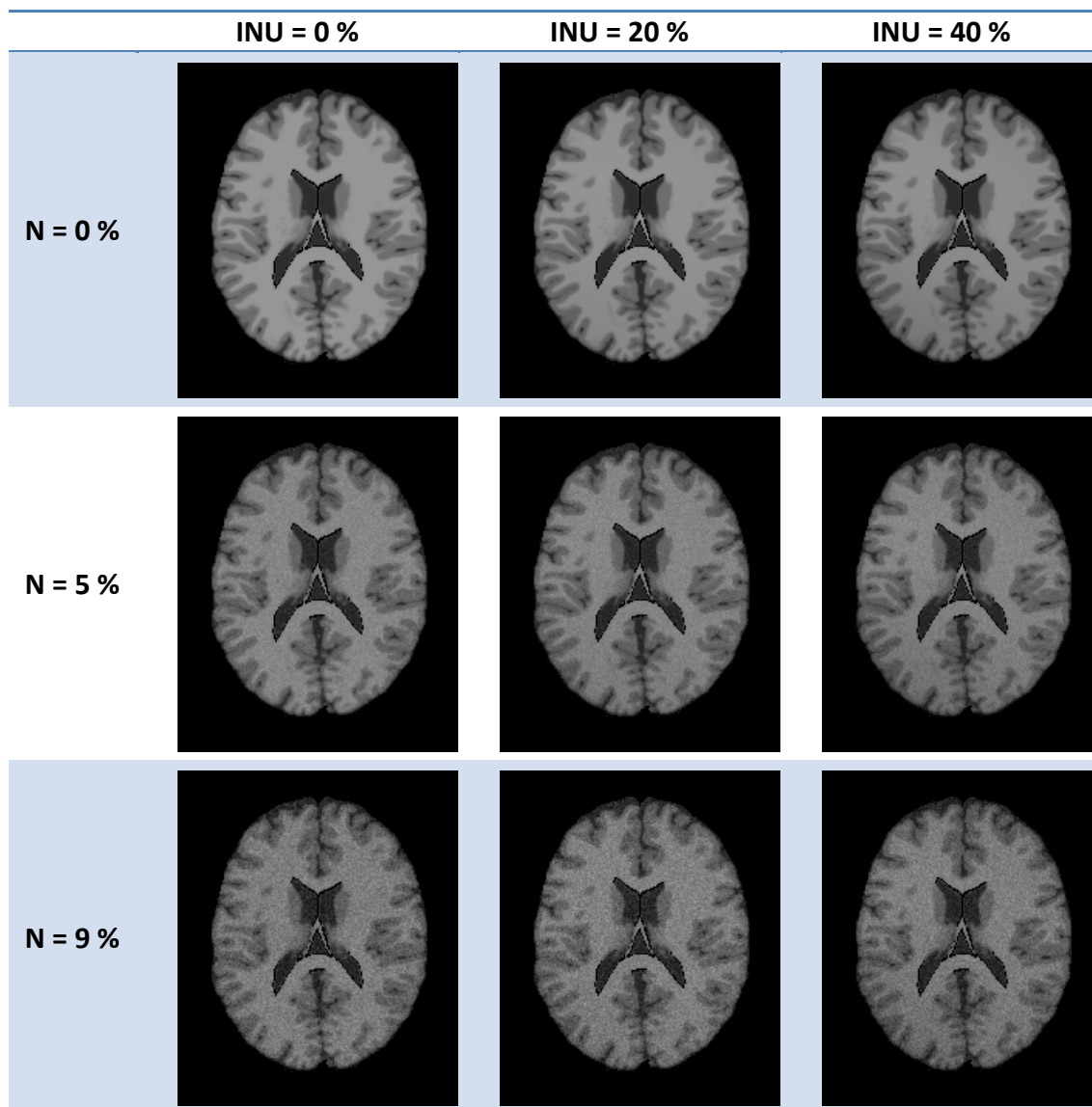
Para crear esta base de datos se obtuvieron imágenes normales y patológicas por simulación de un equipo de RM en tres secuencias: pesadas en  $T_1$ , pesadas en  $T_2$  y pesadas según la **densidad de protones** (PD, *Proton Density*). Pueden elegirse estudios simulados configurando el ancho de corte (distancia entre cortes), el nivel de **ruido** y el nivel de **no-uniformidad de la intensidad** (INU) de la imagen.

El **nivel de ruido** presente en la imagen puede elegirse variando la desviación estándar de ruido gaussiano que se suma a los canales reales e imaginario del simulador de resonador. Se indica como el porcentaje de ruido que multiplica al máximo valor de intensidad que presenta uno de los tejidos tomado como referencia (el rango entonces es de 0 a 100%). El ruido se calcula como ruido pseudoaleatorio gaussiano y se suma a la magnitud final que entrega el simulador.

La **INU** se indica en porcentaje, de modo que, por ejemplo, para un nivel de 20%, el campo de no-uniformidades multiplicativo tiene un rango entre 0.9 y 1.1 sobre el área del cerebro.

El método propuesto se evaluó primeramente con estas imágenes provenientes de esta Base de Datos del *Montréal Neurological Institute*. Esto permitió evaluar su robustez en cuanto al ruido y a la INU, utilizando imágenes con diferentes niveles de estas distorsiones.

En la Figura 30 se puede observar un corte pesado en  $T_1$  con diferentes combinaciones de ruido e INU, pudiéndose apreciar la degradación en la calidad de la imagen, lo que dificulta su segmentación, especialmente en los métodos que trabajan píxel a píxel.



**Figura 30: Imágenes pesadas en  $T_1$  con distintos niveles de distorsión.**

Se observan todas las combinaciones de diferentes niveles de ruido (N) y de no-uniformidades en la intensidad (INU). La imagen sin distorsiones se encuentra en el extremo superior izquierdo y el peor caso en el extremo inferior derecho.

Para evaluar la robustez del método presentado en este trabajo se ha segmentado un volumen completo normal para todas las combinaciones entre nivel de ruido de 0, 1, 3, 5, 7 y 9% y niveles de INU de 0, 20 y 40%.

El estudio analizado consiste en 181 cortes de 217x131 píxeles, tomados cada 1 mm, con un volumen de vóxel de 1 mm<sup>3</sup>.

#### 4.7.3. Imágenes de prueba reales

Para evaluar el método con imágenes reales, se utilizaron estudios obtenidos en dos instituciones diferentes. Se obtuvieron imágenes de diferentes calidades y

preprocesamientos, debido a que los equipos fueron de diferente marca y modelo, lo que contribuyó a evaluar también la adaptación del método propuesto a distintos datos.

Las imágenes se pre-procesaron para remover las regiones de tejidos indeseados, tales como cráneo y meninges. Se aplicó un procedimiento previamente desarrollado en el Laboratorio de Procesos y Medición de Señales (Facultad de Ingeniería, Universidad Nacional de Mar del Plata), basado en Morfología Matemática y distancia geodésica [Pastore et al., 2005].

Con el fin de evaluar cuantitativamente los resultados obtenidos, se siguieron los siguientes pasos:

1. se obtuvieron segmentaciones preliminares mediante el software BRAINS [Iowa Mental Health Clinical Research Center, 2005] y el algoritmo Fuzzy-C-Means, que es no supervisado (no requiere ejemplos de entrenamiento).
2. Se solicitó a especialistas que corrigieran manualmente los errores de segmentación que percibían, en una determinada cantidad de cortes que se detallará en cada caso. Éste fue un proceso lento y complejo.
3. Con esos cortes se calcularon las medidas de calidad de la segmentación, promediando los valores obtenidos, de la misma manera que se encontró en trabajos similares [Song et al., 2007].

Ninguno de los especialistas que modificaron las imágenes segmentadas estuvo involucrado en la construcción de los predicados con que se segmentaron las imágenes mediante el método propuesto.

No se efectuaron procesamientos previos para disminuir el efecto de ruido y de la INU de las imágenes.

#### **4.7.3.1. Conjunto de imágenes de prueba #1**

Se obtuvieron 5 estudios 3D de RM realizados en la Clínica de Demencia del Instituto de Investigaciones Neurológicas “Raúl Carrea” (Buenos Aires, Argentina). Se adquirieron en un equipo de 1.5 Teslas. El protocolo incluyó cortes coronales 3D pesados en T1 con ecos ortogonales a la línea AC-PC (TR/TE= 24/5 ms, ancho de corte =

1.5 mm); y cortes coronales pesados en PD y T2 con ecos de spin rápidos (TR/TE! /TE= 3500/32/96 ms, *echo train length* = 8, ancho de corte = 3 mm). Los estudios se adquirieron en 192 cortes de 256x256 píxeles.

Presentaron gran dificultad para su segmentación. Dos profesionales especialistas en imágenes tomaron 18 cortes representativos de cada estudio. Consensuando entre ellos, modificaron las segmentaciones preliminares, obteniendo así las que fueron utilizadas de referencia.

#### **4.7.3.2. Conjunto de imágenes de prueba #2**

Se obtuvieron en este caso 5 estudios 3D de RM realizados en el Instituto Radiológico (Mar del Plata, Argentina). Los estudios de RM se realizaron en un equipo de 0.5 Tesla Philips Gyroscan NT T5 (*Philips Medical System*). Se utilizó una bobina de cuadratura y se obtuvieron imágenes volumétricas en un plano axial ponderadas en T1 (TR/TE = 632.7/15 ms, DFOV = 24.9 x 24.9 cm) y T2 (TR/TE = 3630.3/90 ms, DFOV = 24.9 x 24.9 cm). Los estudios se adquirieron en 20 cortes de 256x256 píxeles.

Tres especialistas tomaron imágenes representativas, en este caso 15, y optimizaron manualmente las segmentaciones preliminares según su criterio.

---

## 5. Resultados

## 5. Resultados

---

En las siguientes secciones se muestran los resultados obtenidos en imágenes simuladas y en imágenes reales.

Se procesaron imágenes con distintos niveles de distorsiones. Se muestran también los resultados que se obtienen cuando se cambia la operación del sistema por la lógica MAX-MIN en vez de la LDC.

A continuación se efectúa una comparación del desempeño del método presentado con resultados presentados en la bibliografía, cuando se utilizaron las mismas imágenes simuladas. Luego se efectúa una comparación detallada de los resultados obtenidos con el método propuesto y las siguientes técnicas: software *Brains*, algoritmo Fuzzy-C-Means sin modificaciones, algoritmo de las K-medias, clasificador de los k vecinos más próximos, una red neuronal multicapa alimentada hacia delante (*feed forward*) y una red neuronal probabilística.

### 5.1. Con imágenes simuladas

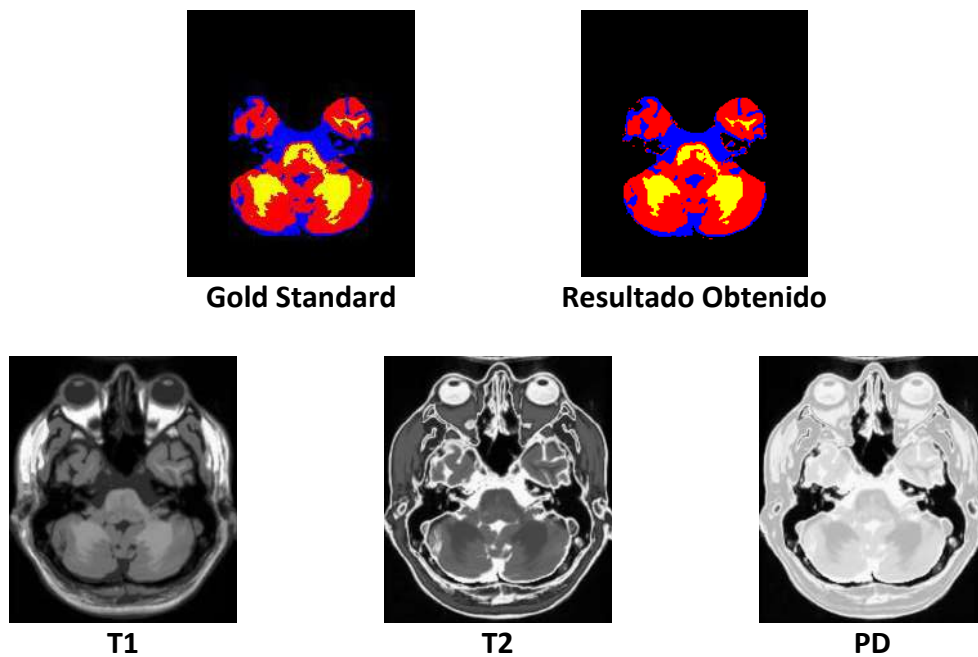
---

#### 5.1.1. Sin distorsión

En la secuencia de la Figura 31 a la Figura 35 pueden observarse los resultados del procesamiento de los cortes #30, #60, #90, #120, #150 de una imagen 3D obtenida por simulación sin distorsión de ningún tipo. El sistema fue optimizado con una muestra de píxeles del corte #90 y el resto de los mismos fue procesado con los parámetros así obtenidos. Se eligieron aleatoriamente 50 píxeles de cada tejido a reconocer.

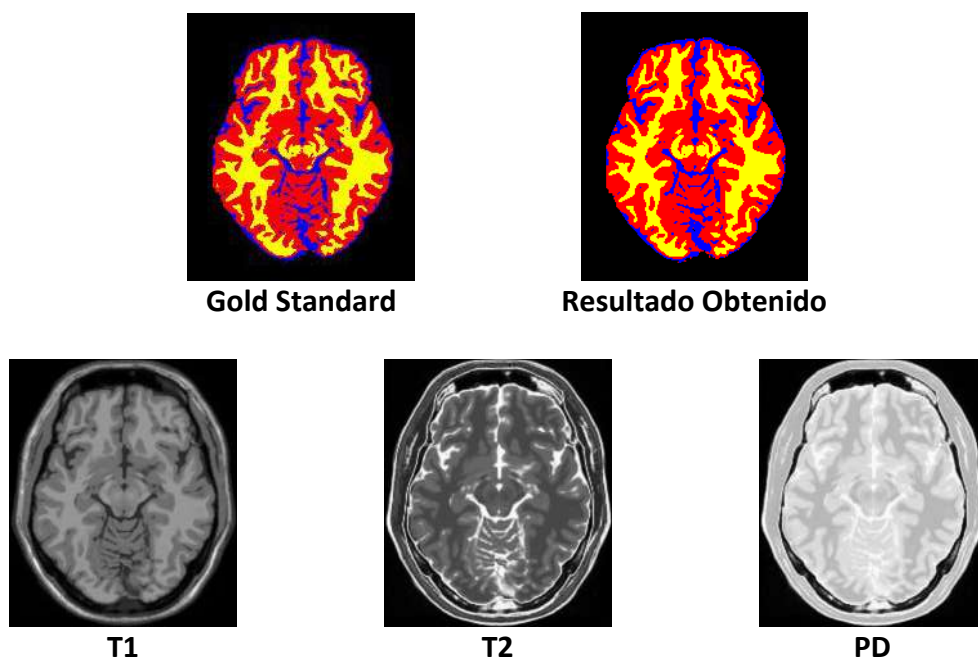
En la fila superior de cada figura se muestra la segmentación objetivo, el Gold Standard, seguido por la segmentación lograda con el método propuesto, que se identificará con el acrónimo **PDLC** (Predicados Difusos y Lógica Compensatoria). En la fila inferior se muestran las imágenes originales sin procesamiento alguno. Como puede apreciarse en todos los casos, comparando con la segmentación de referencia, los resultados obtenidos son visualmente muy similares a los de la simulación.

En todos los casos se muestra la MB en amarillo, la MG en rojo y el LCR en azul.



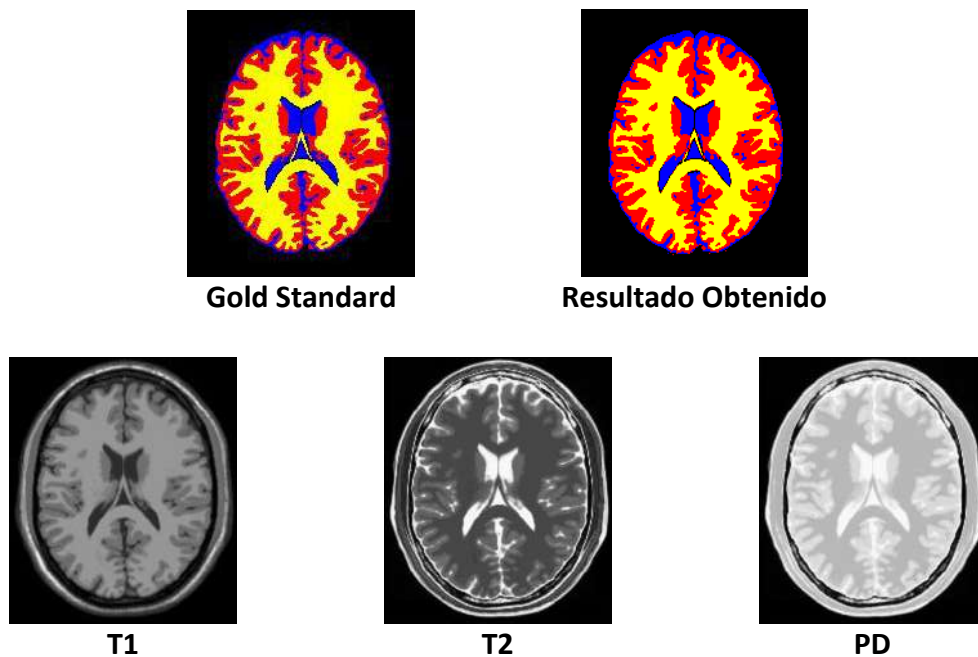
**Figura 31: Resultados obtenidos en el corte #30 de una imagen simulada, sin distorsiones (ruido y no-uniformidades en intensidad).**

En la fila superior se muestra la imagen de referencia que debería obtenerse y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes simuladas originales.



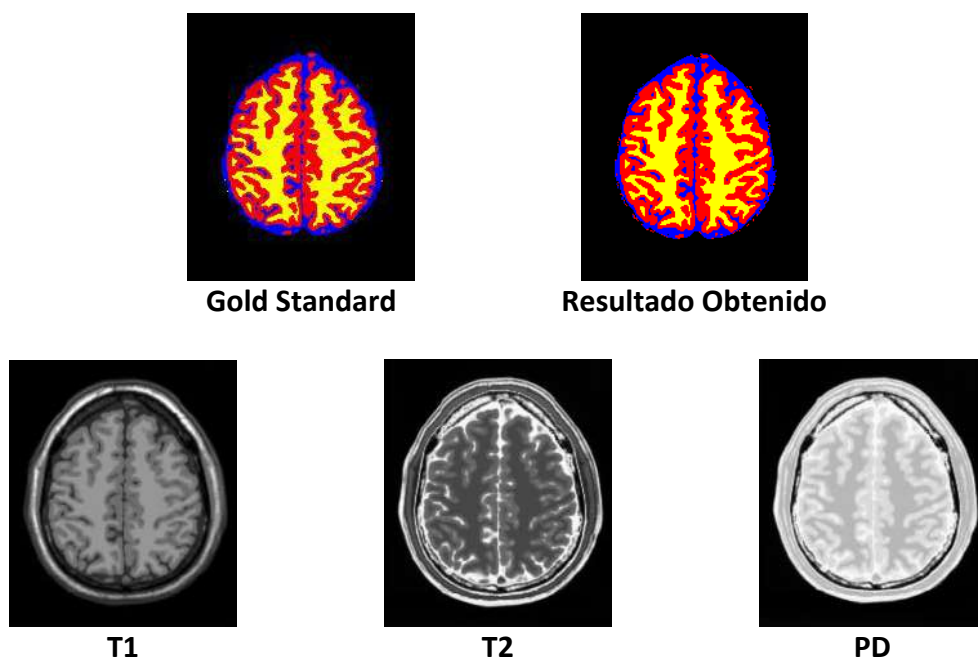
**Figura 32: Resultados obtenidos en el corte #60 de una imagen simulada, sin distorsiones (ruido y no-uniformidades en intensidad).**

En la fila superior se muestra la imagen de referencia que debería obtenerse y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes simuladas originales.



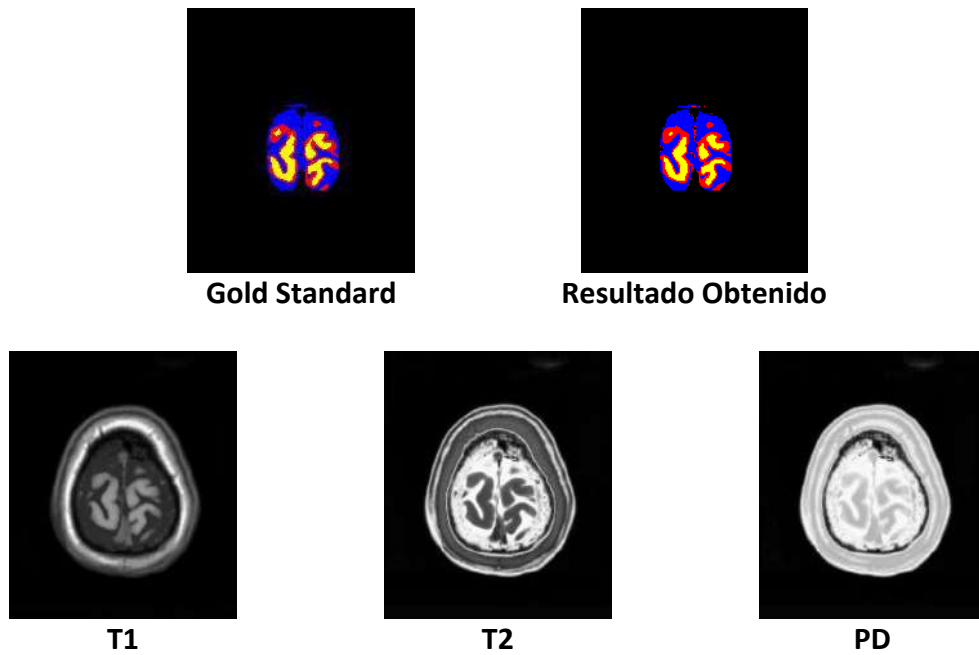
**Figura 33: Resultados obtenidos en el corte #90 de una imagen simulada, sin distorsiones (ruido y no-uniformidades en intensidad).**

En la fila superior se muestra la imagen de referencia que debería obtenerse y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes simuladas originales.



**Figura 34: Resultados obtenidos en el corte #120 de una imagen simulada, sin distorsiones (ruido y no-uniformidades en intensidad).**

En la fila superior se muestra la imagen de referencia que debería obtenerse y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes simuladas originales.



**Figura 35: Resultados obtenidos en el corte #150 de una imagen simulada, sin distorsiones (ruido y no-uniformidades en intensidad).**

En la fila superior se muestra la imagen de referencia que debería obtenerse y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes simuladas originales.

Con el fin de evaluar las segmentaciones obtenidas en estos cortes se han calculado las diferentes medidas de calidad de la segmentación presentadas en la sección 4.7.1, las que se muestran en la Tabla VIII.

**Tabla VIII: Medidas de calidad de la segmentación obtenidas en 5 cortes (imágenes simuladas).**

CT = Coeficiente de Tanimoto, ACC = Exactitud, MCR = Índice de errores y DICE = Coeficiente de DICE.

| CORTE | TC   |      |      | ACC  |      |      | MCR % |      |      | DICE |      |      |
|-------|------|------|------|------|------|------|-------|------|------|------|------|------|
|       | LCR  | MG   | MB   | LCR  | MG   | MB   | LCR   | MG   | MB   | LCR  | MG   | MB   |
| #30   | 0.91 | 0.93 | 0.92 | 0.93 | 0.99 | 0.92 | 6.54  | 1.26 | 7.89 | 0.95 | 0.97 | 0.96 |
| #60   | 0.91 | 0.96 | 0.92 | 0.93 | 0.99 | 0.96 | 6.84  | 0.78 | 3.84 | 0.95 | 0.98 | 0.98 |
| #90   | 0.95 | 0.95 | 0.98 | 0.96 | 0.99 | 0.98 | 3.73  | 0.78 | 2.28 | 0.97 | 0.97 | 0.99 |
| #120  | 0.93 | 0.96 | 0.98 | 0.94 | 0.99 | 0.98 | 5.39  | 0.92 | 2.01 | 0.96 | 0.98 | 0.99 |
| #150  | 0.94 | 0.82 | 0.91 | 0.94 | 0.98 | 0.91 | 5.72  | 1.73 | 8.57 | 0.97 | 0.90 | 0.96 |

#### 5.1.1.1. Medidas de calidad

Los **Coeficientes de Tanimoto (TC)** obtenidos son siempre mayores a 0.91 excepto para la materia gris en el corte #150, en que el valor es de 0.82. Este caso muestra la sensibilidad de este coeficiente cuando el corte evaluado contiene relativamente pocos píxeles correspondientes al tejido. Una baja cantidad de píxeles

erróneamente clasificados hace que el valor del coeficiente se reduzca considerablemente. Este caso se puede observar en la Figura 35, que presenta una pequeña cantidad de píxeles de materia gris, representados en color rojo. Los mayores valores de este coeficiente se obtienen en los cortes centrales, donde se encuentran cantidades mayores de píxeles correspondientes a todos los tejidos.

En general, como puede observarse en la segunda columna de la Tabla VIII, los valores obtenidos de **Exactitud (ACC)** son mayores que los TC. Debe recordarse que no pueden ser comparados entre sí porque responden a diferentes criterios de evaluar la calidad de la segmentación.

Los valores del **Error de Clasificación (MCR)** se observan siempre menores al 8.6%. En la Tabla VIII esta medida se da separadamente para cada sustancia pero suele reportarse promediando los valores obtenidos en todos los tejidos y en todos los cortes considerados.

En todos los casos los resultados obtenidos para el **Coeficiente de Dice (DICE)** muestran valores superiores a 0.90. Debe recordarse que este coeficiente se encuentra algebraicamente relacionado con el TC.

#### 5.1.1.2. Matrices de confusión

Las **matrices de confusión** correspondientes a todos los cortes mostrados se muestran de manera resumida en la Tabla IX. Se aprecia que los valores obtenidos son similares en todos los cortes y que siempre los valores más altos se encuentran en la diagonal principal, como es lo esperado. Debe recordarse que no se dan cantidades de píxeles de cada sustancia sino la relación cantidad de píxeles relacionados a la cantidad total de píxeles de la imagen de referencia.

**Tabla IX: Matrices de confusión obtenidas en cinco cortes en imágenes simuladas sin distorsión.**

| # CORTE | MATRIZ DE CONFUSIÓN |             |             |
|---------|---------------------|-------------|-------------|
| #30     | <b>0.93</b>         | 0.07        | 0.00        |
|         | 0.01                | <b>0.99</b> | 0.00        |
|         | 0.00                | 0.08        | <b>0.92</b> |
| #60     | <b>0.93</b>         | 0.07        | 0.00        |
|         | 0.01                | <b>0.99</b> | 0.00        |
|         | 0.00                | 0.04        | <b>0.96</b> |
| #90     | <b>0.96</b>         | 0.04        | 0.00        |
|         | 0.01                | <b>0.99</b> | 0.00        |
|         | 0.00                | 0.02        | <b>0.98</b> |
| #120    | <b>0.95</b>         | 0.05        | 0.00        |
|         | 0.01                | <b>0.99</b> | 0.00        |
|         | 0.00                | 0.02        | <b>0.98</b> |
| #150    | <b>0.94</b>         | 0.06        | 0.00        |
|         | 0.02                | <b>0.98</b> | 0.00        |
|         | 0.00                | 0.09        | <b>0.91</b> |

Es posible construir una única matriz de confusión para el volumen completo, teniendo en cuenta el total de píxeles para cada sustancia en todos los cortes analizados. En este caso la matriz que se obtiene se muestra en la Tabla X.

En esta tabla se evidencia la similitud entre los valores obtenidos en los diferentes cortes con el resultado general en el volumen completo.

**Tabla X: Matriz de confusión para la clasificación obtenida considerando un volumen completo (imagen simulada, sin distorsión).**

|                 | Clasificado como LCR | Clasificado como MG | Clasificado como MB |
|-----------------|----------------------|---------------------|---------------------|
| <b>LCR real</b> | <b>0.94</b>          | 0.06                | 0.00                |
| <b>MG real</b>  | 0.01                 | <b>0.99</b>         | 0.00                |
| <b>MB real</b>  | 0.00                 | 0.03                | <b>0.97</b>         |

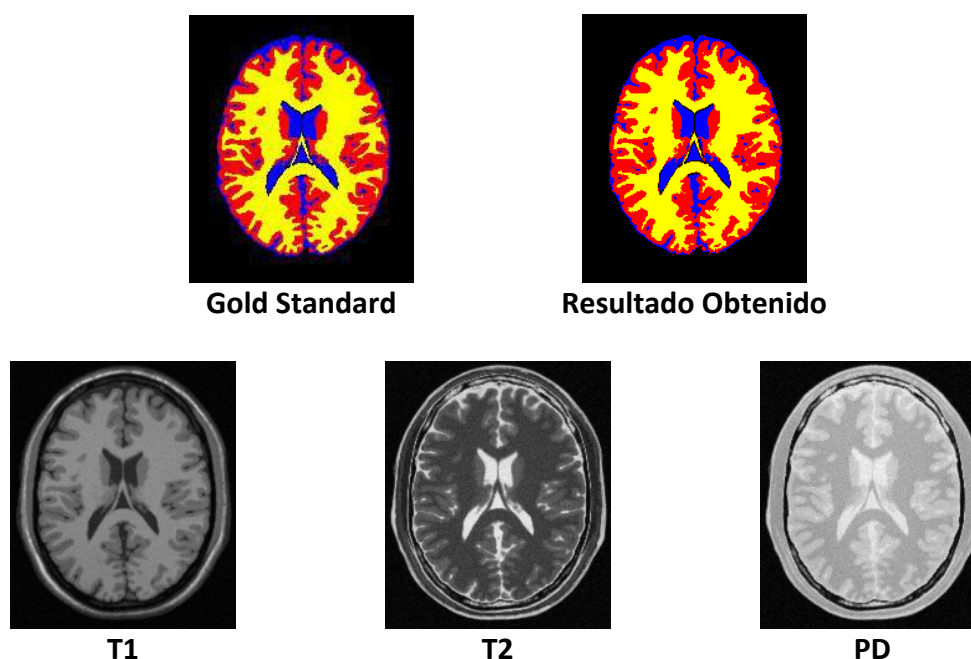
### 5.1.2. Con distorsión

Con el fin de determinar la robustez del método frente a la presencia de distorsiones se procesaron imágenes simuladas con diferentes niveles de ruido (0%, 1%, 3%, 5%, 7% y 9%) y de INU (0%, 20%, 40%).

Se recurre en esta etapa únicamente a la evaluación del TC debido a que constituye una medida exigente de la calidad de la segmentación obtenida, dado que tiene en cuenta consideraciones espaciales de la segmentación y además es muy utilizado en la bibliografía, lo que facilita la comparación con diversos métodos.

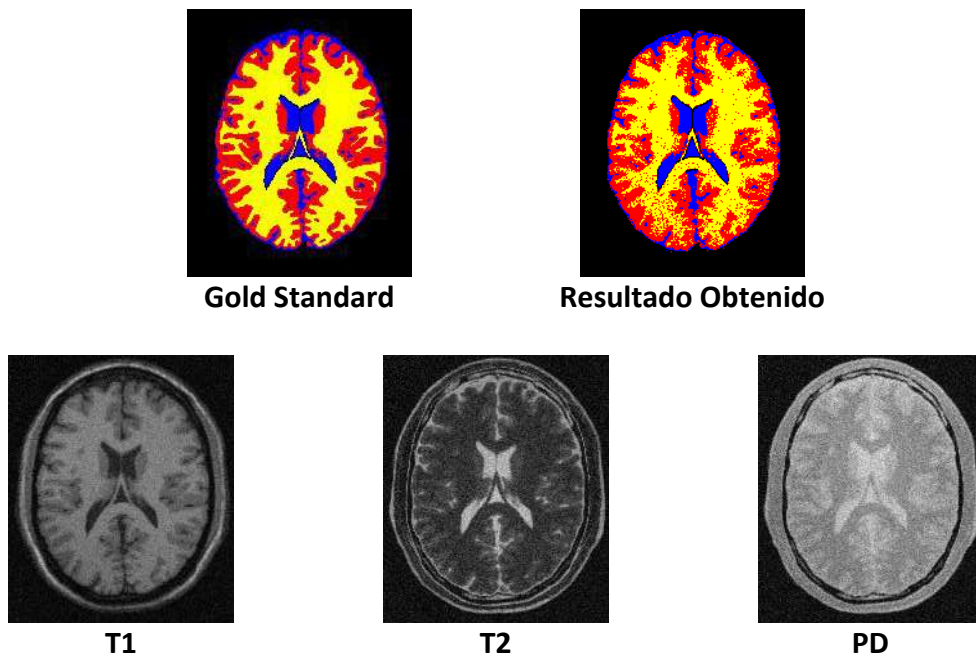
En la Figura 36 se muestran los resultados (utilizando la salida gráfica del software desarrollado) para un nivel medio de ambas distorsiones: nivel de ruido de 3% y de INU de 20%. El corte es el mismo que se mostró en la Figura 33 sin distorsiones. Pueden apreciarse píxeles con errores de clasificación, lo que hace disminuir el valor de los TC obtenidos (aunque siempre se mantienen mayores a 0.91).

El caso extremo en cuanto a distorsiones de la imagen se ve en la Figura 37 en el que se ha procesado un corte con el máximo nivel de ruido (9%) y de INU (40%) disponibles. Los valores del TC han descendido para este corte en particular, siendo aún aceptables (todos mayores a 0.73).



**Figura 36: Resultado obtenido con nivel de ruido de 3% y de INU de 20%.**

Se muestra el corte #90 pero con nivel de ruido de 3% y de INU de 20%. La distorsión es de valores estándar para efectuar pruebas. La detección de los tejidos es aún adecuada: los Coeficientes de Tanimoto siguen siendo de valores elevados, mayores a 0.91.



**Figura 37: Resultado obtenido en el caso de máximo nivel de ruido y de INU disponible en simulación.**

Se muestra el corte #90 pero con nivel de ruido de 9% y de INU de 40%. La distorsión es máxima. Se observa que la detección de los tejidos, en particular de la materia blanca, se ha degradado notablemente. Sin embargo los Coeficientes de Tanimoto toman valores superiores a 0.73 para todas las sustancias.

En la Tabla XI se muestran los valores promedio y desviaciones estándar de los Coeficientes de Tanimoto calculados en 120 cortes, para distintos niveles de ruido y de INU. Puede observarse cómo los valores medios en general disminuyen al incrementar el ruido, desplazándose de arriba hacia abajo en la tabla (por ejemplo, para el LCR, con INU de 20%, el TC comienza en 0.93 para el caso sin ruido y termina en 0.85 para el caso de peor ruido), o al incrementar la INU, desplazándose de izquierda a derecha (por ejemplo, para el LCR, con ruido de 3%, el TC comienza en 0.93 para el caso sin INU y termina en 0.89 para el caso de INU máxima). Los bajos valores de las desviaciones estándar indican una adecuada coherencia de la segmentación entre cortes y también son indicadores de una buena estabilidad del método.

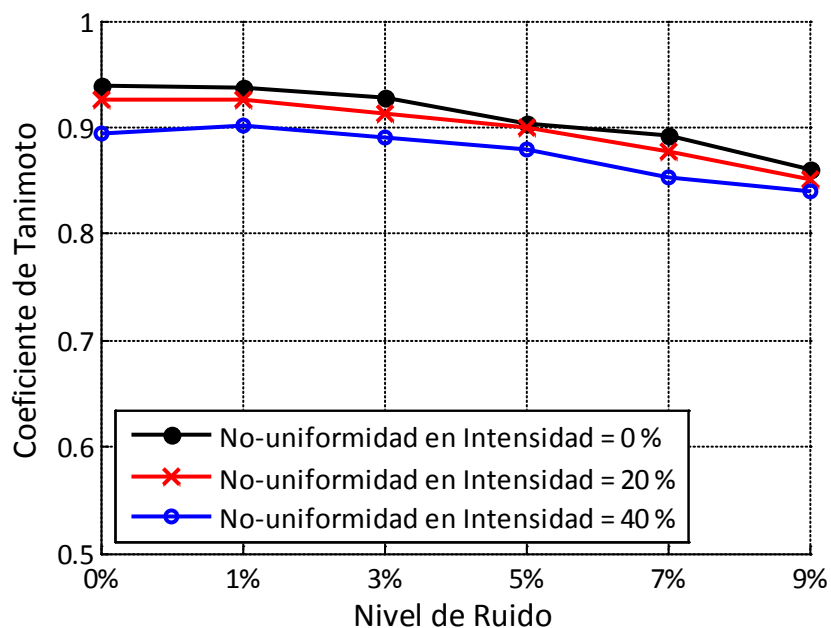
**Tabla XI: Coeficientes de Tanimoto obtenidos en 120 cortes de imágenes 3D simuladas con diferentes distorsiones.**

TCprom = Valores promedio,  $\sigma$  = desviaciones estándar.

|                  |          | INU = 0 %   |             |             | INU = 20 %  |             |             | INU = 40 %  |             |             |
|------------------|----------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
|                  |          | LCR         | MG          | MB          | LCR         | MG          | MB          | LCR         | MG          | MB          |
| <b>RUIDO 0 %</b> | TCprom   | <b>0.94</b> | <b>0.95</b> | <b>0.97</b> | <b>0.93</b> | <b>0.93</b> | <b>0.95</b> | <b>0.89</b> | <b>0.86</b> | <b>0.88</b> |
|                  | $\sigma$ | 0.01        | 0.01        | 0.01        | 0.02        | 0.01        | 0.01        | 0.04        | 0.03        | 0.06        |
| <b>RUIDO 1 %</b> | TCprom   | <b>0.94</b> | <b>0.94</b> | <b>0.95</b> | <b>0.93</b> | <b>0.93</b> | <b>0.95</b> | <b>0.90</b> | <b>0.88</b> | <b>0.91</b> |
|                  | $\sigma$ | 0.01        | 0.01        | 0.01        | 0.02        | 0.01        | 0.02        | 0.04        | 0.02        | 0.05        |
| <b>RUIDO 3 %</b> | TCprom   | <b>0.93</b> | <b>0.93</b> | <b>0.94</b> | <b>0.91</b> | <b>0.91</b> | <b>0.93</b> | <b>0.89</b> | <b>0.86</b> | <b>0.89</b> |
|                  | $\sigma$ | 0.01        | 0.01        | 0.02        | 0.02        | 0.01        | 0.03        | 0.04        | 0.02        | 0.05        |
| <b>RUIDO 5 %</b> | TCprom   | <b>0.94</b> | <b>0.89</b> | <b>0.90</b> | <b>0.90</b> | <b>0.87</b> | <b>0.89</b> | <b>0.88</b> | <b>0.83</b> | <b>0.84</b> |
|                  | $\sigma$ | 0.01        | 0.01        | 0.03        | 0.02        | 0.01        | 0.04        | 0.04        | 0.02        | 0.07        |
| <b>RUIDO 7 %</b> | TCprom   | <b>0.89</b> | <b>0.84</b> | <b>0.84</b> | <b>0.88</b> | <b>0.83</b> | <b>0.84</b> | <b>0.85</b> | <b>0.79</b> | <b>0.80</b> |
|                  | $\sigma$ | 0.02        | 0.02        | 0.04        | 0.03        | 0.02        | 0.04        | 0.05        | 0.02        | 0.07        |
| <b>RUIDO 9 %</b> | TCprom   | <b>0.86</b> | <b>0.77</b> | <b>0.78</b> | <b>0.85</b> | <b>0.78</b> | <b>0.78</b> | <b>0.84</b> | <b>0.76</b> | <b>0.76</b> |
|                  | $\sigma$ | 0.02        | 0.02        | 0.04        | 0.03        | 0.03        | 0.05        | 0.05        | 0.03        | 0.07        |

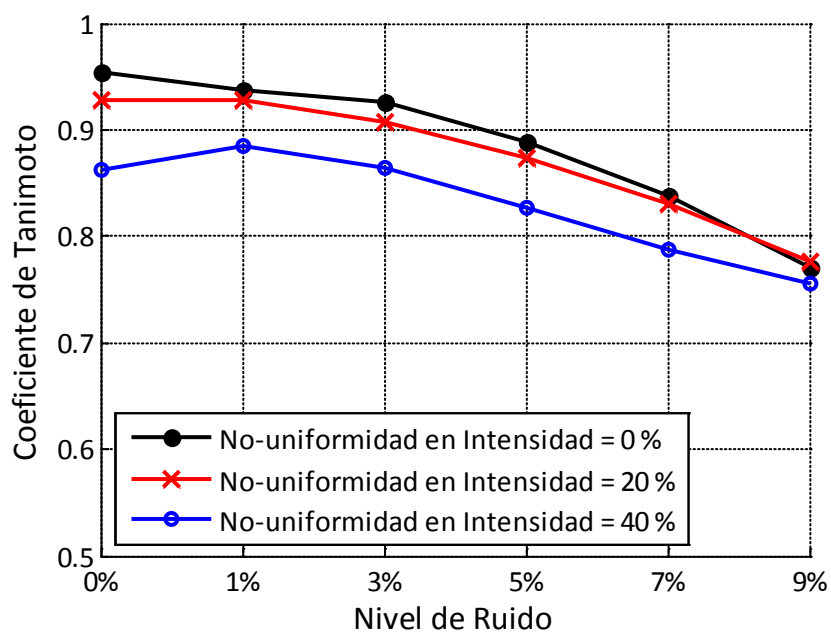
Los valores promedio de los TC presentados en la tabla anterior para diferentes niveles de distorsión se han representado gráficamente en la Figura 38 para el LCR, en la Figura 39 para la MG y en la Figura 40 para la MB. Puede observarse gráficamente la degradación de la segmentación cuando aumentan las distorsiones. Pero aún en los peores casos de ruido y de INU el valor del TC supera el valor 0.75 (peor caso, visible tanto en la MG como en la MB). Para el LCR los resultados son satisfactorios, con valores superiores a 0.84 para la peor condición de distorsión de la imagen.

Se observa una pequeña mejora cuando existen distorsiones de muy poco nivel de ruido (1%) con máximo nivel de INU (40%) que ya ha sido notada en otros análisis [Santalla et al., 2007].



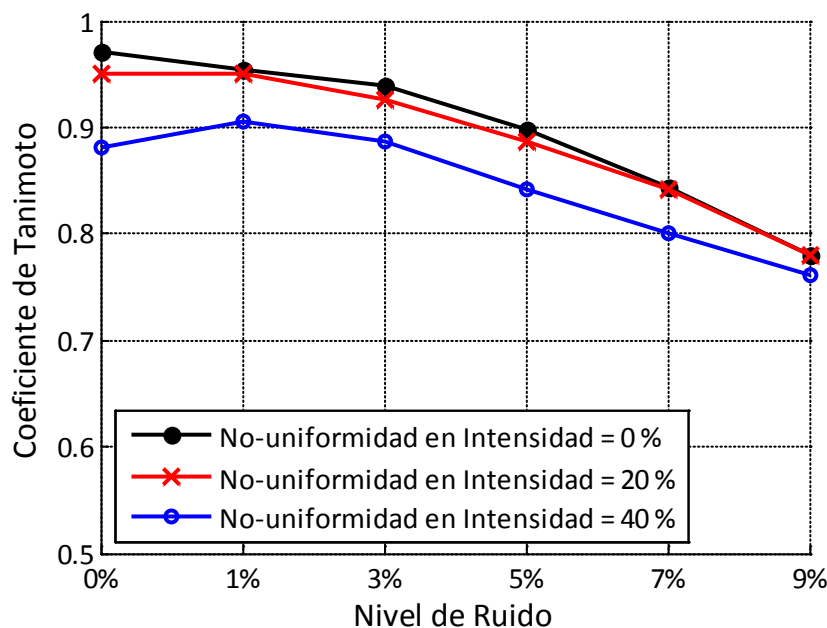
**Figura 38: Coeficientes de Tanimoto obtenidos con el método propuesto para el líquido cefalorraquídeo (LCR).**

Pueden observarse los valores promedio de los Coeficientes de Tanimoto obtenidos con el método propuesto en imágenes simuladas para todas las combinaciones de ruido (1, 3, 5, 7 y 9%) y de INU (0, 20, 40%) disponibles.



**Figura 39: Coeficientes de Tanimoto obtenidos con el método propuesto para la materia gris (MG).**

Pueden observarse los valores promedio de los Coeficientes de Tanimoto obtenidos con el método propuesto en imágenes simuladas para todas las combinaciones de ruido (1, 3, 5, 7 y 9%) y de INU (0, 20, 40%) disponibles.



**Figura 40: Coeficientes de Tanimoto obtenidos con el método propuesto para la materia blanca (MB).**

Pueden observarse los valores promedio de los Coeficientes de Tanimoto obtenidos con el método propuesto en imágenes simuladas para todas las combinaciones de ruido (1, 3, 5, 7 y 9%) y de INU (0, 20, 40%) disponibles.

### 5.1.3. Con los operadores MAX y MIN

Con el fin de demostrar que la utilización de la LDC en el cálculo de los valores de verdad de los predicados mejora altamente la segmentación con respecto a los operadores clásicos, se ha efectuado el procesamiento de todas las imágenes simuladas con las diferentes distorsiones utilizando las operaciones MAX y MIN para la implementación de los conectivos AND y OR respectivamente.

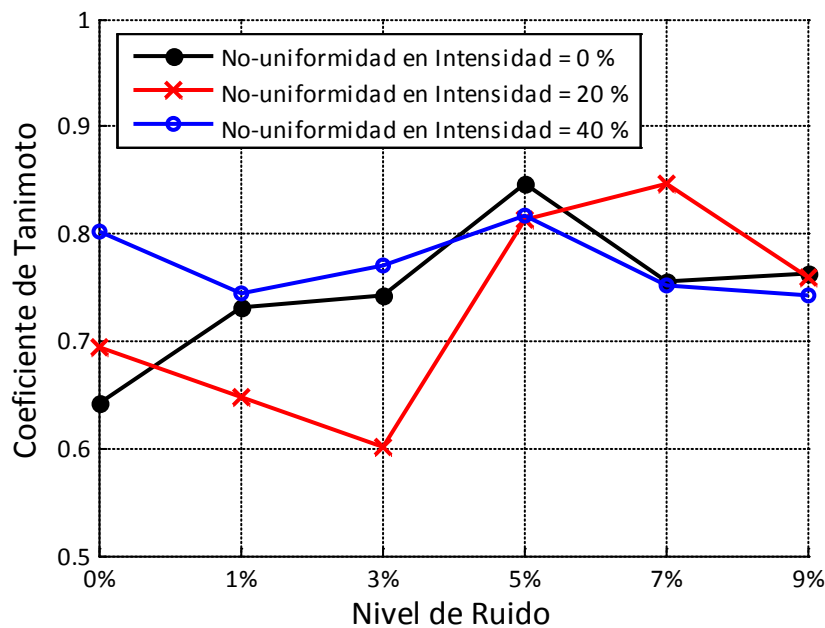
Los resultados para el LCR pueden observarse en la Figura 41. Los valores del TC obtenido son siempre menores que 0.85, lo que no es un valor inadmisibles, pero es superado ampliamente por lo obtenido con LDC, mostrado en la sección anterior.

Similarmente, en la Figura 42 se ven los resultados para la MG, donde el mayor valor es 0.87 y el menor valor es 0.6.

En la Figura 43 se ofrecen los resultados para la MB. En este caso el TC llega a tomar un valor máximo de 0.93, pero los resultados son inadmisibles en la máxima condición de distorsiones, donde apenas se alcanza el valor 0.55.

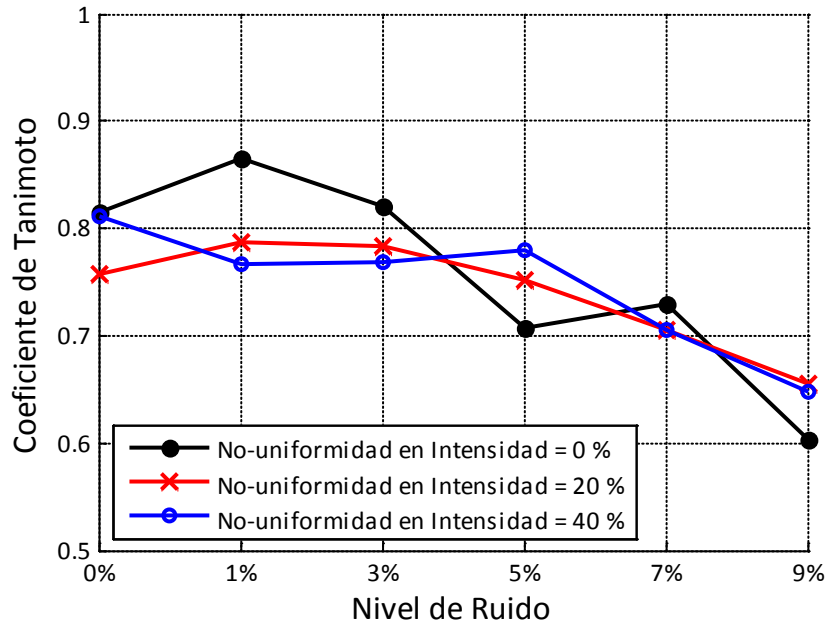
El comportamiento al ir aumentando las distorsiones progresivamente se muestra errático e incoherente, en contraposición al que puede observarse en la sección anterior.

Los resultados mostrados en esta sección justifican el uso de los operadores de la LDC, que mostraron un desempeño superior. Los TC obtenidos con los operadores MAX-MIN muestran valores considerablemente menores que los obtenidos con LDC, especialmente en presencia de ruido y de INU.



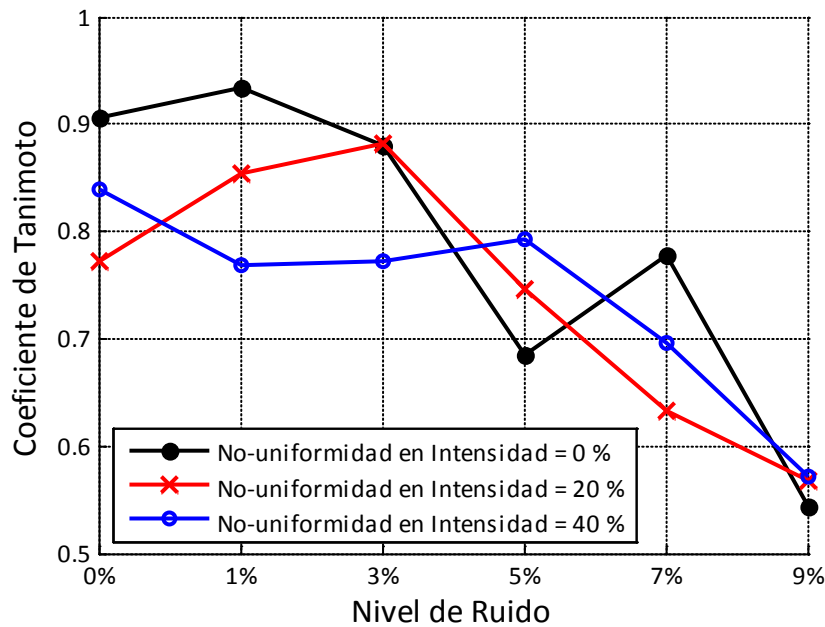
**Figura 41: Coeficientes de Tanimoto obtenidos con los operadores MAX – MIN para el líquido cefalorraquídeo (LCR).**

Pueden observarse los valores promedio de los Coeficientes de Tanimoto obtenidos con el método propuesto pero con operadores MAX y MIN para las operaciones AND y OR respectivamente en imágenes simuladas para todas las combinaciones de ruido (1, 3, 5, 7 y 9%) y de INU (0, 20, 40%) disponibles.



**Figura 42: Coeficientes de Tanimoto obtenidos con los operadores MAX – MIN para la materia gris (MG).**

Pueden observarse los valores promedio de los Coeficientes de Tanimoto obtenidos con el método propuesto pero con operadores MAX y MIN para las operaciones AND y OR respectivamente en imágenes simuladas para todas las combinaciones de ruido (1, 3, 5, 7 y 9%) y de INU (0, 20, 40%) disponibles.



**Figura 43: Coeficientes de Tanimoto obtenidos con los operadores MAX – MIN para la materia blanca (MB).**

Pueden observarse los valores promedio de los Coeficientes de Tanimoto obtenidos con el método propuesto pero con operadores MAX y MIN para las operaciones AND y OR respectivamente en imágenes simuladas para todas las combinaciones de ruido (1, 3, 5, 7 y 9%) y de INU (0, 20, 40%) disponibles.

#### 5.1.4. Comparaciones con otros métodos

En esta sección se ofrece una comparación de los coeficientes de calidad de la segmentación obtenidos en comparación con otros métodos disponibles en la bibliografía, en trabajos en los que se ha utilizado la misma base de IRM simuladas [Montréal Neurological Institute, 2007]. Dado que los autores reportan sus resultados utilizando diferentes mediciones de calidad, se han calculado las mismas para efectuar las comparaciones.

En [Song et al., 2007] se dan resultados obtenidos en estas imágenes para algunas configuraciones específicas de ruido y de INU utilizando el método propuesto en el trabajo y otros métodos. Se presentan los valores medios del MCR y la desviación estándar calculados en 5 cortes centrales (#90, #95, #100, #105, #110). En la Tabla XII se reproducen los valores a los que se le ha agregado los resultados obtenidos con el método propuesto en esta tesis (PDLC). Los métodos evaluados se indican de la siguiente manera:

- **KME:** algoritmo K-medias [Bezdek et al., 1993] para clasificar automáticamente los vectores formados por la intensidad de T1 y la intensidad de T2. Para asegurar la convergencia y conocer a qué tejido corresponde cada grupo clasificado se requiere una adecuada inicialización de los centros de *cluster*.
- **EM:** método de clasificación que utiliza el algoritmo original “*Expectation – Maximization* (EM)” [Wells et al., 1996].
- **FSL:** método que utiliza un modelo de cadenas ocultas de Markov, que tiene en cuenta información especial de la vecindad de los píxeles analizados. También utiliza el algoritmo el algoritmo EM [Zhang et al., 2001]. Se encuentra incluido en la librería FSL, disponible en Internet, bajo el nombre de FAST.
- **PNN:** utiliza como clasificador una Red Neuronal Probabilística (PNN, *Probabilistic Neural Network*) que debe ser entrenada previamente.
- **WPNN:** modificación del algoritmo anterior mediante la utilización de un mapa autoorganizado previamente entrenado.
- **PDLC:** Predicados Difusos con Lógica Compensatoria, el método propuesto.

**Tabla XII: Valores promedio y desviaciones estándar del MCR obtenidos en 5 cortes, para comparación de resultados con otros algoritmos (Reproducido de [Song 2007]).**

|             |          | KME  | EM   | FSL  | PNN  | WPNN | PDL   |
|-------------|----------|------|------|------|------|------|-------|
| Noise = 3 % | MCRprom  | 3.96 | 3.62 | 4.21 | 3.57 | 3.41 | 3.13  |
| INU = 20 %  | $\sigma$ | 0.20 | 0.12 | 0.08 | 0.10 | 0.10 | <0.01 |
| Noise = 9 % | MCRprom  | 8.61 | 7.31 | 6.92 | 7.44 | 6.07 | 6.87  |
| INU = 40 %  | $\sigma$ | 0.36 | 0.27 | 0.16 | 0.29 | 0.15 | <0.01 |

Puede observarse en la Tabla XII que el algoritmo propuesto mejora el MCR a los demás algoritmos para el caso de ruido de menor intensidad (la mejora es de 3.41% para el método WPNN), con significancia estadística ( $p < 0.0001$ ). Se realizó un test de Wilcoxon para comprobar la hipótesis de que el método propuesto mejora los resultados de la segmentación. Para el caso de ruido de mayor nivel el método no puede alcanzar el desempeño del método WPNN pero sí mejora el de todos los otros ( $p < 0.0001$ ). La sensibilidad que muestra el método a tan altos niveles de ruido se debe esencialmente a que el procesamiento es píxel a píxel y no contempla ningún tipo de corrección.

En el trabajo de [Li et al., 2006] se presenta un método basado en contornos activos, utilizando un análisis del histograma para generar las regiones activas. Se evalúa el resultado con la misma definición del TC y los valores que se reportan para el LCR, MG y MB son 0.914, 0.883 y 0.898 respectivamente, utilizando un volumen con 3% de ruido y 20% de INU. Los valores obtenidos en el presente trabajo son 0.913, 0.908 y 0.927 dados en el mismo orden. Este método sólo requiere la imagen pesada en T1 pero su complejidad lleva a tiempos computacionales superiores al propuesto. En términos de la segmentación obtenida, el algoritmo propuesto mejora al presentado en ese trabajo en un 2% en promedio.

En [Jimenez-Alaniz et al., 2006] también se utilizan volúmenes con 3% de ruido y 20% de INU. En este caso los resultados son reportados con Coeficientes de Tanimoto (valores medios  $\pm$  desviaciones estándar) cuyos valores son:  $0.871 \pm 0.031$  para LCR,  $0.9 \pm 0.012$  para MG y  $0.924 \pm 0.038$  para MB. El presente método para este caso reporta los siguientes valores:  $0.913 \pm 0.023$  para LCR,  $0.908 \pm 0.013$  para MG y

0.927  $\pm$  0.025 para MB. Esto representa una mejora en la calidad de la segmentación también en 2% promedio, con estos porcentajes de distorsión.

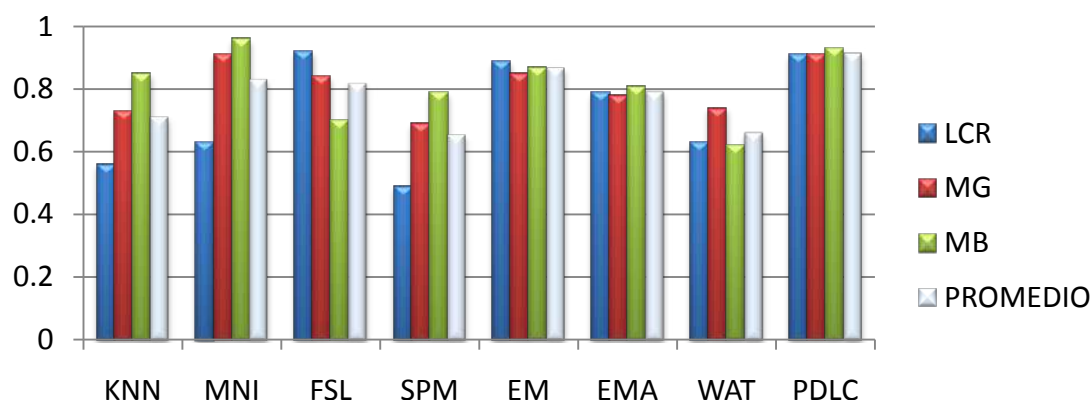
En [Bouix et al., 2007] se realiza un análisis de los TC en diferentes métodos de segmentación en RM de cerebro, que se han indicado con siglas según se detalla a continuación:

- **KNN:** Clasificación estadística realizada con un clasificador de los k vecinos más próximos que ha sido entrenado por registración no lineal de un atlas [Warfield, 1996].
- **MNI:** Clasificador basado en una Red Neuronal de retropropagación entrenada automáticamente por registración afín de un atlas. Este método incluye corrección de las INU [Zijdenbos et al., 2002].
- **FSL:** Mencionado anteriormente.
- **SPM:** Acrónimo de *Statistical Parametric Mapping*. Algoritmo de agrupamiento por mezcla de modelos gaussianos (*mixture model clustering*) que incluye corrección de las INU [Ashburner and Friston, 2003, Ashburner and Friston, 2005].
- **EM:** Mencionado anteriormente.
- **EMA:** algoritmo similar al FAST, incorporando información *a priori* alineada por registración no lineal [Pohl et al., 2004].
- **WAT:** Segmentación basada en la Transformada Watershed, incorporando incorporando información *a priori* alineada por registración no lineal de un atlas [Grau et al., 2004]. La elección del atlas puede producir sesgos en la salida.

Se reportan los TC logrados, como en el caso anterior, en volúmenes con 3% de ruido y 20% de INU, los que se representan gráficamente en la Figura 44, junto a los obtenidos con el método propuesto. Se observan los TC con barras para cada tejido clasificado y una barra adicional para el promedio. Las últimas cuatro barras representan el método propuesto.

Como se aprecia en esa Figura el valor del TC para el método propuesto es similar al del FSL, que es el mejor de los presentados, para la segmentación de LCR.

También iguala al MNI, que es el que mejor segmenta la MG y se ubica en segundo lugar en el desempeño en la segmentación de MB, sólo superado por MNI. Este algoritmo incluye corrección de las INU que no se han aplicado en el método propuesto. El desempeño promedio, considerando los tres tejidos segmentados, el método PDLC supera a todos los demás, mejorando entre 5% (para el EMS) y un 28% (para SPM y WAT, que tienen desempeños similares entre sí).



**Figura 44: Representación gráfica de los Coeficientes de Tanimoto obtenidos a través de diferentes métodos en imágenes simuladas.**

Corresponde a imágenes con 3 % de ruido y 20 % de no-uniformidad (datos provistos por los autores de [Bouix 2007]). La última barra de cada grupo corresponde al desempeño promedio en los tres tejidos. KNN: K-vecinos más próximos; MNI: Redes Neuronales y atlas; FSL: cadenas ocultas de Markov; SPM: mezcla de modelos gaussianos; EM: Expectation-Maximization; EMA: similar al FSL; WAT: Transformada Watershed; PDLC: método propuesto.

### 5.1.5. Análisis del TC mediante test estadístico

Para efectuar un análisis comparativo detallado que permita la realización de tests estadísticos y realizar una posterior comparación con resultados obtenidos por otros métodos en imágenes reales, se evaluó el desempeño de diferentes algoritmos de clasificación, algunos en software disponible y otros programados para esta comparación, que se indicarán de la siguiente manera:

- **BRA:** Acrónimo de “*Brain Research: Analysis of Images, Networks, and Systems*”. Software desarrollado en la Universidad de Iowa [Iowa Mental Health Clinical Research Center, 2005]. Las imágenes han sido segmentadas operando este sistema usuarios entrenados del Instituto Fleni (Buenos Aires, Argentina), quienes entregaron los resultados.

- **FCM:** Algoritmo Fuzzy-C-Means sin modificaciones. Este método es automático y no supervisado pero se deben hacer verificaciones posteriores para lograr que los *clusters* obtenidos correspondan con los esperados.
- **KME:** Algoritmo de las K-medias, configurando como centros iniciales a los valores de gris correspondientes a un píxel de una imagen de referencia.
- **KNN:** Clasificación estadística realizada con un clasificador de los k vecinos más próximos que ha sido entrenado con 30 píxeles de cada sustancia que se obtienen del corte #100.
- **NNE:** Red neuronal multicapa, alimentada hacia delante (*feed forward*). Se estableció una arquitectura de 2 capas de neuronas escondidas y cuatro salidas para identificar cada tejido y se entrenó con el corte #100.
- **PNN:** Red Neuronal Probabilística (PNN, *Probabilistic Neural Network*) entrenada con el corte #100.
- **PDLC:** Predicados Difusos con Lógica Compensatoria, el método propuesto.

Se evaluaron los TC obtenidos en los 100 cortes centrales de una imagen 3D simulada con un nivel de ruido de 3% y una INU de 20%. Como se ha presentado en la sección anterior, este nivel de distorsiones se ha utilizado en numerosos trabajos para efectuar comparaciones. Los valores se muestran en la Tabla XIII.

**Tabla XIII: Coeficientes de Tanimoto obtenidos a través de diferentes métodos en MR simuladas con 3% de ruido y 20% de no-uniformidad (software disponible o programado *ad hoc*).**

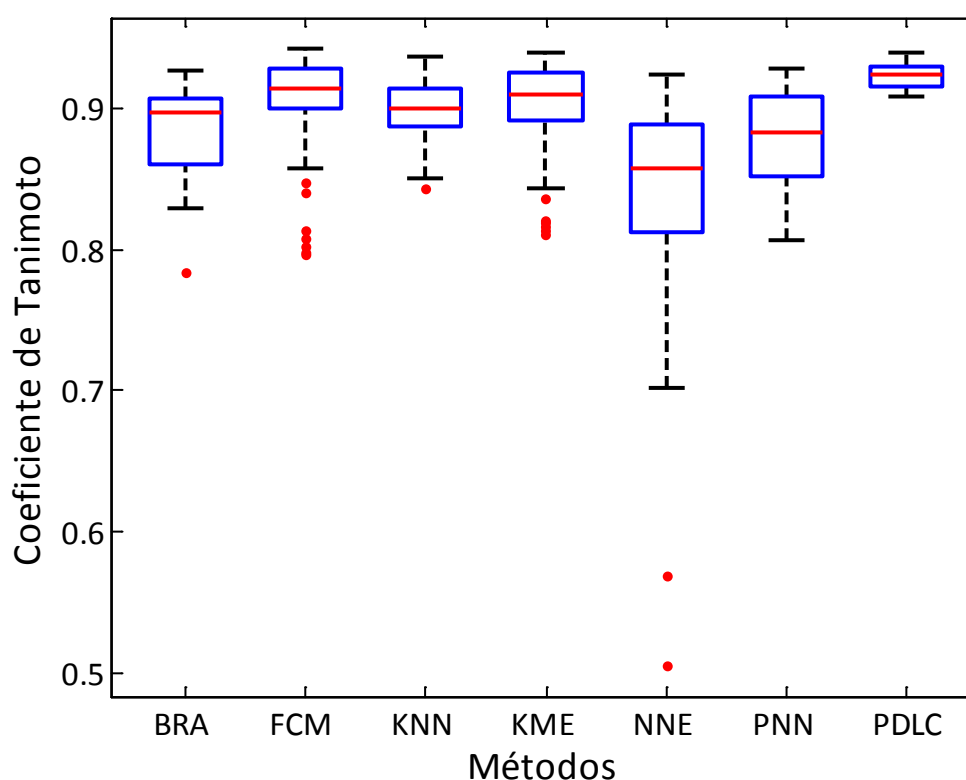
BRA: BRAINS; FCM: Fuzzy-C-Means; KME: K-medias; KNN: k vecinos más próximos; NNE: Red neuronal multicapa; PNN: Red Neuronal Probabilística; PDLC: método propuesto.

|             | LCR         | MG          | MB          | Promedio    |
|-------------|-------------|-------------|-------------|-------------|
| <b>BRA</b>  | 0.89        | 0.87        | 0.89        | <b>0.88</b> |
| <b>FCM</b>  | 0.92        | 0.90        | 0.89        | <b>0.90</b> |
| <b>KNN</b>  | 0.90        | 0.90        | 0.90        | <b>0.90</b> |
| <b>KME</b>  | 0.92        | 0.89        | 0.89        | <b>0.90</b> |
| <b>NNE</b>  | 0.80        | 0.84        | 0.88        | <b>0.84</b> |
| <b>PNN</b>  | 0.90        | 0.88        | 0.86        | <b>0.88</b> |
| <b>PDLC</b> | <b>0.91</b> | <b>0.91</b> | <b>0.93</b> | <b>0.92</b> |

Se observa en esta tabla que los valores obtenidos con el método propuesto (PDLC) son superiores a los demás métodos considerados, excepto para el LCR, donde los métodos FCM y KME lo superan ligeramente. De todas maneras, el desempeño promedio considerando los tres tejidos segmentados es superior a todos los métodos.

Sin embargo, reportando el promedio de los valores calculado sobre todos los cortes no puede analizarse la estabilidad del algoritmo en los sucesivos cortes. Con este fin se muestran en gráficos de caja en la Figura 45 los TC (valores promedio en los tres tejidos).

Los límites superior e inferior de cada “caja” indican los valores de los cuartiles. Las líneas que continúan indican la extensión del resto de los TC. Las líneas centrales corresponden a las medianas de los datos. Los *outliers* se indican con puntos.



**Figura 45: Gráficos de caja que muestra los promedios de los Coeficientes de Tanimoto (calculados en los tres tejidos).**

Se comparan los valores obtenidos con diferentes métodos en 100 cortes de una IRM simulada, con 3% de ruido y 20% de INU. BRA: BRAINS; FCM: Fuzzy-C-Means; KME: K-medias; KNN: k vecinos más próximos; NNE: Red neuronal multicapa; PNN: Red Neuronal Probabilística; PDLC: método propuesto.

Se observa en la figura anterior que el método propuesto presenta TC con mayores valores y menor dispersión. No se observan *outliers*, lo que indica que el

método es estable. Los métodos KME, NNE y FCM presentan gran cantidad de *outliers*, lo que indica cierta inestabilidad al procesar el conjunto completo de cortes.

Dado que no pudo verificarse una distribución normal en los TC, se realizó un test no paramétrico para evaluar si las diferencias entre los resultados obtenidos por el método propuesto y los demás métodos considerados son estadísticamente significativas: el test de rangos signados de Wilcoxon para datos apareados [Gibbons and Chakraborti, 2003]. En la Tabla XIV se reportan los valores de significancia de este test. Valores de  $p$  inferiores a 0.05 indican una diferencia estadísticamente significativa entre los métodos.

**Tabla XIV: Valores de significancia ( $p$ ) que devuelve el test de Wilcoxon para los TC obtenidos en 100 cortes de una imagen simulada, aplicando diferentes métodos en comparación con el método propuesto.**

BRA: BRAINS; KNN: K-vecinos más próximos; KME: K-medias; NNE: Red neuronal multicapa; FCM: Fuzzy-C-Means; PNN: Red Neuronal Probabilística.

|            | <b>p</b>              |
|------------|-----------------------|
| <b>BRA</b> | $7.56 \times e^{-10}$ |
| <b>FCM</b> | $1.07 \times e^{-7}$  |
| <b>KNN</b> | $7.56 \times e^{-10}$ |
| <b>KME</b> | $2.10 \times e^{-9}$  |
| <b>NNE</b> | $7.56 \times e^{-10}$ |
| <b>PNN</b> | $7.56 \times e^{-10}$ |

Como puede observarse, todas las diferencias entre los TC obtenidos son significativas.

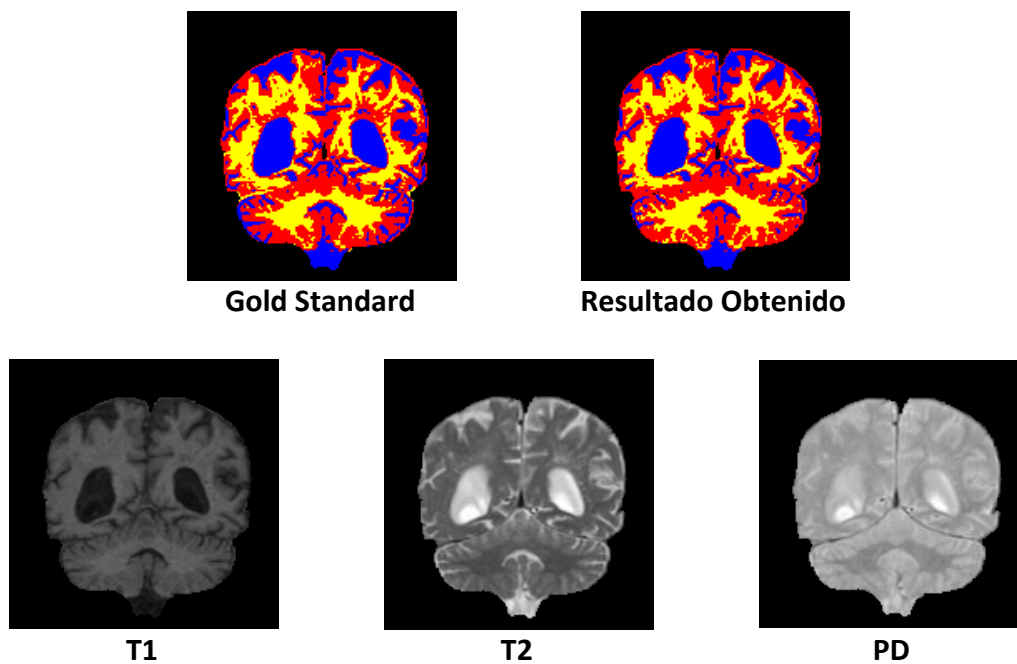
## 5.2. Con imágenes reales

Para entrenar el sistema con imágenes reales se armaron conjuntos de datos de entrenamiento del sistema, seleccionando aleatoriamente 30 píxeles de cada tejido o sustancia a detectar de las imágenes segmentadas manualmente.

## 5.2.1. Conjunto de imágenes #1

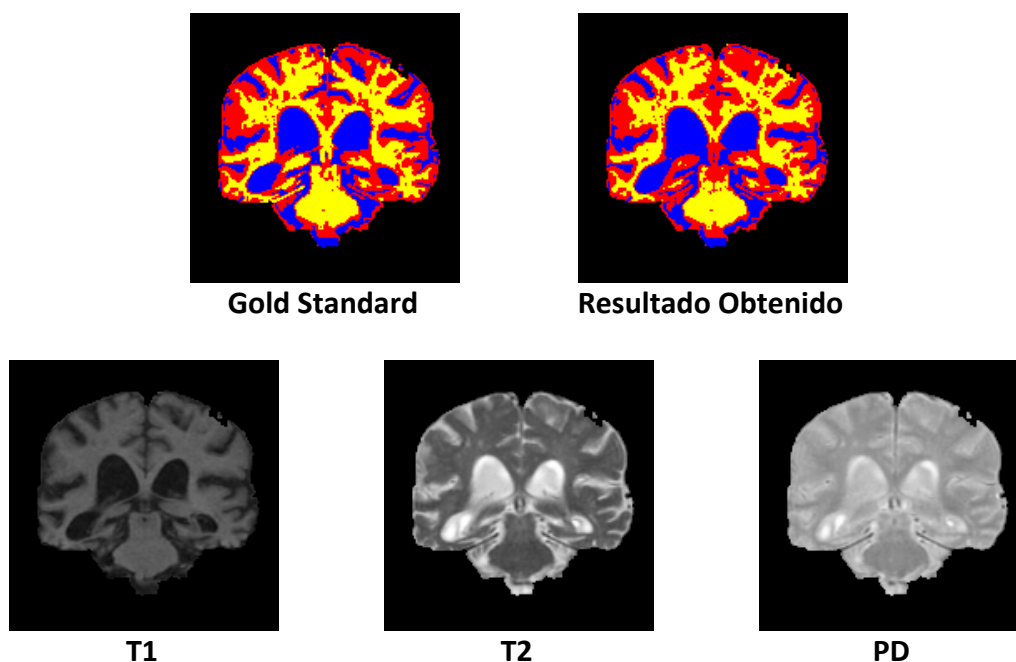
### 5.2.1.1. Mediciones de calidad de la segmentación

El resultado se muestra en la secuencia de imágenes comprendidas desde la Figura 46 hasta la Figura 48, donde también se incluye la segmentación propuesta por los especialistas. Algunas diferencias que se observan correspondientes a regiones pequeñas o píxeles aislados pueden eliminarse con post-procesamientos que no han sido incluidos en esta etapa.



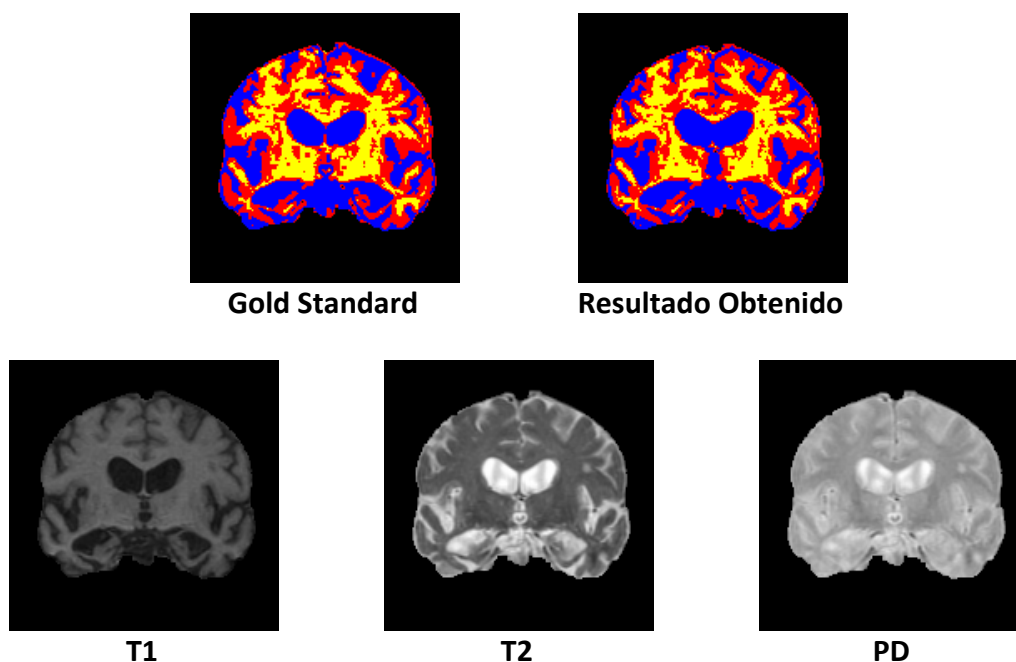
**Figura 46: Corte #59 de una imagen real (Conjunto #1).**

En la fila superior se muestra la imagen de referencia que debería obtenerse (segmentación manual realizada por especialistas) y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes originales luego de haber sido procesadas para remover las regiones de píxeles que no se clasificarán.



**Figura 47: Corte #79 de una imagen real (Conjunto #1).**

En la fila superior se muestra la imagen de referencia que debería obtenerse (segmentación manual realizada por especialistas) y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes originales luego de haber sido procesadas para remover las regiones de píxeles que no se clasificarán.



**Figura 48: Corte #99 de una imagen real (Conjunto #1).**

En la fila superior se muestra la imagen de referencia que debería obtenerse (segmentación manual realizada por especialistas) y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes originales luego de haber sido procesadas para remover las regiones de píxeles que no se clasificarán.

De la misma manera que en las imágenes simuladas, en la Tabla XV se presentan diferentes mediciones de calidad obtenidas. Los TC son en promedio inferiores a los obtenidos en imágenes simuladas. Igualmente los valores son siempre superiores a 0.85. La medición de exactitud (ACC) muestra que siempre el 88% de los píxeles son correctamente clasificados. Se muestra también la medición del MCR y el coeficiente de DICE a fines de su comparación con los de imágenes simuladas. El MCR muestra la dificultad que surge en las imágenes reales, con valores mayores que en las simulaciones. El caso de mayor error se aprecia en el corte #79 en la detección de la MB. Análogamente, el coeficiente de DICE disminuye sus valores en promedio.

**Tabla XV: Valores de diferentes mediciones de calidad en 3 cortes de una imagen real (Conjunto #1) obtenidos con el método propuesto.**

CT = Coeficiente de Tanimoto, ACC = Exactitud, MCR = Índice de errores y DICE = Coeficiente de DICE.

| CORTE | TC   |      |      | ACC  |      |      | MCR % |      |       | DICE |      |      |
|-------|------|------|------|------|------|------|-------|------|-------|------|------|------|
|       | LCR  | MG   | MB   | LCR  | MG   | MB   | LCR   | MG   | MB    | LCR  | MG   | MB   |
| #59   | 0.93 | 0.88 | 0.89 | 0.93 | 0.94 | 0.95 | 7.26  | 5.53 | 4.78  | 0.96 | 0.94 | 0.94 |
| #79   | 0.88 | 0.83 | 0.87 | 0.92 | 0.95 | 0.88 | 8.09  | 4.58 | 11.31 | 0.94 | 0.91 | 0.93 |
| #99   | 0.92 | 0.85 | 0.87 | 0.95 | 0.94 | 0.91 | 5.24  | 6.38 | 8.76  | 0.96 | 0.92 | 0.93 |

En la Tabla XVI pueden verse las matrices de confusión para los cortes que se han mostrado en las figuras. Como es de esperarse, nuevamente los mayores valores se ubican en la diagonal principal de las matrices.

**Tabla XVI: Matrices de confusión obtenidas en 3 cortes de una imagen real (Conjunto #1) obtenidas con el método propuesto.**

| # CORTE | MATRIZ DE CONFUSIÓN |             |             |
|---------|---------------------|-------------|-------------|
| #59     | <b>0.93</b>         | 0.07        | 0.00        |
|         | 0.00                | <b>0.95</b> | 0.05        |
|         | 0.00                | 0.05        | <b>0.95</b> |
| #79     | <b>0.92</b>         | 0.08        | 0.00        |
|         | 0.03                | <b>0.96</b> | 0.01        |
|         | 0.00                | 0.11        | <b>0.89</b> |
| #99     | <b>0.95</b>         | 0.05        | 0.00        |
|         | 0.03                | <b>0.94</b> | 0.03        |
|         | 0.00                | 0.09        | <b>0.91</b> |

### 5.2.1.2. Comparación con otros métodos de segmentación

Para efectuar una comparación con resultados obtenidos por otros métodos, se probaron diferentes algoritmos de clasificación, algunos en software disponible y otros programados para esta comparación.

Se dan los valores de los TC obtenidos a fines de comparación en la Tabla XVII. Como puede observarse en dicha tabla, el método propuesto supera siempre los demás métodos, mostrando siempre coeficientes de mayor valor. En este caso la mejora se produce en todos los tejidos a clasificar y también en el desempeño promedio, en que se logra una mejora de 28% respecto al método KNN, que es el que mejores resultados proveyó. En la última columna se dan los valores de significancia del test de Wilcoxon para evaluar las diferencias entre cada método y el propuesto. En todos los casos se obtuvo  $p < 0.05$ .

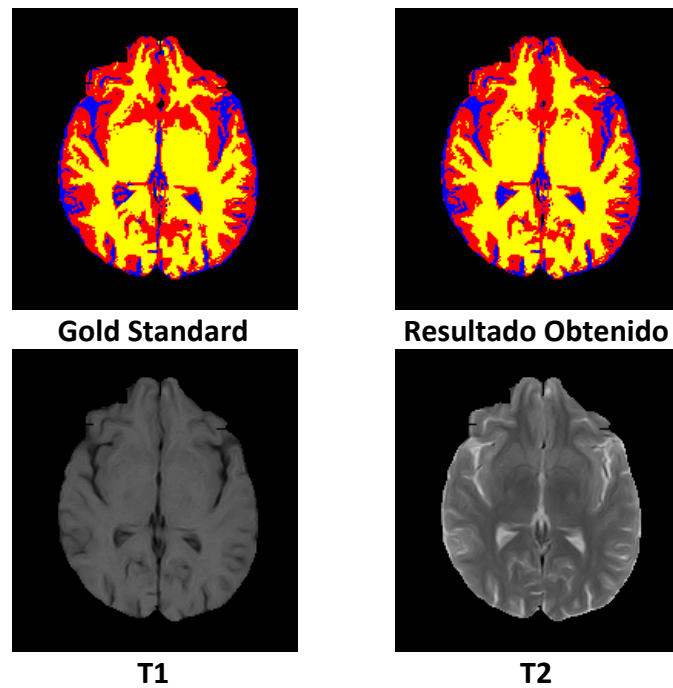
**Tabla XVII: Coeficientes de Tanimoto promedio obtenidos a través de diferentes métodos en un estudio de MR real (Conjunto #1) y significancias obtenidas con el test de Wilcoxon, comparando con los obtenidos mediante el método propuesto (PDLC).**

|             | LCR         | MG          | MB          | Promedio    | p     |
|-------------|-------------|-------------|-------------|-------------|-------|
| <b>BRA</b>  | 0.73        | 0.56        | 0.73        | <b>0.67</b> | 0.004 |
| <b>FCM</b>  | 0.72        | 0.40        | 0.61        | <b>0.58</b> | 0.004 |
| <b>KME</b>  | 0.56        | 0.24        | 0.51        | <b>0.43</b> | 0.004 |
| <b>KNN</b>  | 0.79        | 0.61        | 0.67        | <b>0.69</b> | 0.008 |
| <b>NNE</b>  | 0.72        | 0.54        | 0.64        | <b>0.63</b> | 0.005 |
| <b>PNN</b>  | 0.81        | 0.75        | 0.77        | <b>0.78</b> | 0.012 |
| <b>PDLC</b> | <b>0.91</b> | <b>0.82</b> | <b>0.88</b> | <b>0.88</b> |       |

## 5.2.2. Conjunto de imágenes #2

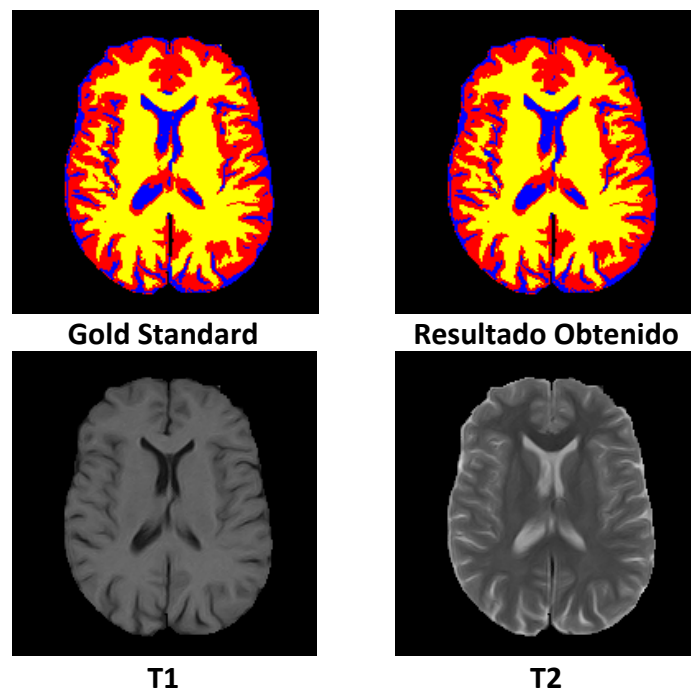
### 5.2.2.1. Mediciones de calidad de la segmentación

El resultado con este conjunto de imágenes puede observarse para algunos cortes desde la Figura 49 hasta la Figura 51. En este caso las imágenes presentaron mayor calidad, lo que facilitó la segmentación, pudiéndose observar gran similitud entre la imagen que se tomó como referencia y la imagen obtenida.



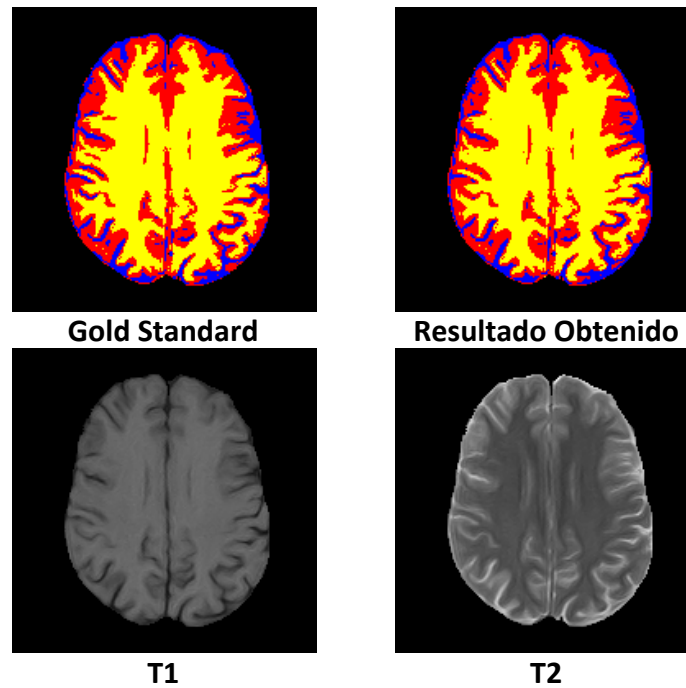
**Figura 49: Corte #80 de una imagen real (Conjunto #2).**

En la fila superior se muestra la imagen de referencia que debería obtenerse (segmentación manual realizada por especialistas) y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes originales luego de haber sido procesadas para remover las regiones de píxeles que no se clasificarán.



**Figura 50: Corte #90 de una imagen real (Conjunto #2).**

En la fila superior se muestra la imagen de referencia que debería obtenerse (segmentación manual realizada por especialistas) y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes originales luego de haber sido procesadas para remover las regiones de píxeles que no se clasificarán.



**Figura 51: Corte #100 de una imagen real (Conjunto #2).**

En la fila superior se muestra la imagen de referencia que debería obtenerse (segmentación manual realizada por especialistas) y la segmentación obtenida con el método propuesto. En la fila inferior se muestran las imágenes originales luego de haber sido procesadas para remover las regiones de píxeles que no se clasificarán.

Las impresiones visuales se pueden constatar a través de los valores de calidad de la segmentación hallados, que se exponen en la Tabla XVIII. Puede observarse que los valores de ACC son muy cercanos a 1, y no ocurre lo mismo con los valores del TC. Esto permite valorar la información que aporta el TC que queda enmascarada en otras mediciones de calidad. Los valores de MCR se mantuvieron siempre menores a 10% (lo que es coherente con una ACC de 90%). Los coeficientes de Dice (DICE) muestran siempre valores superiores a 0.86. El corte #90 presentó más dificultades en la segmentación del LCR, lo que se refleja en el TC (y por consiguiente en DICE, ya que están relacionados) y no en ACC ni en MCR.

**Tabla XVIII: Valores de diferentes mediciones de calidad en 3 cortes de una imagen real (Conjunto #2) obtenidos con el método propuesto.**

| CORTE | TC   |      |      | ACC  |      |      | MCR % |      |      | DICE |      |      |
|-------|------|------|------|------|------|------|-------|------|------|------|------|------|
|       | LCR  | MG   | MB   | LCR  | MG   | MB   | LCR   | MG   | MB   | LCR  | MG   | MB   |
| #80   | 0.94 | 0.93 | 0.99 | 0.94 | 0.94 | 1.00 | 5.79  | 6.22 | 0.00 | 0.96 | 0.96 | 0.99 |
| #90   | 0.75 | 0.89 | 0.99 | 0.99 | 0.90 | 0.99 | 0.12  | 9.84 | 0.00 | 0.86 | 0.94 | 0.99 |
| #100  | 0.83 | 0.96 | 0.99 | 0.94 | 0.99 | 0.99 | 6.17  | 1.14 | 0.00 | 0.91 | 0.98 | 0.99 |

En la Tabla XIX pueden verse las matrices de confusión para los cortes que se han mostrado en las figuras. Los casos en que la suma de valores de la fila no alcanza a ser 1 se deben a que algunos píxeles han sido erróneamente clasificados como fondo. El fondo no fue incluido en las matrices de confusión para no aumentar su dimensión, pues éste no agrega información a la calidad de la segmentación obtenida.

**Tabla XIX: Matrices de confusión obtenidas en 3 cortes de una imagen real (Conjunto #2) obtenidas con el método propuesto.**

| # CORTE | MATRIZ DE CONFUSIÓN |             |             |
|---------|---------------------|-------------|-------------|
| #80     | <b>0.94</b>         | 0.01        | 0.00        |
|         | 0.00                | <b>0.94</b> | 0.00        |
|         | 0.00                | 0.00        | <b>1.00</b> |
| #90     | <b>0.99</b>         | 0.00        | 0.00        |
|         | 0.09                | <b>0.91</b> | 0.01        |
|         | 0.00                | 0.00        | <b>0.99</b> |
| #100    | <b>0.94</b>         | 0.04        | 0.00        |
|         | 0.01                | <b>0.99</b> | 0.03        |
|         | 0.00                | 0.00        | <b>0.99</b> |

#### 5.2.2.2. Comparación con otros métodos de segmentación

La misma comparación efectuada con las imágenes del conjunto anterior fue repetida en estas y se obtuvieron los resultados que se muestran en la Tabla XX. En la última columna se dan los valores de significancia del test de Wilcoxon para evaluar las diferencias entre cada método y el propuesto. En todos los casos se obtuvo  $p < 0.05$ . Se logra una mejora del TC de 6% respecto al FCM, que muestra un buen desempeño y la mejora es del 13% para el caso de los métodos KNN y NNE.

**Tabla XX:** Coeficientes de Tanimoto obtenidos a través de diferentes métodos en un estudio de MR real (Conjunto #2) y significancias obtenidas con el test de Wilcoxon, comparando con los obtenidos mediante el método propuesto (PDLC).

|             | LCR         | MG          | MB          | Promedio    | p     |
|-------------|-------------|-------------|-------------|-------------|-------|
| <b>BRA</b>  | 0.61        | 0.66        | 0.89        | <b>0.72</b> | 0.004 |
| <b>FCM</b>  | 0.78        | 0.86        | 0.93        | <b>0.86</b> | 0.019 |
| <b>KME</b>  | 0.77        | 0.84        | 0.91        | <b>0.84</b> | 0.004 |
| <b>KNN</b>  | 0.76        | 0.75        | 0.86        | <b>0.79</b> | 0.008 |
| <b>NNE</b>  | 0.65        | 0.68        | 0.86        | <b>0.73</b> | 0.004 |
| <b>PNN</b>  | 0.89        | 0.85        | 0.87        | <b>0.78</b> | 0.003 |
| <b>PDLC</b> | <b>0.87</b> | <b>0.93</b> | <b>0.99</b> | <b>0.93</b> |       |

### 5.3. Discusión

Los valores obtenidos para el TC en las imágenes simuladas son altos para todos los tejidos a detectar, aun en las peores condiciones de nivel de ruido y de INU (el mínimo valor es 0.76, obtenido en este caso), lo que permite asegurar en principio que el método funciona con robustez. En el caso de máxima distorsión la segmentación obtenida en imágenes simuladas se degrada, pero la imagen de la que se parte es realmente de bajísima calidad y supera por mucho la situación que suele presentarse en casos reales.

Cuando se probó el método con imágenes simuladas sin distorsiones el TC tomó siempre valores superiores a 0.94. Este resultado es ligeramente superior al reportado en trabajos anteriores en los que se presentan esquemas de procesamiento mucho más costosos computacionalmente o extremadamente complejos para su comprensión.

La utilización de los operadores de la LDC proveyó resultados estables y coherentes según se observó en los gráficos que muestran la disminución del TC cuando se disminuye la calidad de las imágenes a segmentar (aumento del ruido y las INU). Los resultados con los operadores tradicionales MAX-MIN no son adecuados.

Con los algoritmos programados, operando en un equipo con un microprocesador Pentium IV de 3 GHz, mono procesador, con 1 Gbyte de RAM y Disco

Rígido de 120 Gbytes, se observaron los siguientes tiempos de cálculo: 5.37 segundos por corte (16 minutos para un estudio completo) para el FCM, 1.24 segundos (3.72 minutos para un estudio completo) para el KME, 2.12 segundos (6.36 minutos para un estudio completo) para el KNN, 3.45 segundos (10.3 minutos para un estudio completo) para el NNE, 1.63 segundos (4.83 minutos para un estudio completo) para el PNN. El método propuesto requirió 0.96 segundos para procesar un corte, lo que implican 2.52 minutos para un estudio 3D completo.

Los predicados que definen los tejidos a reconocer podrían involucrar otras características más allá de los niveles de gris en las imágenes. Podrían incorporarse, por ejemplo, característicos de texturas tales como rugosidad o inclusive nivel de ruido o cualquiera que pueda ser cuantificable. Pero en este caso debe destacarse que el tiempo insumido en el cálculo del característico necesario para ingresar en el sistema incrementaría (en algunos casos considerablemente) los tiempos de procesamiento. Aun así, hay característicos ampliamente utilizados que se calculan sin gran costo computacional, como por ejemplo el gradiente o las componentes de la transformada de Fourier de una pequeña región de la imagen que caracterizaría a cada píxel.

Dada la relativa simplicidad de las operaciones involucradas, el sistema de procesamiento puede ser programado en cualquier lenguaje sin necesidad de motores de cálculo ni cualquier otra librería específica. Cualquier lenguaje que manipule eficientemente vectores y matrices será adecuado. Dado que no era el objetivo, en esta etapa aún no se ha optimizado el código. Sin embargo sólo con compilarlo en un lenguaje de más bajo nivel que MatLab® se lograría una apreciable reducción de los tiempos de procesamiento. No se deja de tener en cuenta la posibilidad de implementarlo en hardware, por ejemplo a través de software embebido en un microprocesador, de manera de obtener una herramienta en tiempo real.

Las líneas que se presentarán como trabajo futuro constituirán un considerable aporte en la exactitud y/o adaptabilidad del sistema. También brindarán muchas más posibilidades a la hora de elegir un adecuado esquema de procesamiento en problemas de segmentación que pueden expresarse en forma lingüística.

## 6. Conclusiones

## 6. Conclusiones

---

En el procesamiento de imágenes médicas, la segmentación es una etapa sumamente importante para realizar posteriores cuantificaciones. La objetividad y repetitividad del análisis permite determinar patrones de evolución o cambios en los pacientes y dar así un paso hacia la elaboración de guías para la asistencia al diagnóstico temprano de las enfermedades, su tratamiento y su evolución.

También es una etapa relevante para mejorar los informes de rutina y poder discernir entre valores normales y patológicos de las diferentes estructuras presentes.

Gran cantidad de estudios de enfermedades de origen cerebral se basan en la información que se obtiene de neuroimágenes, en particular de RM, y específicamente en la diferenciación de los tejidos y su cuantificación. Esto hace que la tarea de segmentar este tipo de imágenes sea un tema de constante interés y desarrollo.

La segmentación manual es un método que queda en desuso debido a la gran cantidad de tiempo que requiere y al cansancio visual que produce, teniendo en cuenta que un volumen puede involucrar mucho más de cien imágenes. Los métodos de segmentación basados en computadora propuestos van solucionando diferentes aspectos de este tema, como por ejemplo asegurar la repetitividad de los resultados obtenidos, pero ninguno ha mostrado ser el ideal.

Es interesante que, además de un procesamiento matemático, un sistema de segmentación permita de alguna manera la participación de un experto especialista en diagnóstico por imágenes, incluyendo su conocimiento. En este sentido, los métodos basados en sistemas de procesamiento del lenguaje han dado una respuesta.

En esta tesis se presentó como aporte al tema un método para segmentar tejidos en IRM basado en un sistema híbrido compuesto por:

- Un conjunto de predicados que son analizados por LD, aplicando los operadores de la LDC para el cálculo de operaciones difusas.

- Un AG que, basándose en píxeles previamente correctamente segmentados, mejora las definiciones de los conjuntos difusos que son utilizados en los predicados.

Se estableció un procesamiento esencialmente diferente a los enfoques dados por los sistemas de inferencia difusa de Mamdani y/o de Sugeno.

Se propuso un modelo con las siguientes características:

- conserva la **interpretabilidad** del modelo de Mamdani, en el que las reglas pueden interpretarse lingüísticamente. En este caso no se consideran reglas sino predicados o expresiones difusas;
- al mismo tiempo, conserva la **exactitud** del modelo de Sugeno, pues la optimización con AG involucra datos numéricos de entrada y salida.

Aplicando este método de segmentación basado en predicados con valores de verdad difusos se ha logrado discriminar exitosamente la MG, la MB y el LCR en estudios de RM de cerebro.

El sistema propuesto implementa en forma computacional el conocimiento que brindan los expertos a la hora de interpretar las IRM en forma de predicados lógicos compuestos.

Una vez determinados los predicados y optimizado el sistema, la operación es sistemática y objetiva, lo que puede constituir una gran ayuda en centros de diagnóstico donde se procesa gran cantidad de imágenes.

Como las operaciones involucradas son relativamente sencillas, los tiempos de cálculo son cortos con respecto a otros métodos descritos en la bibliografía, lo que hace al método ser altamente adecuado para estudios completos en los que deben segmentarse gran cantidad de imágenes.

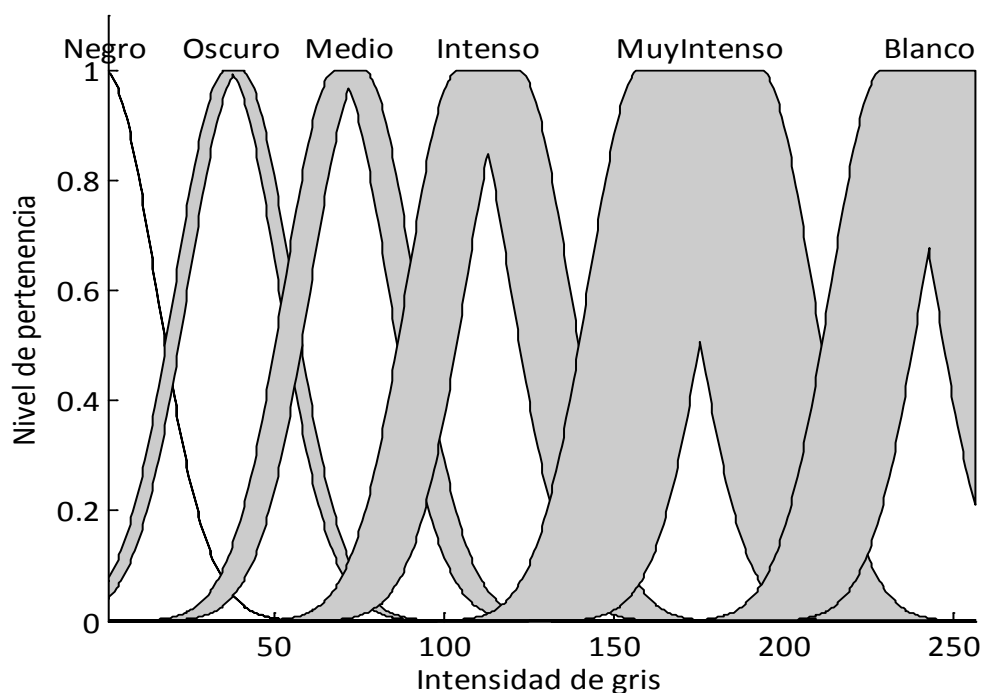
Se ha probado el sistema con imágenes simuladas disponibles en bases de datos en Internet, lo que permitió comprobar un mejor desempeño con respecto a otros algoritmos. Los experimentos realizados con estas imágenes permitieron concluir que el método resulta robusto respecto al ruido y a las INU, dificultades inherentes a las IRM.

Dada la relativa sencillez de los cálculos a efectuar, este sistema puede ser programado en cualquier lenguaje sin requerir grandes motores de cálculo ni librerías especiales. Cualquier lenguaje que pueda trabajar eficientemente con vectores resulta adecuado.

Como **línea de trabajo futuro**, se sugiere evaluar exhaustivamente las posibles mejoras que pueden ofrecer los **conjuntos difusos de tipo 2** (Ver *Apéndice C*). Estos conjuntos pueden agregar incertidumbre en la definición de las funciones de pertenencia, lo que aumenta los grados de libertad y brinda nuevas posibilidades de soluciones satisfactorias en imágenes complejas. Pueden aplicarse estos conjuntos difusos al cuantificar el grado de verdad de los predicados.

Este enfoque parece acertado en esta aplicación en la que las definiciones de las funciones de pertenencia iniciales surgen de los histogramas de intensidades de gris para las distintas sustancias en las distintas imágenes. Como se recomienda en la bibliografía, para diseñar un conjunto difuso tipo 2 se comienza con la región de incertidumbre y a partir de ella se determinan las funciones de pertenencia superior e inferior. En este sistema pueden definirse preliminarmente, confiando en el AG la posterior optimización.

La Figura 52 muestra una posible definición inicial de los conjuntos difusos tipo 2. Cada función de pertenencia requiere más parámetros que los utilizados en los conjuntos tipo 1. Estos parámetros deben ser optimizados por el AG: cada conjunto difuso se determinará por dos funciones gaussianas combinadas, cada una con su centro y ambas con la misma dispersión. Quedan definidas regiones que caracterizan la incertidumbre en la determinación de valores de verdad.



**Figura 52: Funciones de pertenencia tipo 2 preliminares, a ser posteriormente optimizadas.**

Estas funciones de pertenencia agregan grados de libertad al Algoritmo Genético para su optimización y pueden ser una opción para casos de segmentación de imágenes complejas.

Entonces los conjuntos difusos tipo 2 proveen más grados de libertad para el AG, lo que en principio constituye una potencial mejora en los resultados obtenidos. Sin embargo, hasta el presente no hay una teoría cerrada que asegure que esto ocurrirá en todos los casos [Mendel, 2003].

El grado de verdad para cada predicado simple no será un único valor sino un intervalo de valores. Las operaciones lógicas entre valores de verdad numéricos tendrán que reemplazarse entonces por las definidas para la LD de tipo 2 por intervalos.

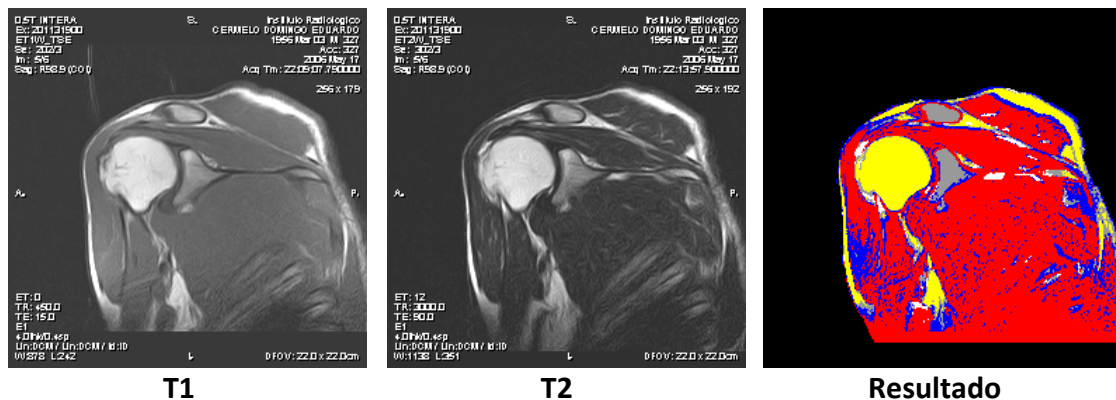
Queda pendiente también la exploración del método presentado para ser utilizado en **otros tipos de imágenes**, basándose en otras variables que no necesariamente serían las intensidades. El modelo de procesamiento propuesto, resulta prometedor para aplicaciones en otros tipos de RM con una mayor cantidad de sustancias a reconocer. También sería adecuado para procesar un conjunto de características extraídas de una imagen para su segmentación, como por ejemplo las derivadas del análisis de texturas.

Tal como se hizo con las imágenes de cerebro, se obtuvieron resultados preliminares utilizando el sistema con médicos especialistas en IRM de hombro. En este caso los cortes son anatómicamente más complejos para su comprensión. Sin llegar a segmentar una imagen de prueba completa, se le pidió a los especialistas que delimitaran manualmente algunas regiones de los tejidos a detectar. En este caso se procuró reconocer: hueso cortical (gris), músculo (rojo), grasa (amarillo), tendón (azul), líquido (blanco) y fondo (negro).

Los predicados utilizados fueron los siguientes:

- P1 = “El **Hueso** es claro en T1 y en T2”
- P2 = “El **Músculo** es gris medio en T1 y muy gris en T2”
- P3 = “La **Grasa** es muy claro en T1 y en T2”
- P4 = “El **Tendón** es claro en T1 y oscuro en T2”
- P5 = “El **Líquido** es oscuro en T1 y muy claro en T2”
- P6 = “El **Fondo** es muy oscuro en T1 y en T2”

Se procedió de la misma manera que con las imágenes de cerebro, sin involucrar a PD, secuencia de la cual no se tienen imágenes en el caso de hombros. Los resultados preliminares se muestran en la Figura 53.



**Figura 53: Resultado preliminar para una Resonancia Magnética de hombro.**

Los resultados no son concluyentes pero confirman que el método propuesto es factible de ser adaptado para otro tipo de imágenes.

Finalmente, como otra alternativa de mejora en la segmentación, se planea ampliar los alcances de este método para la **segmentación de volúmenes completos** considerando la relación entre los sucesivos cortes del estudio 3D y no sólo la segmentación de cada corte por separado.

La utilización de expresiones difusas basadas en el conocimiento de especialistas aplicadas al problema de la clasificación constituye un interesante enfoque en la segmentación de imágenes. Las propiedades de la Lógica Difusa Compensatoria, ya demostradas en otros ámbitos, y los resultados con ella obtenidos han justificado su utilización en este contexto.

El método propuesto constituye un enfoque simple para representar objetivamente la manera en que los especialistas interpretan las imágenes. Su utilización en el procesamiento automatizado lo convierte en una herramienta útil. Se destacan sus potencialidades en cuanto a su adaptabilidad para diversos tipos de imágenes.

Los resultados presentados permiten concluir que el método es robusto, no requiere pre-procesamientos ni complejos algoritmos matemáticos y por lo tanto no demanda un equipo informático exigente para poder implementarse.

En el análisis de imágenes médicas, la segmentación es un requerimiento siempre vigente, considerando que la tecnología para la obtención de las mismas varía constantemente. Todo nuevo avance en este sentido constituye un aporte para el análisis y la sistematización del descubrimiento del conocimiento que las imágenes médicas contienen. Estos avances ofrecen mejoras en el estudio, el diagnóstico y el seguimiento de diversas patologías, lo que redunda directa o indirectamente en una mejor calidad de vida que todos merecemos.

---

# Apéndice A: Algoritmos Genéticos

# Apéndice A: Algoritmos Genéticos

---

## 6.1. Introducción

---

Un problema que presentan algunos métodos de optimización, en los cuales se intenta hallar un conjunto de parámetros tales que produzcan un mínimo global de una determinada función de evaluación (o función de costo), tales como el algoritmo de *Back-propagation* (usado en el entrenamiento de RN) o la optimización por mínimos cuadrados es que al transcurrir algunas iteraciones pueden converger a mínimos locales de la función de evaluación. Las funciones de costo suelen ser altamente alineales y esta característica hace que no sean adecuadas para los métodos basados en derivadas, como los recién citados.

Los AG no requieren derivadas en los cálculos realizados en las iteraciones y son un método de optimización estocástico, lo que hace que sea menos probable que la solución quede atrapada en un mínimo local.

En el contexto de la Inteligencia Computacional, los AG se han utilizado con éxito para optimizar la estructura y los parámetros de RN [Pedrycz et al., 2006]. También se han aplicado para optimizar la forma de las funciones de pertenencia que se utilizan en la base de reglas de sistemas de inferencia difusa [Herrera, 2005].

Los AG imitan la evolución de generaciones como fuera planteado en la teoría de la evolución de Darwin.

Se consideran problemas cuyo objetivo es la minimización de alguna función de costo (o función de evaluación, *fitness function*) paramétrica. La solución estará dada por el conjunto de parámetros que hacen que esa función tome el mínimo valor posible (mínimo global), evitando que en las iteraciones la búsqueda de la solución quede atrapada en mínimos locales.

Las posibles soluciones se indicarán como vectores del espacio de parámetros. A cada posible solución se la llamará “individuo” que formará parte de un conjunto de

soluciones, que será la “población” o “generación”. La población tendrá una determinada cantidad de individuos que deberá indicarse.

La bondad de cada una de las posibles soluciones (el desempeño o *performance* de cada individuo) se medirá según el valor que entregue la función de evaluación para esa solución.

A continuación se resume el funcionamiento básico de un AG:

- En una etapa inicial, se generan soluciones al problema sugeridas al azar, tantas como individuos se indicó tenga la población. Cada solución será, como se explicó, un vector de parámetros.
- A todas esas soluciones se les mide su desempeño, es decir, cuán buena es cada una de ellas.
- Se selecciona una parte de esas mejores soluciones y las otras son eliminadas (sobreviven solamente las mejores). Este proceso será debidamente explicado a continuación.
- Las soluciones seleccionadas se mezclan a través de diferentes procesos (reproducción, cruzamiento, mutación) y así se genera una nueva generación de posibles soluciones, que se espera sea mejor que la generación anterior.
- Se generan así nuevas poblaciones hasta que se logre la convergencia o hasta que se cumpla algún criterio de detención indicado previamente.

Estos algoritmos buscan la solución en todo el espacio de posibles soluciones. El costo computacional suele ser alto y es dependiente de la complejidad de la función de evaluación y de la cantidad de parámetros a optimizar.

En las secciones siguientes se explican conceptos específicos del contexto de AG. Cuando se hace referencia a parámetros de configuración del algoritmo, se especifican los nombres utilizados por el programa MatLab® R2006b.

## 6.2. Nociones de Algoritmos Genéticos

### 6.2.1. Codificación

El conjunto de parámetros se representa como una cadena de bits, que se denomina *cromosoma* o *individuo*. Cada elemento del cromosoma, compuesto por varios bits, se denomina *gen*. Por ejemplo, el punto del plano bidimensional (13, 3) podría representarse como se muestra en la Figura 54.

|   |   |   |   |   |   |   |   |
|---|---|---|---|---|---|---|---|
| 1 | 1 | 0 | 1 | 0 | 0 | 1 | 1 |
|---|---|---|---|---|---|---|---|

**Figura 54: Posible representación del cromosoma que identifica los números 13 y 3 en un Algoritmo Genético.**

Este caso se compone de dos genes compuestos por 4 bits. Los primeros cuatro bits corresponden al número 13 y los restantes cuatro bits representan al número 3.

En este caso cada gen estaría compuesto por 4 bits. Los primeros cuatro bits corresponden al número 13 y los restantes cuatro bits al número 3.

Se utilizan otros paradigmas de codificación según se trabaje con números enteros, negativos, de punto flotante, etc.

### 6.2.2. Población Inicial

El AG comienza creando una población inicial aleatoria. La población tiene tantas posibles soluciones como individuos se haya indicado (parámetro “*population size*”).

Si se tiene alguna información preliminar del problema o de la función de evaluación, puede indicarse el rango de valores que los parámetros tomarán durante la búsqueda, a fin de acelerar la convergencia. Pero esto no es un requisito para que el algoritmo entregue una solución adecuada.

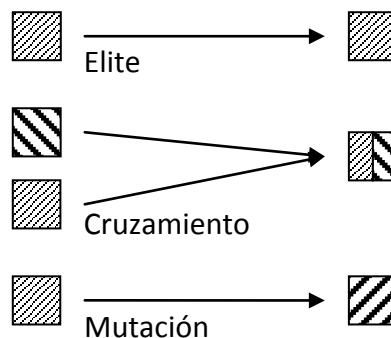
También es posible determinar una población inicial basada en cierto conjunto de parámetros que no es el óptimo pero que puede ser un buen comienzo para la búsqueda estocástica.

### 6.2.3. Creación de una nueva generación

En cada iteración, el AG usa la población actual para crear los *hijos* que compondrán la nueva generación. Se selecciona un grupo de individuos de la población actual, que toman el nombre de *padres*. Los *padres* aportarán sus genes para formar sus *hijos*. Generalmente se seleccionan como *padres* a los individuos que presentan el mejor valor de la función de optimización, pero este criterio puede ser diferente (parámetro “*selection function*”).

El algoritmo puede crear *hijos* por tres metodologías diferentes (Figura 55):

- **Elite:** son los individuos de la generación actual con mejores valores de la función de evaluación, los cuales automáticamente “sobrevivirán” y pasarán a la siguiente generación.
- **Cruzamiento** (*crossover*): genera los hijos combinando los genes de un par de *padres*. Una forma consiste en armar el nuevo individuo eligiendo aleatoriamente, gen a gen, el que corresponde al *padre* o a la *madre*.
- **Mutación:** genera los hijos introduciendo cambios aleatorios o mutaciones en los genes de un *padre* único.



**Figura 55: Diferentes métodos de generación de hijos para una nueva población en un Algoritmo Genético.**

Los individuos de la generación actual pueden pasar directamente a la siguiente generación, combinarse o mutarse para originar la siguiente generación.

### 6.2.4. Detención del Algoritmo

El algoritmo puede detenerse por diversos motivos, que deben configurarse:

- **Cantidad de generaciones:** se detiene cuando se ha alcanzado un número prefijado de generaciones (parámetro “*generations*”).
- **Límite de tiempo:** se detiene cuando ha transcurrido una determinada cantidad de tiempo de cálculo (parámetro “*time limit*”, en segundos).
- **Valor de la función de evaluación:** se detiene cuando el valor de la función de evaluación para el mejor individuo de la población actual es menor que un determinado valor (parámetro “*fitness limit*”).
- **Mejora relativa entre generaciones:** se detiene cuando el cambio del valor promedio de la función de evaluación de una población a otra es menor que un cierto límite (La condición es que el cociente entre el valor promedio y el valor de “*stall generations*” sea menor que el parámetro “*function tolerance*”).
- **Mejora relativa a lo largo del tiempo:** se detiene si no hay mejoras significativas del valor promedio de la función de evaluación durante un determinado tiempo de cálculo (parámetro “*stall time limit*”).

### 6.2.5. Diversidad de la población

Es uno de los factores más importantes que determinan el desempeño de un AG. Se espera que en las sucesivas iteraciones el promedio del valor de la función de evaluación para las poblaciones tenga una tendencia marcada a disminuir, hasta llegar a un valor adecuado.

Si la distancia promedio entre individuos es alta, la diversidad será alta; si la distancia promedio es pequeña, la diversidad es baja. La elección adecuada de la diversidad de la población se efectúa básicamente por prueba y error.

La diversidad de la población puede controlarse básicamente mediante la adecuada elección del rango de valores que tomarán los parámetros de la población inicial. La cantidad de mutación que se permite y la cantidad de individuos que componen la población también afectan a la diversidad.

### 6.2.6. Aptitud de los individuos

En este proceso se convierten los valores de la función de evaluación obtenidos para cada individuo a un rango adecuado para una función de selección (proceso posterior), que seleccionará los padres para una nueva generación. Esta función deberá asignar una alta probabilidad de selección a los individuos con mayores valores en la nueva escala. El valor asignado a cada individuo se denomina “aptitud”.

El rango de los valores de aptitud obtenidos es determinante del modo en que el algoritmo evoluciona. Si los valores varían demasiado, los individuos con mayores valores siempre serán seleccionados y se reproducirán siempre, repitiendo los mismos genes una y otra vez, lo que puede hacer que el algoritmo no busque la solución en otros lugares del espacio de soluciones. Por el contrario, si los valores de aptitud varían poco, todos los individuos tendrían la misma probabilidad de reproducirse y la búsqueda evolucionaría muy lentamente.

Una forma común de proceder a la asignación de aptitud es basándose en el ranking o posición que ocupa cada individuo según su valor de función de evaluación. El ranking del mejor individuo es 1, el siguiente es 2 y así sucesivamente. Entonces:

- El valor de la función escalada para un individuo con ranking  $N$  es proporcional al valor  $1/\sqrt{N}$ .
- La suma de los valores de aptitud sobre la población entera es igual a la cantidad de padres que se requieren para crear la siguiente generación.

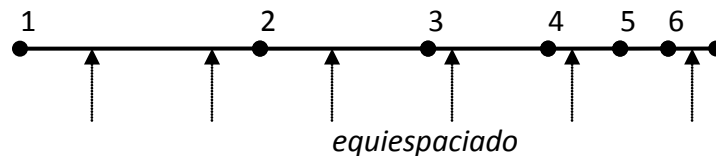
### 6.2.7. Selección

Luego de evaluar todos los individuos de una población, una función de selección elige los padres que producirán la siguiente generación. Un individuo puede ser elegido como padre más de una vez, en cuyo caso sus genes contribuirán a formar más de un hijo. Hay varias opciones para el proceso de selección.

Por ejemplo, en la **selección estocástica uniforme** se traza una línea en la cual cada padre corresponde a una sección de la línea proporcional a su valor de aptitud. El algoritmo utilizado se mueve a lo largo de la línea en pasos de igual tamaño y en el lugar donde se producen los pasos se eligen los padres. Como ejemplo, si se requieren

6 padres para crear la nueva población, los valores podrían ser 2.2, 1.4, 1.1, 0.6, 0.4, 0.3, cuya representación en una recta sería como se muestra en la Figura 56.

El individuo 1 sería elegido dos veces, los individuos 2, 3, 4 y 6 serán elegidos una única vez y el individuo 5 no será elegido nunca.



**Figura 56: Selección estocástica uniforme.**

Representación de los valores de aptitud para una población de 6 individuos y selección de padres. El individuo 1 sería elegido dos veces, los individuos 2, 3, 4 y 6 serán elegidos una única vez y el individuo 5 no será elegido nunca.

Otra posibilidad para el proceso de selección, más determinística, se basa en los valores de aptitud (función de selección “*remainder*”). Se tienen en cuenta la parte entera y fraccionaria de los mismos. Con la parte entera se asigna la cantidad de veces que el individuo será elegido como padre. Con la parte fraccionaria se efectúa un proceso de selección estocástico uniforme como el explicado en la sección anterior.

Por ejemplo, si un individuo tiene valor 2.4, la función selecciona a ese individuo 2 veces como padre. Luego de que los individuos han sido así elegidos, el resto de los padres se eligen estocásticamente: la probabilidad de ser elegido como padre es proporcional a la parte fraccionaria de su valor de aptitud, siguiendo el proceso de selección estocástico uniforme.

### 6.2.8. Opciones para la reproducción

Para determinar el proceso de reproducción se debe indicar:

- Cuántos individuos (los que mejor valor de función de evaluación presenten) pasarán directamente a la siguiente generación (parámetro “*population size*”). Generalmente son unos pocos. Si fueran demasiados, entonces se corre el riesgo de preservar los mismos individuos por varias generaciones, lo que haría disminuir la diversidad de la población, y por tanto la búsqueda sería menos efectiva.
- La fracción de individuos de la siguiente generación que serán creados por cruzamiento (parámetro “*crossover fraction*”). Si es 1, entonces todos los

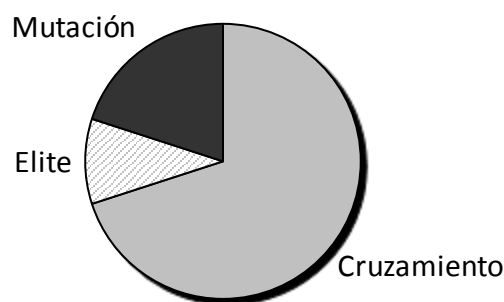
individuos de no elite serán creados por cruzamiento. Si es 0, entonces todos los individuos de no elite serán creados por mutación.

Por ejemplo, si en MatLab® se determinan los siguientes parámetros:

- Population size = 20
- Elite count = 2
- Crossover fraction = 0.8

la próxima generación tendrá las siguientes características:

- Tiene 20 individuos en total.
- 2 de ellos pasarán directamente de una generación a la otra.
- 18 individuos no serán de elite, entonces el algoritmo redondeará:  $0.8 \times 18 = 14.4$  a 14 para hallar la cantidad de individuos creados por cruzamiento, y el resto (4 individuos) serán creados por mutación. La proporción de individuos creados en la nueva generación se representa gráficamente en la Figura 57.



**Figura 57: Posible configuración de la proporción de individuos en una nueva generación generados por diferentes métodos.**

Se muestra el caso en que en MatLab® se indica *Population size=20*, *Elite count=2*; *Crossover fraction=0.8*.

### 6.2.9. Mutación

El efecto principal de la mutación es asegurar la diversidad de la población. Generalmente se configura con un valor inicial y luego la probabilidad de mutación va disminuyendo conforme avanzan las iteraciones.

Existen diversas funciones de mutación. Por ejemplo, la función de mutación gaussiana, en la cual se agrega un valor extraído de una distribución gaussiana a cada elemento del individuo.

El efecto de la mutación se ve reflejado en el valor de la distancia vectorial promedio entre individuos conforme avanzan las iteraciones. Si el valor elegido para la tasa con la que la mutación decrece (parámetro “*shrink*”) es alto, cercano a 1, entonces los individuos serán cada vez más similares entre sí y se apreciará una notoria disminución de la distancia entre ellos. Esto no se observa con valores menores.

#### 6.2.9.1. Caso especial – Cruzamiento sin mutación

En este caso, el algoritmo tomaría los genes de los individuos de la población inicial y simplemente los recombinaría. No se crearían nuevos genes al no haber mutación.

Muy probablemente se llegaría a un valor óptimo para una población inicial dada y los genes del mejor individuo se recombinarían entre sí.

El algoritmo se detendría al no percibir mejoras en la función de evaluación del mejor individuo. El valor medio de la función de evaluación para las sucesivas poblaciones tenderá al valor final del mejor individuo y ya no cambiará.

#### 6.2.9.2. Caso especial – Mutación sin cruzamiento

En este caso, el algoritmo nunca mejora el valor obtenido por el mejor individuo de la primera generación, pues sus genes no son utilizados para obtener nuevos individuos (al no haber cruzamiento).

Como en el caso anterior, el algoritmo se detendría al no percibir mejoras en la función de evaluación. El valor medio de la función de evaluación para las sucesivas poblaciones será errático, aunque tienda a disminuir.

La Figura 58 es una representación gráfica resumida del proceso iterativo de un AG con todos los pasos que intervienen.

## 6.3. Aplicación

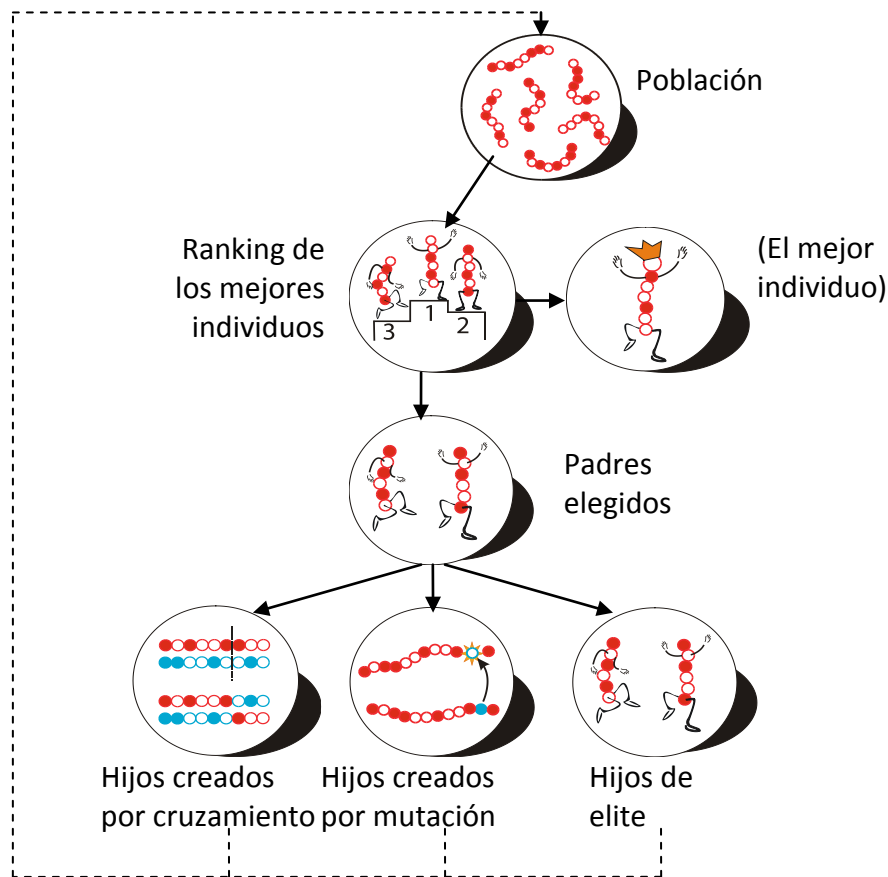
---

El problema quedará definido al indicar, como mínimo:

- La cantidad de parámetros que se desea obtener.

- El cálculo involucrado para hallar el valor de la función de evaluación (dependiente, por supuesto, de los parámetros a hallar).

Los AG se han utilizado con éxito en forma conjunta con otras herramientas en tareas de extracción del conocimiento.



**Figura 58: Representación esquemática que resume el funcionamiento iterativo de un Algoritmo Genético.**

De la población inicial se eligen padres para crear la nueva población, la que es nuevamente considerada como inicial para originar la siguiente. Cuando termina el algoritmo se elige el mejor individuo como solución.

Se ha mencionado en ocasiones que el uso de este tipo de paradigmas en el descubrimiento de reglas tiende a producir resultados con capacidad de generalización baja. No obstante, la creciente cantidad de trabajos publicados demuestra que en el contexto de la Inteligencia Computacional tienen un gran potencial aún por abordar.

Los AG realizan una búsqueda global e independiente del dominio, lo que los convierte en una herramienta robusta y aplicable en distintas etapas de procesos de extracción de conocimiento.

En la sección 4.5 se utiliza un AG para la optimización de funciones de pertenencia y su función de evaluación estará indicada como un predicado del cual se intenta maximizar su valor de verdad.

---

# Apéndice B: Interfaz gráfica de desarrollo

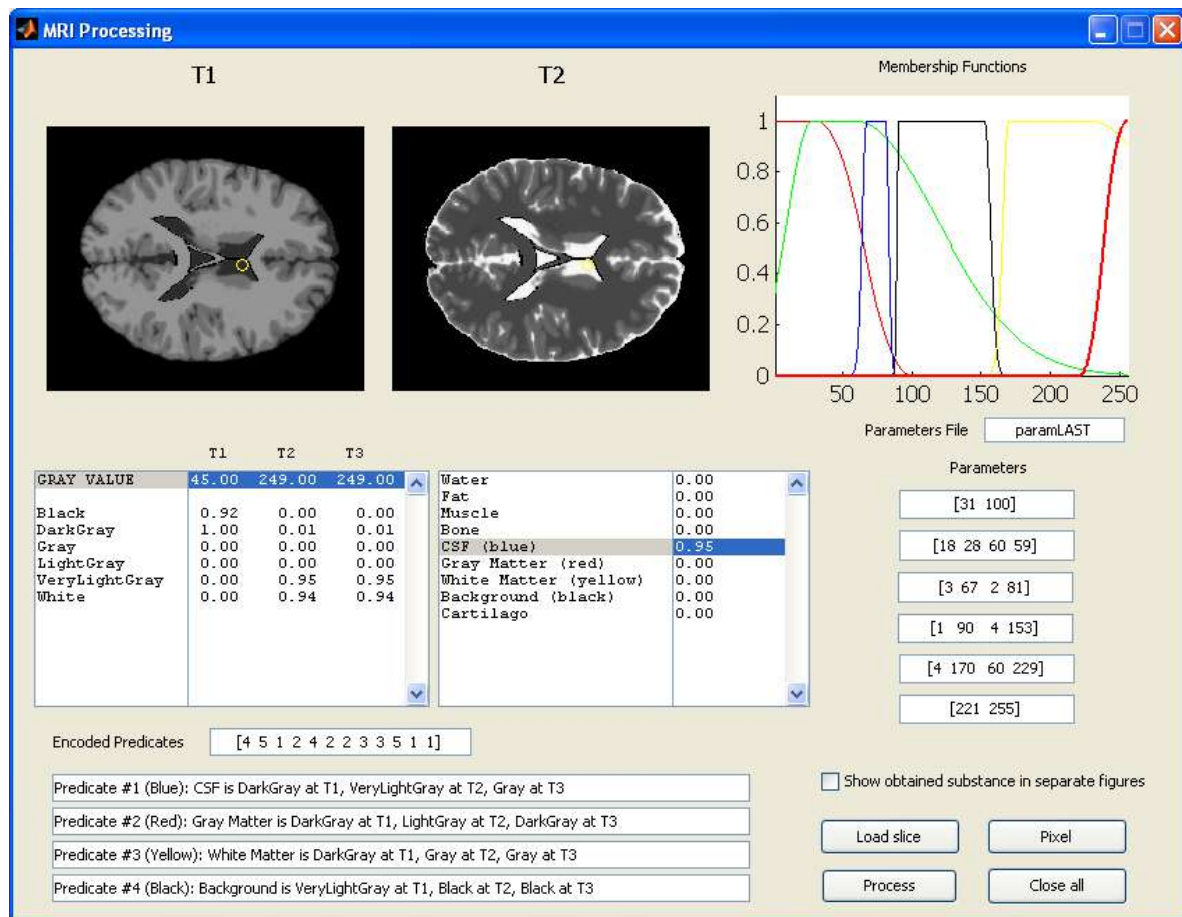
## Apéndice B: Interfaz gráfica de desarrollo

---

En la etapa de prueba se diseñó, en el entorno de MatLab® 2006b, una interfaz gráfica (Figura 59) para testear, eligiendo un píxel por un operador: los valores de gris, las pertenencias a los diferentes conjuntos difusos y los valores de verdad de los predicados que definen a cada sustancia. Luego se generalizó el proceso para una imagen completa.

Esta interfaz permitió probar el sistema de forma visual y operativa, a modo de prototipo. Fue diseñada en forma general para poder ser ampliada con facilidad, aunque incluye elementos específicos para el procesamiento de IRM de cerebro. Con el propósito de describir su funcionalidad, se enuncian algunas de sus prestaciones que fueron sumamente valiosas en la etapa de investigación y desarrollo.

En la interfaz gráfica (Figura 59), a la derecha, en la parte superior, se muestran las funciones de pertenencia utilizadas para el procesamiento. Los parámetros de estas funciones se almacenan en un archivo cuyo nombre por defecto es paramLAST.mat. El nombre de este u otro archivo es el que se muestra en el cuadro de texto que se indica como *"Parameters File"*. Debajo de éste se encuentran los parámetros obtenidos para cada una de las funciones de pertenencia de los conjuntos difusos. Si se quiere probar con otros parámetros diferentes, almacenados en otro archivo, se cambia este nombre y automáticamente se actualizan las gráficas de las funciones de pertenencia con los nuevos parámetros.



**Figura 59: Interfaz gráfica utilizada para pruebas.**

La interfaz mostrada fue utilizada en el diseño y prueba del método propuesto. Permite, entre otras funciones, evaluar los valores obtenidos de pertenencias a todos los conjuntos difusos para un píxel indicado.

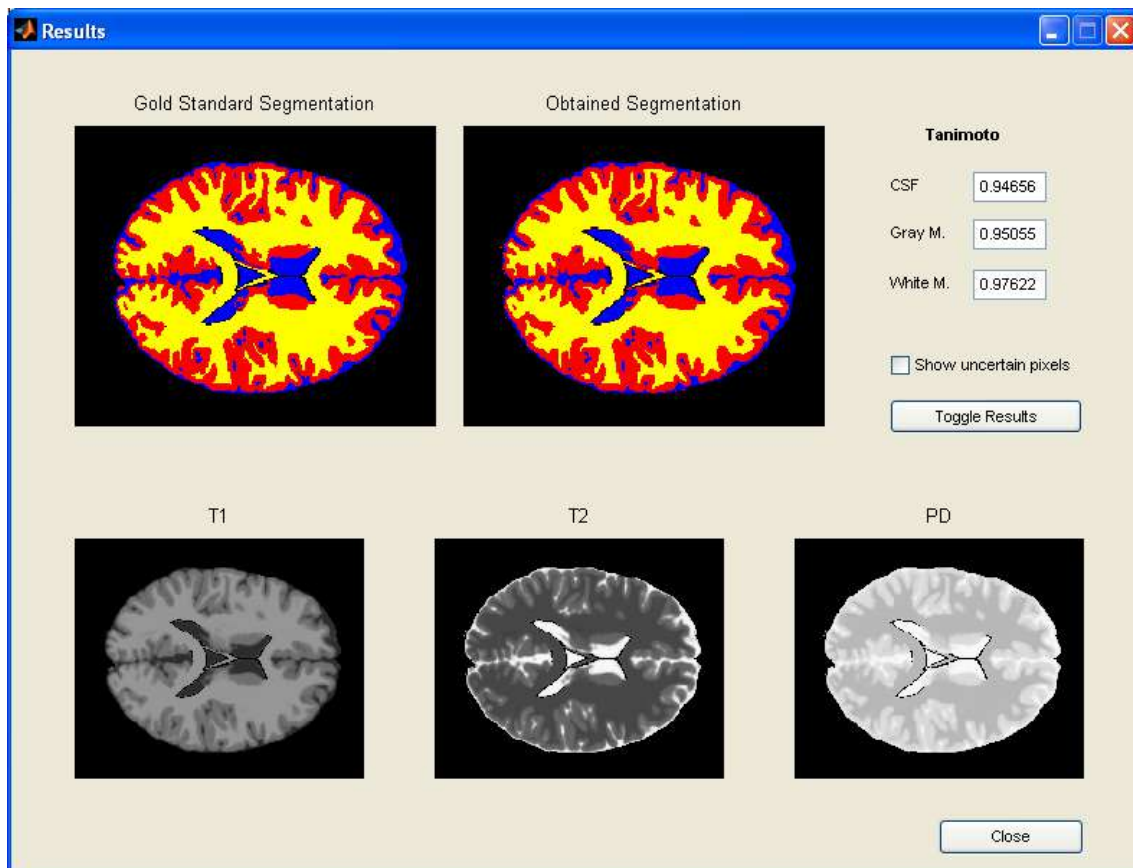
Cuando se pulsa el botón “Pixel” es posible seleccionar un píxel de la imagen, que por supuesto coincide en ubicación en las imágenes T1 y T2, mostradas en la parte superior de la pantalla. Una vez seleccionado:

- En la tabla de la izquierda se muestran los valores de gris del píxel seleccionado para las tres imágenes que intervienen en el procesamiento (que se han indicado como T1, T2 y PD). Debajo de ellos puede verse la pertenencia de esos valores de gris a cada conjunto difuso que ha sido definido.
- En la tabla de la derecha se muestran los valores de verdad difusos de los predicados que definen diferentes tejidos. Si no han sido definidos predicados para determinados tejidos, los valores serán siempre 0.

A través del botón “Load Slice” puede cargarse una nueva imagen que se mostrará inmediatamente reemplazando la anterior, a fines de su procesamiento.

Cuando se pulsa el botón “Process” se procesa la imagen completa con los parámetros mostrados. Luego se muestra una nueva ventana con todas las imágenes procesadas y los resultados obtenidos (Figura 60):

- Se observa la segmentación utilizada como *Gold Standard* y la segmentación obtenida con el método propuesto. Las imágenes mostradas en esta figura se han incluido a modo de ejemplo, dado que en el capítulo 5 se analizarán cuantitativamente los resultados.



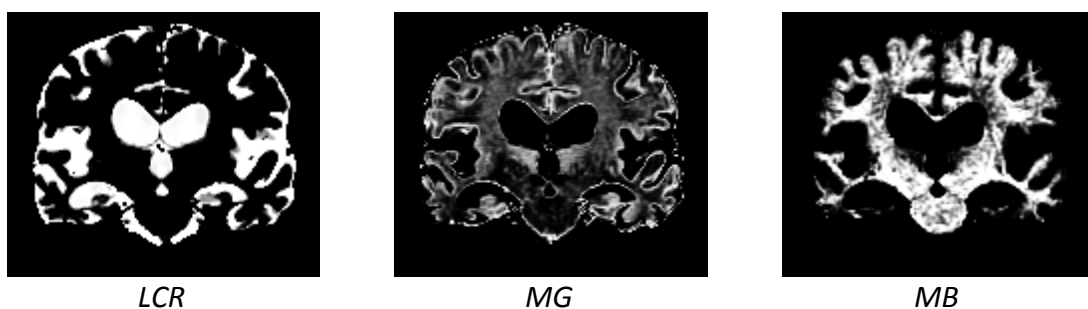
**Figura 60: Interfaz gráfica para visualización de resultados.**

Esta pantalla se despliega para mostrar las imágenes que han sido procesadas, la segmentación de referencia y la segmentación obtenida y los Coeficientes de Tanimoto resultantes.

- A la derecha se observan los Coeficientes de Tanimoto calculados para cada uno de los tejidos. Si la casilla “Show obtained substance in separate figures” se encuentra seleccionada, entonces también se mostrarán imágenes para cada una de los tejidos por separado. En este caso se pueden observar los valores de verdad de los diferentes predicados para cada píxel, representados en el rango de grises (Figura 61) y desplazar el cursor por sobre la imagen para observar el

valor obtenido. También se muestran en imágenes separadas los píxeles que han sido asignados como criterio final a cada uno de los tejidos (Figura 62).

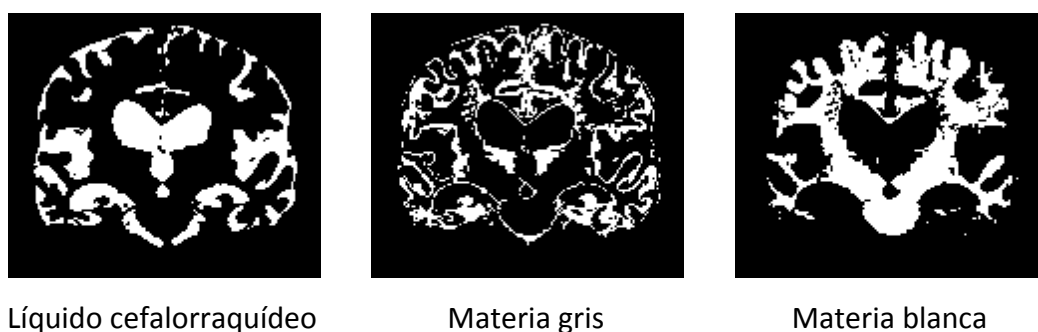
- El botón *“Toggle Results”* cambia alternativamente la imagen original en el lugar de la obtenida para comparar de una manera visual las diferencias que presentan.
- Si se selecciona la casilla *“Show uncertain pixels”* se muestran en color gris los píxeles que presentaron valores de verdad menores a un cierto umbral (0.5 por defecto) en la totalidad de los predicados involucrados.



**Figura 61: Representación de los valores de verdad obtenidos.**

Imágenes separadas que muestran los valores de verdad obtenidos píxel a píxel para los tres tejidos a detectar. Los valores extremos se observan como blanco (valor de verdad 1) y negro (valor de verdad 0). Los valores intermedios se representan con diferentes intensidades de grises.

Los elementos de la región inferior izquierda de la interfaz (Figura 59) no se utilizan en esta etapa. Se aplicaría en caso de probar diferentes opciones en la automatización de la definición de los predicados involucrados en la detección de sustancias.



**Figura 62: Imágenes de los tejidos asignados a cada píxel.**

Luego de proceder a la decisión, cada píxel es asignado al tejido cuyo predicado posee el mayor grado de verdad. Las imágenes binarias muestran en blanco los píxeles asignados a los tres tejidos por separado.

## Apéndice C: Conjuntos Difusos Tipo 2

# Apéndice C: Conjuntos Difusos

## Tipo 2

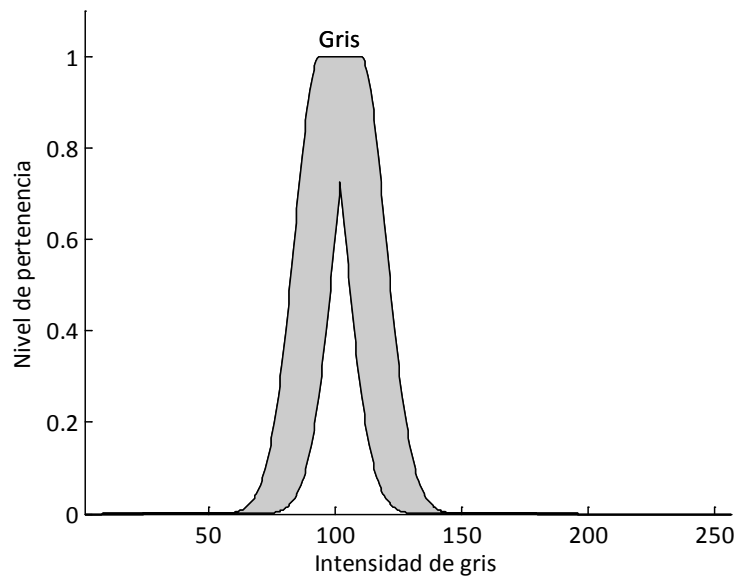
---

Cuando no es posible determinar los grados de pertenencia como 0 (no pertenece) o 1 (pertenece) se utilizan los conjuntos difusos definidos en la Sección 3.3.2, conocidos como “de tipo 1”, que permiten que sus elementos tengan grados de pertenencia continuos en  $[0, 1]$ .

Si la definición de un concepto es tan difusa que ni siquiera es posible determinar la pertenencia como un elemento único del intervalo  $[0, 1]$ , entonces un posible enfoque o metodología adecuada podría ser la utilización de conjuntos difusos tipo 2 [Mendel, 2001, Mendel, 2003].

Los conjuntos difusos tipo 2 pueden pensarse para su representación como un conjunto tipo 1 el cual no tiene un límite preciso en su definición (no alcanza una única curva para ser representado). Por ejemplo, si una curva de forma gaussiana se generaliza para representar un conjunto difuso tipo 2, se lograría el conjunto que se observa en la Figura 63, aplicado a la intensidad de gris presente en un píxel. Los conjuntos tipo 2 presentan grados de pertenencia que son ellos mismos también difusos.

Entonces, para un valor de intensidad ya no se tiene un valor único de grado de pertenencia, sino un intervalo de valores. Estos valores podrían no estar pesados de la misma manera, lo que daría un valor difuso de pertenencia a un conjunto difuso. Sin embargo, la teoría de conjuntos difusos tipo 2 ha sido reducida y simplificada para el caso en que todos estos valores de “pertenencia a la pertenencia” pesan todos 1. Así surge la LD tipo 2 **de intervalos** (*Interval Type-2 Fuzzy Logic*). El concepto de conjuntos difusos de tipo 2 fue introducido primero por Zadeh en 1975 [Zadeh, 1975] y su aplicación en sistemas de inferencia fue presentada en 1998 y publicada en 1999 [Karnik et al., 1999].



**Figura 63: Representación gráfica de un conjunto difuso de tipo 2.**

Para un valor de intensidad determinado ya no se tiene un valor único de grado de pertenencia, sino un intervalo de valores.

Hay una extensa nueva terminología para referirse a conjuntos difusos tipo 2. Se denomina **FOU** (*Footprint of Uncertainty*, zona de incertidumbre) a la región comprendida entre los límites del conjunto difuso. Se llama **función de pertenencia superior** al límite superior de esta región y **función de pertenencia inferior** al límite inferior (ver Figura 63) [Karnik and Mendel, 2001, Mendel, 2007a, Mendel et al., 2006].

Los resultados obtenidos con este tipo de lógica han demostrado ser prometedores tanto en Sistemas de Inferencia como en sistemas basados en expresiones lingüísticas [Mendel, 2007b]. Si bien no se han utilizado conjuntos difusos tipo 2 en el método presentado en esta tesis, se presenta su utilización como una interesante línea de trabajo futuro.

# Agradecimientos

**Gustavo Javier Meschino**

Modelos híbridos de Inteligencia Computacional aplicados  
en la Segmentación de Imágenes de Resonancia Magnética

# Agradecimientos

---

Quiero agradecer a gran cantidad de personas:

A Emilce Moler, quien me dio la oportunidad de iniciarme en el mundo de escribir trabajos científicos y quien valoró mis pequeños aportes con gran estímulo, mostrándome con su vida que todo puede contribuir a ayudar a los demás.

A Virginia Ballarin, con quien, además de trabajar muy duro en el desarrollo concreto de este y otros trabajos, vivimos la vida universitaria compartiendo los valores que difundimos cantando.

A Fernando Clara, quien hizo prender en mí la inquietud por la investigación. Mi ejemplo de silenciosa humildad y grandeza académica.

A Isabel Passoni, Adriana Scandurra y Emilio Maldonado, mis compañeros de todos los días, compañeros en la ciencia y en los detalles de la vida diaria, quienes alegran y divierten cada jornada laboral con permanente buena onda.

A Rafael Espin Andrade, quien desde su Cuba natal aportó no solamente brillantes ideas que enriquecieron imprevistamente este trabajo, sino que sumó grandes párrafos de aliento en sus e-mails que hicieron posible que nunca parara de trabajar.

A Guillermo Abras, Teresa Codagnone, Carmen Balsamo y Alicia Tarditti, mis primeros cálidos Jefes de cátedra, de los cuales aprendí todo lo mejor que un docente debe tener, cada uno con su estilo.

A Marcel Brun, Juani Pastore y Agustina Bouchet por suplir con sus valiosas contribuciones mis falencias en la expresión formal de la matemática.

A Jorge Martínez Arca, Lucy Dai Pra, Azul Gonzalez y Eduardo Blotta, valiosas personas con las que trabajar es siempre la mejor experiencia de equipo.

A Victoria D'Onofrio, mi primer contacto y apoyo oficial de la facultad, cuando ella era presidenta del Centro de Estudiantes.

A Manuel González, quien permanentemente me ha demostrado su confianza para colaborar en sus tareas de conducción.

Al Dr. Anibal Introzzi, al Dr. Carlos Capiel, al Dr. Sebastián Costantino, al Dr. Gerardo Tusman y al Dr. Omar Montes, médicos con gran inquietud por la investigación, que me mostraron que la medicina es una ciencia que fascina e hicieron importantísimos aportes para este trabajo.

Y en otro orden de cosas... gracias también:

A mi Dios y Señor, que me ha dado mucho más de lo que merezco.

A mi Mamá y mi Papá, Elba y Jorge, el origen de todo y mi línea de referencia.

A mi esposa Ana y mi hija Magdalena, mi visión del amor y la paciencia, las que hacen de mis días una larga experiencia de cariño. Perdón por mis ausencias, físicas y de las otras, muchas de las cuales esta tesis es culpable.

A mi hermana Rosana, la otra mamá que he recibido en esta vida, y a mi cuñado Luis, ejemplo de coherencia, fuerza y silenciosa lucha.

A mis sobrinos Iván y Gerónimo, mi experiencia de dar y recibir montañas de ternura, de ver nacer y crecer la vida.

A mi familia política: Maruca, mi abuela postiza, Alicia, mi suegra que está siempre presente con su inagotable energía, y Nerea, Walter, Pancho, Juliana y Nacho, verdaderos hermanos adicionales que me ha regalado esta vida.

A mis infinitos amigos, primos y tíos, eternos compañeros de camino con que he sido bendecido: Silvina, Coco, Tía Dorita, Tía Dora, Mariela, Mario, Gati, Pablo, Valentín, Cyntia, Gustavo, Laura, Miguel, Susana, Carlitos, Inés, Mariana, Horacito, Amalia, Ana, Laura, Germán, Silvia, Pichi, Myriam, Gustavo, Jorgelina, Hernán, Mariela, Eduardo, Pablo, Andrea, Sebastián... y los que quizá olvido de mencionar explícitamente aquí.

Gracias, infinitas gracias a todos.

Gustavo

## Referencias

# Referencias

---

- Abras G. N., Ballarin V. L. (2005) A Weighted K-means Algorithm applied to Brain Tissue Classification. *Journal of Computer Science & Technology*, Vol. 5 (3), pp. 121-126.
- Alcalá R., Alcalá Fernández J., Casillas J., Cordón O., Herrera F. (2006) Hybrid learning models to get the interpretability - accuracy trade-off in fuzzy modeling. *Soft Computing*, Vol. 10 (9), pp. 717-734. doi: 10.1007/s00500-005-0002-1.
- American Academy of Neurology (1994) Practice parameter for diagnosis and evaluation of dementia (summary statement) - Report of the Quality Standards Subcommittee of the American Academy of Neurology. *Neurology*, Vol. 44 (11), pp. 2203-2206.
- Angelini E. D., Song T., Mensh B. D., Laine A. F. (2007) Brain MRI segmentation with multiphase minimal partitioning: a comparative study. *International Journal of Biomedical Imaging*, Vol. 2007, pp. 10526. doi: 10.1155/2007/10526.
- Ashburner J., Friston K. J. (2003) Spatial normalization using basis functions. IN Frackowiak, R. S. J., Friston, K. J., Frith, C., Dolan, R., Price, C. J., Zeki, S., Ashburner, J., Penny, W. (Eds.) *Human Brain Function*. 2nd. ed., Academic Press.
- Ashburner J., Friston K. J. (2005) Unified segmentation. *Neuroimage*, Vol. 26 (3), pp. 839-851. doi: 10.1016/j.neuroimage.2005.02.018.
- Ayer A. J. (1940) *The foundations of empirical knowledge*, London, Macmillan. ISBN 0333-0-0572-4.
- Ballarin V. L. (2001) *Procesamiento Digital de Imágenes aplicado a la Clasificación de Tejido Cerebral*. Tesis de Doctorado en Ciencias Biológicas (orientación Bioingeniería). Universidad Nacional de Tucumán, Tucumán.
- Ballarin V. L., Meschino G. J., Abras G. N., Passoni L. I. (2005) Segmentación de Imágenes Cerebrales de Resonancia Magnética Basada en Redes Neuronales de Regresión Generalizada. En resúmenes de *XV Congreso Argentino de Bioingeniería*, Vol. Publicado en CD, pp. 74. Septiembre, Paraná, Argentina.

- Baraldi A., Blonda P. (1999) A survey of fuzzy clustering algorithms for pattern recognition. II. *IEEE Transactions on Systems, Man and Cybernetics, Part B (Cybernetics)*, Vol. 29, pp. 786-801.
- Barra V., Lundervold A. (2007) A Collaborative Software Tool for the Evaluation of MRI Brain Segmentation Methods. En resúmenes de *International Symposium on Information Technology Convergence (ISITC 2007)*, pp. 235-239.
- Berkeley G. (1710) *A Treatise concerning the Principles of Human Knowledge*, Dublin, Printed by Aaron Rhames for Jeremy Pepyat. ISBN 978-0-1987-5161-8.
- Bezdek J. C., Hall L. O., Clarke L. P. (1993) Review of MR image segmentation techniques using pattern recognition. *Medical Physics*, Vol. 20 (4), pp. 1033-1048.
- Blink E. J. (2004) MRI: Physics. En <http://www.mri-physics.com/html/textuk.html>.
- Bonissone P. P. (1997) Soft computing: the convergence of emerging reasoning technologies. *Soft Computing*, Vol. 1 (1), pp. 6-18.
- Bonissone P. P., Subbu R., Eklund N., Kiehl T. R. A. (2006) Evolutionary algorithms + domain knowledge = real-world evolutionary computation. *IEEE Transactions on Evolutionary Computation*, Vol. 10 (3), pp. 256-280.
- Bonnet N., Cutrona J., Herbin M. (2002) A 'no-threshold' histogram-based image segmentation method. *Pattern Recognition*, Vol. 35 (10), pp. 2319-2322. doi: 10.1016/S0031-3203(02)00057-2.
- Bouix S., Martin-Fernandez M., Ungar L., Nakamura M., Koo M. S., Mccarley R. W., Shenton M. E. (2007) On evaluating brain tissue classifiers without a ground truth. *Neuroimage*, Vol. 36 (4), pp. 1207-1224. doi: 10.1016/j.neuroimage.2007.04.031.
- Bradley W. G., Jr., Waluch V., Yadley R. A., Wycoff R. R. (1984) Comparison of CT and MR in 400 patients with suspected disease of the brain and cervical spinal cord. *Radiology*, Vol. 152 (3), pp. 695-702.
- Brunetti A., Postiglione A., Tedeschi E., Ciarmiello A., Quarantelli M., Covelli E. M., Milan G., Larobina M., Soricelli A., Sodano A., Alfano B. (2000) Measurement of global brain atrophy in Alzheimer's disease with unsupervised segmentation of spin-echo MRI studies. *Journal of Magnetic Resonance Imaging*, Vol. 11 (3), pp.

- 260-266. doi: 10.1002/(SICI)1522-2586(200003)11:3<260::AID-JMRI4>3.0.CO;2-I.
- Castellanos R., Mitra S. (2000) Segmentation of Magnetic Resonance Images Using a Neuro-Fuzzy Algorithm. En resúmenes de *13th IEEE Symposium on Computer-Based Medical Systems (CBMS '00)*, pp. 207, Washington, DC, USA.
- Clark M. C., Hall L. O., Goldgof D. B., Clarke L. P., Velthuizen R. P., Silbiger M. S. (1994) MRI segmentation using fuzzy clustering techniques. *IEEE Engineering in Medicine and Biology Magazine*, Vol. 13 (5), pp. 730-742.
- Cocosco C. A., Zijdenbos A. P., Evans A. C. (2003) A fully automatic and robust brain MRI tissue classification method. *Medical Image Analysis*, Vol. 7 (4), pp. 513-527.
- Cordon O. (2004) *Genetic fuzzy systems: new developments*, Amsterdam, Elsevier B.V. ISBN 8-4973-2433-1.
- Cordon O., Gomide F., Herrera F., Hoffmann F., Magdalena L. (2004) Ten years of genetic fuzzy systems: current framework and new trends. *Fuzzy Sets and Systems*, Vol. 141 (1), pp. 5-31.
- Courchesne E., Chisum H. J., Townsend J., Cowles A., Covington J., Egaas B., Harwood M., Hinds S., Press G. A. (2000) Normal brain development and aging: quantitative analysis at in vivo MR imaging in healthy volunteers. *Radiology*, Vol. 216 (3), pp. 672-682.
- Courchesne E., Karns C. M., Davis H. R., Ziccardi R., Carper R. A., Tigue Z. D., Chisum H. J., Moses P., Pierce K., Lord C., Lincoln A. J., Pizzo S., Schreibman L., Haas R. H., Akshoomoff N. A., Courchesne R. Y. (2001) Unusual brain growth patterns in early life in patients with autistic disorder: an MRI study. *Neurology*, Vol. 57 (2), pp. 245-254.
- Courchesne E., Pierce K., Schumann C. M., Redcay E., Buckwalter J. A., Kennedy D. P., Morgan J. (2007) Mapping early brain development in autism. *Neuron*, Vol. 56 (2), pp. 399-413. doi: 10.1016/j.neuron.2007.10.016.
- Coussement A. (2000) El canto de los protones (RM ¿Sin esfuerzo?). En <http://coussement.unice.fr/>.
- Cox E. A. (1994) *The Fuzzy Systems Handbook*, Cambridge, MA, USA, Academic Press Professional. ISBN 0-1219-4270-8.

- Crum W. R., Camara O., Hill D. L. G. (2006) Generalized Overlap Measures for Evaluation and Validation in Medical Image Analysis. *IEEE Transactions on Medical Imaging*, Vol. 25 (11), pp. 1451-1461.
- Chang C. W., Ying H., Kent T. A., Yen J., Ketonen L. M., Reynolds M. L., Hillman G. R. (2002) A New Method for Two-stage Hybrid Fuzzy Segmentation of MR Images of Human Brains with Lesions. *International Journal of Fuzzy Systems*, Vol. 4 (4), pp. 873-882.
- Chen S. S., Keller J. M., Crownover R. M. (1993) On the calculation of fractal features from images. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, Vol. 15 (10), pp. 1087-1090.
- Chua S. E., Cheung C., Cheung V., Tsang J. T., Chen E. Y., Wong J. C., Cheung J. P., Yip L., Tai K. S., Suckling J., McAlonan G. M. (2007) Cerebral grey, white matter and csf in never-medicated, first-episode schizophrenia. *Schizophrenia Research*, Vol. 89 (1-3), pp. 12-21. doi: 10.1016/j.schres.2006.09.009.
- Chulho W., Dong Hun K. (2007) Region growing method using edge sharpness for brain ventricle detection. En resúmenes de *SICE Annual Conference*, pp. 1930-1933, doi: 10.1109/SICE.2007.4421302.
- Chun D. N., Yang H. S. (1996) Robust image segmentation using genetic algorithm with a fuzzy measure. *Pattern Recognition*, Vol. 29 (7), pp. 1195-1211.
- Dalton C. M., Chard D. T., Davies G. R., Miszkief K. A., Altmann D. R., Fernando K., Plant G. T., Thompson A. J., Miller D. H. (2004) Early development of multiple sclerosis is associated with progressive grey matter atrophy in patients presenting with clinically isolated syndromes. *Brain*, Vol. 127 (Pt 5), pp. 1101-7. doi: 10.1093/brain/awh126.
- Davis P. C., Gray L., Albert M., Wilkinson W., Hughes J., Heyman A., Gado M., Kumar A. J., Destian S., Lee C., Et Al. (1992) The Consortium to Establish a Registry for Alzheimer's Disease (CERAD). Part III. Reliability of a standardized MRI evaluation of Alzheimer's disease. *Neurology*, Vol. 42 (9), pp. 1676-1680.
- Delon J., Desolneux A., Lisani J. L., Petro A. B. (2005) Color Image Segmentation Using Acceptable Histogram Segmentation. *Pattern Recognition and Image Analysis*, Vol. 3523, pp. 239-246. doi: 10.1007/b136831.

- Denkowski M., Chlebiej M., Mikolajczak P. (2004) Segmentation of human brain MR images using rule-based fuzzy logic inference. *Studies in Health Technology and Informatics*, Vol. 105, pp. 264-272.
- Di Bona S., Niemann H., Pieri G., Salvetti O. (2003) Brain volumes characterisation using hierarchical neural networks. *Artificial Intelligence in Medicine*, Vol. 28 (3), pp. 307-322.
- Dickey C. C., Mccarley R. W., Shenton M. E. (2002) The brain in schizotypal personality disorder: a review of structural MRI and CT findings. *Harvard Review of Psychiatry*, Vol. 10 (1), pp. 1-15.
- Dougherty E. R. (1993) *Mathematical Morphology in Image Processing*, CRC Press. ISBN 0-8247-8724-2.
- Dounias G., Linkens D. (2004) Adaptive systems and hybrid computational intelligence in medicine. *Artificial Intelligence in Medicine*, Vol. 32 (3), pp. 151-155.
- Dubois D., Prade H. (1985) Review of fuzzy set aggregation connectives. *Information Sciences*, Vol. 36 (1-2), pp. 85-121.
- Dubois D., Prade H. M. (1980) *Fuzzy Sets and Systems: Theory and Applications*, San Diego, CA, USA, Academic Press Inc. ISBN 978-0-1222-2750-9.
- Dubois D., Prade H. M. (2000) *Fundamentals of fuzzy sets*, Boston, Kluwer Academic. ISBN 0-7923-7732-X.
- Duda R. O., Hart P. E., Stork D. G. (2000) *Pattern Classification*, New York, Wiley Interscience. ISBN 978-0-4710-5669-0.
- Edelman R. R., Hesselink J. R., Zlatkin M. (1996) *Clinical Magnetic Resonance Imaging*, W.B. Saunders Company. ISBN 978-0-7216-2241-5.
- Egmont-Petersen M., De Ridder D., Handels H. (2002) Image processing with neural networks--a review. *Pattern Recognition*, Vol. 35 (10), pp. 2279-2301.
- Elkan C., Berenji H. R., Chandrasekaran B., De Silva C. J. S., Attikiouzel Y., Dubois D., Prade H., Smets P., Freksa C., Garcia O. N., Klir G. J., Bo Yuan A., Mamdani E. H., Pelletier F. J., Ruspini E. H., Turksen B., Vadiée N., Jamshidi M., Pei-Zhuang Wang A., Sie-Keng Tan A., Shaohua Tan A., Yager R. R., Zadeh L. A. (1994) The paradoxical success of fuzzy logic. *IEEE Expert*, Vol. 9 (4), pp. 3-49.
- Espín Andrade R. A., Alberto C. L., Carignano C. (2007) Comparación de métodos de evaluación de eficiencia: una aplicación a la educación superior en Argentina.

En resúmenes de *Tercer Encuentro de la Red Iberoamericana de Evaluación y Decisión Multicriterio (RED-M 2007)*, Vol. Publicado en CD. 5-8 de noviembre, Culiacán, Sinaloa, México.

Espin Andrade R. A., Marx Gómez J., Mazcorro Téllez G., Fernández González E. (2003) Compensatory Logic: A Fuzzy Normative Model for Decision Making. En resúmenes de *10th Congress of International Association for Fuzzy-Set Management and Economy - Emergent Solutions for the Information and Knowledge Economy (SIGEF 2003)*, pp. 135-146, León, España.

Espin Andrade R. A., Marx Gómez J., Mazcorro Téllez G., Fernández González E., Lecich M. I. (2004) Compensatory Logic: A Fuzzy Approach to Decision Making. En resúmenes de *4th International Symposium on Engineering of Intelligent Systems - EIS 2004*, pp. Publicado en CD. 29 de febrero al 2 de marzo, Madeira, Portugal.

Fan J., Zeng G., Body M., Hacid M.-S. (2005) Seeded region growing: an extensive and comparative study. *Pattern Recognition Letters*, Vol. 26 (8), pp. 1139-1156.

Fletcher L. M., Hornak J. P. (1994) Multispectral Image Segmentation in Magnetic Resonance. IN Dougherty, E. R. (Ed.) *Digital Image Processing Methods*. New York, Marcel Dekker.

Fox J. (1994) On the Necessity of Probability: Reasons to Believe and Ground for Doubt. IN Wright, G., Ayton, P. (Eds.) *Subjective Probability*. England, Wiley & Sons.

French S. (1986) *Decision theory: an introduction to the mathematics of rationality*, New York, NY, USA, Halsted Press. ISBN 0-4702-0308-0.

Garciadiego A. R. (1992) *Bertrand Russell and the origins of the set-theoretic "paradoxes"*, Basel, Birkhäuser. ISBN 3-7643-2669-7.

Ge Y., Grossman R. I., Babb J. S., Rabin M. L., Mannon L. J., Kolson D. L. (2002) Age-related total gray matter and white matter changes in normal adult brain. Part I: volumetric MR imaging analysis. *American Journal of Neuroradiology*, Vol. 23 (8), pp. 1327-1333.

Germond L., Dojat M., Taylor C., Garbay C. (2000) A cooperative framework for segmentation of MRI brain scans. *Artificial Intelligence in Medicine*, Vol. 20 (1), pp. 77-93.

- Gibbons J. D., Chakraborti S. (2003) *Nonparametric Statistical Inference*, CRC Press. ISBN 978-0-8247-4052-8.
- Gilmore J. H., Lin W., Prastawa M. W., Looney C. B., Vetsa Y. S., Knickmeyer R. C., Evans D. D., Smith J. K., Hamer R. M., Lieberman J. A., Gerig G. (2007) Regional gray matter growth, sexual dimorphism, and cerebral asymmetry in the neonatal brain. *The Journal of Neuroscience*, Vol. 27 (6), pp. 1255-1260. doi: 10.1523/JNEUROSCI.3339-06.2007.
- Glasbey C. A. (1993) An analysis of histogram-based thresholding algorithms. *CVGIP: Graph. Models Image Process.*, Vol. 55 (6), pp. 532-537. doi: 10.1006/cgip.1993.1040.
- Goldberg D. E. (1989) *Genetic Algorithms in Search, Optimization and Machine Learning*, Boston, MA, USA, Addison-Wesley Longman Publishing Co., Inc. ISBN 0-2011-5767-5.
- Gonzalez M. A., Meschino G. J., Ballarin V. L. (2007) Automatic fuzzy inference system development for marker based watershed segmentation. *Journal of Physics: Conference Series* Vol. 90. doi: 10.1088/1742-6596/90/1/012061.
- Grau V., Mewes A. U. J., Alcaniz M., Kikinis R. A. K. R., Warfield S. K. A. W. S. K. (2004) Improved watershed transform for medical image segmentation using prior information. *IEEE Transactions on Medical Imaging*, Vol. 23 (4), pp. 447-458.
- Haroun R., Boumghar F., Hassas S., Hamami L. (2006) A Massive Multi-agent System for Brain MRI Segmentation. IN Springer (Ed.) *Massively Multi-Agent Systems I*. Berlin / Heidelberg.
- Hata Y., Kobashi S., Hirano S., Kitagaki H., Mori E. (2000) Automated segmentation of human brain MR images aided by fuzzy information granulation and fuzzy inference. *IEEE Transactions on Systems, Man, and Cybernetics, Part C: Applications and Reviews*, Vol. 30 (3), pp. 381- 395.
- Held K., Kops E. R., Krause B. J., Wells W. M., Kikinis R. A., Muller-Gartner H. W. (1997) Markov random field segmentation of brain MR images. *IEEE Transactions on Medical Imaging*, Vol. 16 (6), pp. 878-886.
- Hermann B., Jones J., Sheth R., Dow C., Koehn M., Seidenberg M. (2006) Children with new-onset epilepsy: neuropsychological status and brain structure. *Brain*, Vol. 129 (Pt 10), pp. 2609-2619. doi: 10.1093/brain/awl196.

- Herrera F. (2005) Genetic Fuzzy Systems: Status, Critical Considerations and Future Directions. *International Journal of Computational Intelligence Research*, Vol. 1 (1), pp. 59-67.
- Ho B.-C. (2007) MRI brain volume abnormalities in young, nonpsychotic relatives of schizophrenia probands are associated with subsequent prodromal symptoms. *Schizophrenia Research*, Vol. 96 (1), pp. 1-13.
- Hogarth R. (1991) *Judgement and Choice*, Wiley, USA. ISBN 978-0-4719-1479-2.
- Hornak J. P. (2007) The Basics of MRI. En <http://www.cis.rit.edu/htbooks/mri/>.
- Hulshoff Pol H. E., Kahn R. S. (2008) What happens after the first episode? A review of progressive brain changes in chronically ill patients with schizophrenia. *Schizophrenia Bulletin*, Vol. 34 (2), pp. 354-366. doi: 10.1093/schbul/sbm168.
- Hulshoff Pol H. E., Schnack H. G., Bertens M. G., Van Haren N. E., Van Der Tweel I., Staal W. G., Baare W. F., Kahn R. S. (2002) Volume changes in gray matter in patients with schizophrenia. *American Journal of Psychiatry*, Vol. 159 (2), pp. 244-250.
- Ibrahim M., John N., Kabuka M., Younis A. (2006) Hidden Markov models-based 3D MRI brain segmentation. *Image and Vision Computing*, Vol. 24 (10), pp. 1065-1079.
- Iowa Mental Health Clinical Research Center (2005) BRAINS Software Package. En <http://www.psychiatry.uiowa.edu/mhcrc/IPLpages/BRAINS.htm>.
- Jian Y. (2005) General C-means clustering model. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, Vol. 27 (8), pp. 1197-1211.
- Jimenez-Alaniz J. R., Medina-Banuelos V., Yanez-Suarez O. (2006) Data-driven brain MRI segmentation supported on edge confidence and a priori tissue information. *IEEE Transactions on Medical Imaging*, Vol. 25 (1), pp. 74-83.
- Jinn-Yi Y., Fu J. C. (2008) A hierarchical genetic algorithm for segmentation of multi-spectral human-brain MRI. *Expert Systems with Applications*, Vol. 34 (2), pp. 1285-1295. doi: 10.1016/j.eswa.2006.12.012.
- Kahneman D., Tversky A. (2005) Subjective Probability: A Judgement of Representativeness. IN Kahneman D., S. P., Tversy A. (Ed.) *Judgement under Uncertainty. Heuristics and Biases*. 21st printing from 1982 ed. Cambridge, USA, Cambridge University Press.

- Kang J., Min L., Luan Q., Li X., Liu J. (2008) Novel modified fuzzy c-means algorithm with applications. *Digital Signal Processing*, Vol. In Press, Corrected Proof. doi: 10.1016/j.dsp.2007.11.005.
- Karnik N. N., Mendel J. M. (2001) Operations on type-2 fuzzy sets. *Fuzzy Sets and Systems*, Vol. 122, pp. 327-348.
- Karnik N. N., Mendel J. M., Qilian L. (1999) Type-2 fuzzy logic systems. *IEEE Transactions on Fuzzy Systems*, Vol. 7 (6), pp. 643-658.
- Kecman V. (2001) *Learning and Soft Computing - Support Vector Machines, Neural Networks and Fuzzy Logic Models*, Cambridge, Massachusetts, The MIT Press. ISBN 0-2621-1255-8.
- Knopman D. S., Dekosky S. T., Cummings J. L., Chui H., Corey-Bloom J., Relkin N., Small G. W., Miller B., Stevens J. C. (2001) Practice parameter: diagnosis of dementia (an evidence-based review). Report of the Quality Standards Subcommittee of the American Academy of Neurology. *Neurology*, Vol. 56 (9), pp. 1143-1153.
- Kobashi S., Fujiki Y., Matsui M., Inoue N., Kondo K., Hata Y., Sawada T. (2006) Interactive segmentation of the cerebral lobes with fuzzy inference in 3T MR images. *IEEE Transactions on Systems, Man and Cybernetics, Part B (Cybernetics)*, Vol. 36 (1), pp. 74-86.
- Kohavi R., Provost F. (1998) Glossary of Terms. *Machine Learning*, Vol. 30 (2/3), pp. 271-274.
- Kohonen T. (2001) *Self-Organizing Maps*, New York, USA, Springer. ISBN 978-3-540-67921-9.
- Kovalev V. A., Kruggel F., Gertz H. J., Von Cramon D. Y. A. (2001) Three-dimensional texture analysis of MRI brain datasets. *IEEE Transactions on Medical Imaging*, Vol. 20 (5), pp. 424-433.
- Krause P. J., Clark D. A. (1994) Uncertainty and subjective probability in AI systems. IN Wright, G., Ayton, P. (Eds.) *Subjective Probability*. England, Wiley & Sons.
- Kuehn M. (2001) *Kant: a biography*, Cambridge, Cambridge University Press. ISBN 0-5215-2406-7 y 0-5214-9704-3.
- Kuncheva L. I. (1995) Editing for the k-nearest neighbors rule by a genetic algorithm. *Pattern Recognition Letters*, Vol. 16 (8), pp. 809-814.

- Li H.-X., Yen V. C. (1995b) *Fuzzy Sets and Fuzzy Decision-Making*, N.W. Boca Raton, CRC Press. ISBN 0-8493-8931-3.
- Li H., Yezzi A., Cohen L. D. (2006) 3D Brain Segmentation Using Dual-Front Active Contours with Optional User Interaction. *International Journal of Biomedical Imaging*, Vol. 2006, pp. 1-17. doi: 10.1155/IJBI/2006/53186.
- Li S. Z. (1995a) *Markov random field modeling in computer vision*, Springer-Verlag. ISBN 4-4317-0145-1.
- Liang Z.-P., Lauterbur P. C., IEEE Engineering in Medicine and Biology Society. (2000) *Principles of magnetic resonance imaging : a signal processing perspective*, Bellingham, Wash, New York, SPIE Optical Engineering Press; IEEE Press. ISBN 0-7803-4723-4.
- Liew A. W. C., Yan H. (2006) Current Methods in the Automatic Tissue Segmentation of 3D Magnetic Resonance Brain Images. *Current Medical Imaging Reviews*, Vol. 2, pp. 91-103.
- Likar B., Derganc J., Pernus F. (2002) Segmentation-based retrospective correction of intensity nonuniformity in multispectral MR images. En resúmenes de *SPIE 2002; Medical Imaging 2002: Image Processing*, pp. 1531-1540.
- Lindley D. (1994) *Subjective Probability*, England, Wiley & Sons. ISBN 978-0-5215-3668-4.
- Lukasiewicz J. (1970) *Selected works*, Amsterdam, North-Holland Pub. Co. ISBN 0-7204-2252-3.
- Macovski A. (1996) Noise in MRI. *Magnetic Resonance in Medicine*, Vol. 36 (3), pp. 494-497.
- Mamdani E. H. (1974) Application of fuzzy algorithms for control of a simple dynamic plant. *Proceedings of IEE*, Vol. 121 (12), pp. 1585-1588.
- Manes F. (2000) Resonancia magnética nuclear en la Enfermedad de Alzheimer. *Revista Neurológica Argentina*, Vol. 25, pp. 29-37.
- Mcinerney T., Terzopoulos D. (1996) Deformable models in medical images analysis: a survey. *Medical Image Analysis*, Vol. 1 (2), pp. 91-108.
- Mechtcheriakov S., Brenneis C., Egger K., Koppelstaetter F., Schocke M., Marksteiner J. (2007) A widespread distinct pattern of cerebral atrophy in patients with alcohol addiction revealed by voxel-based morphometry. *Journal of Neurology*,

- Neurosurgery & Psychiatry*, Vol. 78 (6), pp. 610-614. doi: 10.1136/jnnp.2006.095869.
- Mendel J. M. (2001) *Uncertain Rule-Based Fuzzy Logic Systems: Introduction and New Directions*, Upper Saddle River, NJ, USA, Prentice-Hall PTR. ISBN 978-0-1304-0969-0.
- Mendel J. M. (2003) Type-2 Fuzzy Sets: Some Questions and Answers. *IEEE Connections, Newsletter of the IEEE Neural Networks Society*, Vol. 1, pp. 10-13.
- Mendel J. M. (2007a) Type-2 Fuzzy Sets and Systems: an Overview. *IEEE Computational Intelligence Magazine*, Vol. 2, pp. 20-29.
- Mendel J. M. (2007b) Advances in Type-2 Fuzzy Sets and Systems. *Information Sciences*, Vol. 177, pp. 84-110.
- Mendel J. M., John R. I., Liu F. (2006) Interval Type-2 Fuzzy Logic Systems Made Simple. *IEEE Transactions on Fuzzy Systems*, Vol. 14, pp. 808-821.
- Meschino G. J., Espin Andrade R. A., Ballarin V. L. (2007) Reconocimiento de tejidos en imágenes cerebrales de Resonancia Magnética a través de valores de verdad de predicados difusos. *Mecánica Computacional*, Vol. 26, pp. 865-878.
- Meschino G. J., Moler E. G. (2004a) Semiautomated Segmentation of Bone Marrow Biopsies by Texture Features and Mathematical Morphology. *Analytical and Quantitative Cytology and Histology*, Vol. 26 (1), pp. 31-38.
- Meschino G. J., Passoni L. I., Moler E. G. (2004b) Semiautomated Segmentation of Bone Marrow Biopsies Images Based on Texture Features and GRNN. En resúmenes de *X Congreso Argentino de Ciencias de la Computación (CACIC 2004)*, Vol. Publicado en CD. 4 al 8 de octubre, San Justo, Buenos Aires, Argentina.
- Meschino G. J., Passoni L. I., Scandurra A. G., Ballarin V. L. (2006) Representación Automática Pseudo Color de Imágenes Médicas mediante Mapas Autoorganizados. En resúmenes de *35º Jornadas Argentinas de Informática e Investigación Operativa - SIS, Simposio de Informática y Salud*, Vol. Publicado en CD, pp. 105-115. 4 al 8 de septiembre, Mendoza, Argentina.
- Moler E. G. (1998) *Análisis del status epidemiológico de los algoritmos: implicancias en los desarrollos científico-tecnológicos*. Tesis de Magister Scientiae en Filosofía de la Ciencia. Universidad Nacional de Mar del Plata, Mar del Plata.

- Moler E. G. (2003) Técnicas de procesamiento digital de imágenes aplicadas al cálculo de parámetros histomorfométricos no-euclidianos. *Revista Argentina de Bioingeniería*, Vol. 9 (1), pp. 11-15.
- Montréal Neurological Institute (2007) MNI's BrainWeb dataset: <http://www.bic.mni.mcgill.ca/brainweb/>. Montréal Neurological Institute, McGill University.
- Narr K. L., Sharma T., Woods R. P., Thompson P. M., Sowell E. R., Rex D., Kim S., Asuncion D., Jang S., Mazziotta J., Toga A. W. (2003) Increases in regional subarachnoid CSF without apparent cortical gray matter deficits in schizophrenia: modulating effects of sex and age. *American Journal of Psychiatry*, Vol. 160 (12), pp. 2169-2180.
- Pal T., Pal N. R. (2003) SOGARG: A self-organized genetic algorithm-based rule generation scheme for fuzzy controllers. *IEEE Transactions on Evolutionary Computation*, Vol. 7 (4), pp. 397-415.
- Passino K. M., Yurkovich S. (1998) *Fuzzy Control*, Boston, MA, USA, Addison-Wesley Longman Publishing Co., Inc. ISBN 0-2011-8074-X.
- Pastore J. I., Meschino G. J., Moler E. (2006) Segmentación de biopsias de médula ósea mediante filtros morfológicos y rotulación de regiones homogéneas. *Revista Brasileira de Engenharia Biomédica*, Vol. 21 (1), pp. 81-88.
- Pastore J. I., Moler E. G., Ballarin V. L. (2005) Segmentation of brain magnetic resonance images through morphological operators and geodesic distance. *Digital Signal Processing*, Vol. 15 (2), pp. 153-160.
- Pedrycz W., Reformat M., Li K. (2006) OR/AND neurons and the development of interpretable logic models. *IEEE Transactions on Neural Networks*, Vol. 17 (3), pp. 636-658.
- Peirce C. S. (1883) *Studies in logic*, Amsterdam, Benjamin. ISBN 9-0272-3271-7.
- Pham D. L., Xu C., Prince J. L. (2000) Current methods in medical image segmentation. *Annual Review of Biomedical Engineering*, Vol. 2, pp. 315-37. doi: 2/1/315 [pii] 10.1146/annurev.bioeng.2.1.315.
- Pohl K. M., Bouix S., Kikinis R., Grimson W. E. L. (2004) Anatomical guided segmentation with non-stationary tissue class distributions in an expectation-

- maximization framework. En resúmenes de *IEEE International Symposium on Biomedical Imaging: Nano to Macro 2004*, pp. 81-84.
- Pyun K. P., Johan L., Chee Sun W., Gray R. M. A. G. R. M. (2007) Image Segmentation Using Hidden Markov Gauss Mixture Models. *IEEE Transactions on Image Processing*, Vol. 16 (7), pp. 1902-1911.
- Richard N., Dojat M., Garbay C. (2004) Automated segmentation of human brain MR images using a multi-agent approach. *Artificial Intelligence in Medicine*, Vol. 30 (2), pp. 153-175.
- Rokach L. (2008) Genetic algorithm-based feature set partitioning for classification problems. *Pattern Recognition*, Vol. 41 (5), pp. 1676-1700.
- Sahoo P. K., Soltani S., Wong A. K. C., Chen Y. C. (1988) A survey of thresholding techniques. *Computer Vision, Graphics, and Image Processing*, Vol. 41 (2), pp. 233-260. doi: 10.1016/0734-189X(88)90022-9.
- Salvado O., Bourgeat P., Tamayo O. A., Zuluaga M. A. Z. M., Ourselin S. A. O. S. (2007) Fuzzy classification of brain MRI using a priori knowledge: weighted fuzzy C-means. En resúmenes de *IEEE 11th International Conference on Computer Vision 2007 (ICCV 2007)*, pp. 1-8. 14 al 20 de Octubre, Rio de Janeiro, Brasil.
- Santalla H., Meschino G. J., Ballarin V. L. (2007) Effects on MR images compression in tissue classification quality. *Journal of Physics: Conference Series*, Vol. 90. doi: 10.1088/1742-6596/90/1/012061.
- Sasikala M., Kumaravel N., Ravikumar S. (2006) Segmentation of Brain MR Images using Genetically Guided Clustering. En resúmenes de *28th Annual International Conference of the IEEE Engineering in Medicine and Biology Society (EMBS '06)*, pp. 3620-3623, New York, USA.
- Scandurra A. G., Meschino G. J., Passoni L. I., Dai Pra A. L., Introzzi A. R., Clara F. M. (2007) Optimization of arterial age prediction models based in pulse wave. *Journal of Physics: Conference Series*, Vol. 90. doi: 10.1088/1742-6596/90/1/012080.
- Schwarz D., Kasperek T. (2007) Brain Tissue Classification with Automated Generation of Training Data Improved by Deformable Registration. *Lecture Notes in Computer Science*, Vol. 4673 (1), pp. 301-308.

- Serra J. (1992) *Image Analysis and Mathematical Morphology*, Academic Press, Inc. ISBN 0-1263-7240-3.
- Shen S., Sandham W., Granat M., Sterr A. (2005) MRI fuzzy segmentation of brain tissue using neighborhood attraction with neural-network optimization. *IEEE Transactions on Information Technology in Biomedicine*, Vol. 9 (3), pp. 459-467.
- Sijbers J., Poot D., Den Dekker A. J., Pintjens W. (2007) Automatic estimation of the noise variance from the histogram of a magnetic resonance image. *Physics in Medicine and Biology*, Vol. 52, pp. 1335-1348.
- Siromoney A., Raghuram L., Siromoney A., Korah I., Prasad G. N. (2000) Inductive logic programming for knowledge discovery from MRI data. *IEEE Engineering in Medicine and Biology Magazine*, Vol. 19 (4), pp. 72-77.
- Siyal M. Y., Yu L. (2005) An intelligent modified fuzzy c-means based algorithm for bias estimation and segmentation of brain MRI. *Pattern Recognition Letters*, Vol. 26 (13), pp. 2052-2062.
- Smith V. L. (2000) Rational Choice: The contrast between Economics and Psychology. IN Smith, V. L. (Ed.) *Bargaining and Market Behavior*. Cambridge, UK, Cambridge University Press.
- Soille P., Pesaresi M. (2002) Advances in mathematical morphology applied to geoscience and remote sensing. *IEEE Transactions on Geoscience and Remote Sensing*, Vol. 40 (9), pp. 2042-2055.
- Song T., Angelini E. D., Mensh B. D., Laine A. (2004) Comparison study of clinical 3D MRI brain segmentation evaluation. En resúmenes de *Annual International Conference of the IEEE Engineering in Medicine and Biology Society*, Vol. 3, pp. 1671-1674, San Francisco, CA, USA.
- Song T., Jamshidi M. M., Lee R. R., Huang M. (2007) A Modified Probabilistic Neural Network for Partial Volume Segmentation in Brain MR Image. *IEEE Transactions on Neural Networks*, Vol. 18 (5), pp. 1424-1432.
- Sousa J. M. C., Kaymak U. (2003) *Fuzzy Decision Making in Modeling and Control*, World Scientific. ISBN 9-8102-4877-6.
- Storey P. (2006) Introduction to magnetic resonance imaging and spectroscopy. *Methods in Molecular Medicine*, Vol. 124, pp. 3-57.

- Sugeno M., Kang G. T. (1988) Structure Identification of Fuzzy Model. *Fuzzy Sets and Systems*, Vol. 28 (1), pp. 15-33.
- Supot S., Thanapong C., Chuchart P., Manas S. (2007) Segmentation of Magnetic Resonance Images Using Discrete Curve Evolution and Fuzzy Clustering. En resúmenes de *IEEE International Conference on Integration Technology 2007 (ICIT 07)*, pp. 697-700, Shenzhen, China.
- Szilágyi L., Szilágyi S. M., Benyó Z. (2007) A Modified Fuzzy C-Means Algorithm for MR Brain Image Segmentation. *Image Analysis and Recognition*, Vol. 4633, pp. 866-877. doi: 10.1007/978-3-540-74260-9.
- Takagi T., Sugeno M. (1985) Fuzzy identification of systems and its application to modeling and control. *IEEE Transactions on Systems, Man, and Cybernetics*, Vol. 15 (1), pp. 116-132.
- Tian D., Fan L. (2007) A Brain MR Images Segmentation Method Based on SOM Neural Network. En resúmenes de *The 1st International Conference on Bioinformatics and Biomedical Engineering 2007 (ICBBE 2007)*, pp. 686-689.
- Todd R. R., Buf J. M. H. D. (1993) A review of recent texture segmentation and feature extraction techniques. *CVGIP: Image Understanding*, Vol. 57 (3), pp. 359-372. doi: 10.1006/ciun.1993.1024.
- Tusman G., Suarez-Sipmann F., Bohm S. H., Pech T., Reissmann H., Meschino G. J., Scandurra A. G., Hedenstierna G. (2006) Monitoring dead space during recruitment and PEEP titration in an experimental model. *Intensive Care Medicine*, Vol. 32 (11), pp. 1863-1871.
- Tversky A., Kahneman D. (2005) Judgement under Uncertainty. Heuristics and Biases. IN Kahneman D., S. P., Tversy A. (Ed.) *Judgement under Uncertainty. Heuristics and Biases*. 21st printing from 1982 ed. Cambridge, USA, Cambridge University Press.
- Verdegay J. L. (2005) Una revisión de las metodologías que integran la Soft Computing. En resúmenes de *Simposio sobre Lógica Fuzzy y Soft Computing LFSC 2005 (EUSFLAT)*, pp. 151-156. Septiembre 2005, Granada, España.
- Vincent L., Soille P. (1991) Watersheds in Digital Spaces: An Efficient Algorithm Based on Immersion Simulations. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, Vol. 13 (6), pp. 583-598. doi: 10.1109/34.87344.

- Vovk U., Pernus F., Likar B. (2007) A Review of Methods for Correction of Intensity Inhomogeneity in MRI. *IEEE Transactions on Medical Imaging*, Vol. 26 (3), pp. 405-421.
- Wang D., Doddrell D. M. (2005) Method for a detailed measurement of image intensity nonuniformity in magnetic resonance imaging. *Medical Physics*, Vol. 32 (4), pp. 952-960.
- Warfield S. (1996) Fast k-NN classification for multichannel image data. *Pattern Recognition Letters*, Vol. 17 (7), pp. 713-721.
- Wells W. M., Grimson W. L., Kikinis R., Jolesz F. A. (1996) Adaptive segmentation of MRI data. *IEEE Transactions on Medical Imaging*, Vol. 15 (4), pp. 429-442. doi: 10.1109/42.511747.
- Wittgenstein L. (1958) *Philosophical investigations*, Oxford, Blackwell. ISBN 978-0-6312-3127-1.
- Yang M. S., Lin K. C., Liu H. C., Lirng J. F. (2007) Magnetic resonance imaging segmentation techniques using batch-type learning vector quantization algorithms. *Journal of Magnetic Resonance Imaging*, Vol. 25 (2), pp. 265-277.
- Yuan Y., Zhuang H. (1996) A genetic algorithm for generating fuzzy classification rules. *Fuzzy Sets and Systems*, Vol. 84 (1), pp. 1-19. doi: 10.1016/0165-0114(95)00302-9.
- Zadeh L. A. (1965) Fuzzy sets. *Information and Control*, Vol. 8, pp. 338-353.
- Zadeh L. A. (1975) The Concept of a Linguistic Variable and Its Application to Approximate Reasoning-1. *Information Sciences*, Vol. 8, pp. 199-249.
- Zhang Y., Brady M., Smith S. (2001) Segmentation of brain MR images through a hidden Markov random field model and the expectation-maximization algorithm. *IEEE Transactions on Medical Imaging*, Vol. 20 (1), pp. 45-57.
- Zhou Y., Bai J. (2007) Atlas-Based Fuzzy Connectedness Segmentation and Intensity Nonuniformity Correction Applied to Brain MRI. *IEEE Transactions on Biomedical Engineering*, Vol. 54 (1), pp. 122-129.
- Zijdenbos A. P., Forghani R., Evans A. C. (2002) Automatic "pipeline" analysis of 3-D MRI data for clinical trials: application to multiple sclerosis. *IEEE Transactions on Medical Imaging*, Vol. 21 (10), pp. 1280-1291.

Zimmermann H. J. (2001) *Fuzzy Set Theory and Its Applications*, Boston, Kluwer Academic Publishers. ISBN 978-0-7923-7435-0.

---