

El Reconocimiento de Patrones y su Aplicación a las Señales Digitales

María del Pilar Gómez Gil

Editora



ACADEMIA MEXICANA DE COMPUTACIÓN, A, C.

Segunda edición: 2019

Academia Mexicana de Computación, A. C.

Todos los derechos reservados conforme a la ley.

ISBN: 978-607-97357-7-7

Corrección de estilo: María del Pilar Gómez Gil, Enrique Sucar Succar.

Diseño de portada: Mario Alberto Vélez Sánchez.

Este libro se realizó con el apoyo del CONACyT, proyecto I1200/28/2019.

Queda prohibida la reproducción parcial o total, directa o indirecta, del contenido de esta obra, sin contar con autorización escrita de los autores, en términos de la Ley Federal del Derecho de Autor y, en su caso, de los tratados internacionales aplicables.

Hecho en México.

Made in Mexico.

El reconocimiento de patrones y su aplicación a las señales digitales.

Editado por: María del Pilar Gómez Gil

Autores participantes (en orden alfabético):

Vicente Alarcón Aquino

Dustin Carrión-Ojeda

Mario Ignacio Chacón Murguía

Rigoberto Fonseca Delgado

Elisa Isabel Gómez

Pilar Gómez Gil

Adolfo Guzmán Arenas

Hiram Habid López

Javier Herrera Vega

Pablo Ibargüengoitya

Felipe Orihuela Espina

J. Manuel Ramírez Cortés

Juan Alberto Ramírez Quintana

Alberto Reyes Ballesteros

Eduardo Rivas Posada

Juan Carlos Sánchez

Hermilo Sánchez Cruz

Guillermo Santamaría

Agradecimientos

Este libro no podría haberse realizado sin la participación entusiasta y desinteresada de un buen número de personas e instituciones. En primera instancia, agradecemos a la Academia Mexicana de la Computación por la oportunidad de difundir nuestra pasión por la ciencia. Gracias a su comité directivo, por la confianza depositada en nosotros para la realización de este proyecto.

No menos importante es nuestro profundo agradecimiento a los autores que participaron en la preparación de los capítulos, por su trabajo altruista dedicado a difundir la ciencia. Asimismo, agradecemos a todos los estudiantes, asistentes administrativos, colegas y amigos que de alguna manera cooperaron con sus comentarios, consejos y apoyo técnico. Nuestro agradecimiento también a las instituciones educativas y de investigación en donde los autores estamos adscritos.

Prólogo

A finales de 2017, la Academia Mexicana de la computación (Amexcomp) decidió incluir entre sus actividades de 2018, la creación de una serie de libros electrónicos encargados de difundir el conocimiento básico asociado a las disciplinas que la conforman. Este libro es parte del resultado de dicho trabajo, en lo que respecta a la sección académica de “Análisis de Señales y Reconocimiento de Patrones (AS/RP)”.

Nuestro quehacer diario actual está totalmente impregnado de tecnologías de información (TI), debido al avance impresionante en el desarrollo de diversas áreas del conocimiento, en especial de las ciencias de la computación y la electrónica. Estos avances han llevado a una importante disminución en costos de producción tanto del hardware como del software, lo que a su vez está apoyando el incremento continuo en las velocidades de cómputo y capacidades de almacenamiento de los equipos que utilizamos. La gran dependencia tecnológica de nuestras actividades de trabajo, diversión, comunicación y cuidado de la salud también es innegable. Como componente fundamental de una gran mayoría de estas herramientas de soporte tecnológico, se encuentran aplicaciones asociadas al análisis de señales digitales en el contexto de reconocimiento de patrones (AS/RP), el área que nos ocupa en este libro. Esta presencia ubicua de AS/RP lleva a los usuarios de TI a querer familiarizarse con conceptos que apoyan el

diseño de aplicaciones de AS/RP, así como enterarse sobre que se investiga en México con respecto a estas áreas; aquí intentaremos cubrir parte de este requerimiento.

Durante la preparación de este libro, los autores vivimos el interesante reto de intentar expresar en palabras simples algunos de los conceptos que sostienen nuestra investigación. Además, hemos tenido que decidir que puntos, aplicaciones y ejemplos incluir. AS/RP es un área inmensamente grande, por lo que resulta imposible cubrirla completamente en un libro. Finalmente, el libro quedó formado por 8 capítulos. El capítulo 1 inicia a los lectores en el proceso básico de reconocer objetos a través de la percepción digital realizada por algún transductor electrónico y analizada matemáticamente. Este capítulo también comenta brevemente sobre los primeros reconocedores, los tipos principales de algoritmos y las instituciones educativas y de investigación del país que han sobresalido en investigaciones asociadas a AS/RP. El capítulo 2 detalla los conceptos fundamentales que apoyan a las técnicas de reconocimiento de patrones, incluyendo algunas definiciones de métricas de desempeño y descripciones sobre los algoritmos de reconocimiento mas populares. El capítulo 3 trata sobre los conceptos de representación de mediciones digitales y las operaciones básicas que permiten manipularlas, incluyendo una breve explicación de que son los sistemas digitales lineales, así como las principales técnicas de análisis en espacios de frecuencia y tiempo. Asimismo, se dan algunos ejemplos de proyectos de investigación relacionados al procesamiento digital de señales. El capítulo 4 nos introduce en el camino que ha seguido la visión por computadora, incluyendo ejemplos de técnicas de procesamiento de imágenes, aplicaciones de la visión computacional en la bio-informática, vehículos autónomos, análisis de imágenes astronómicas e interpretación de escritura, compresión de imágenes

y otras aplicaciones desarrolladas tanto a nivel internacional como nacional. El capítulo 5 trata sobre la ciencia de datos y su aporte en la búsqueda de patrones en bases de datos. Asimismo se explica el proceso básico de minería de datos y se incluyen explicaciones de los diferentes enfoques, usos y aplicaciones en México relacionados a esta especialidad. El capítulo 6 nos introduce al uso de AS/RP para apoyar la obtención, administración y aprovechamiento de la energía, incluyendo la descripción de interesantes aplicaciones en México como el diagnóstico de transformadores y la predicción de velocidad del viento para apoyar la producción de energía eólica. Asimismo, dicho capítulo incluye una opinión sobre cuales son las tareas pendientes a resolverse en este campo asociadas a AS/RP. El capítulo 7 nos presenta un breve panorama sobre el uso de AS/RP para aplicaciones en la medicina y biología, explicando conceptos asociados a la reconstrucción, procesamiento, análisis e interpretación de imágenes bio-médicas y dando ejemplos de aplicaciones para la prevención, diagnóstico y tratamiento de enfermedades, así como para el análisis biológico a nivel celular, de organismos y de eco-sistemas. El capítulo 8 nos introduce en la situación actual del diseño de interfaces cerebro-computadora, incluyendo estrategias de adquisición de señales cerebrales, modelos matemáticos de procesamiento y clasificación, así como algunas aplicaciones desarrolladas en nuestro país y en otros lugares, enfocadas a apoyar personas con capacidades diferenciadas, así como aplicaciones de entretenimiento.

Concluimos este prólogo recordando la famosa frase de Sócrates: “el conocimiento os hará libres.” Deseamos que esta modesta aportación sirva para motivar a nuestros lectores a conocer más sobre estos temas y a participar activamente en el desarrollo del conocimiento, de manera que sigamos construyendo una sociedad cada día mas libre.

María del Pilar Gómez Gil
pgomez@inaoep.mx

Abreviaturas

<i>ADALINE</i>	<i>Adaptive Linear Neuron</i> (Neurona Adaptiva Lineal)
<i>ADC</i>	<i>Analog-to-digital converter</i> (Convrtidor Analógico-digital)
<i>ADN</i>	Ácido Desoxirribonucleico
<i>ANN</i>	<i>Artificial Neural Network</i> (Red Neuronal Artificial)
<i>AR</i>	<i>Auto-regressive models</i> (Modelos Auto-regresivos)
<i>ARMA</i>	<i>Auto-regressive moving average</i> (Modelos Auto-regresivos de Promediado móvil)
<i>ART</i>	<i>Adaptive Resonance Theory</i> (Teoría de Resonancia Adaptiva)
<i>ASCII</i>	<i>American Standard Code for</i> <i>Information Interchange</i> (Código Estándar Americano para Intercambio de Información)
<i>AS/RP</i>	Análisis de Señales y Reconocimiento de Patrones
<i>BN</i>	<i>Bayesian Network</i> (Red Bayesiana)
<i>BCI</i>	<i>Brain Computer Interface</i> (Interfaz Cerebro-Computadora)
<i>BMP</i>	<i>Bitmap</i> (mapa de bits)

<i>CD</i>	<i>Compact Disk</i> (Disco Compacto)
<i>CEMIE</i>	Centro Mexicano de Innovación en Energía
<i>CFE</i>	Comisión Federal de Electricidad
<i>CIMAT</i>	Centro de Investigación en Matemáticas ,
<i>CICESE</i>	Centro de Investigación Científica y de Educación Superior de Ensenada
<i>CINVESTAV</i>	Centro de Investigación y de Estudios Avanzados del Instituto Politécnico Nacional
<i>CPU</i>	<i>Central processing unit</i> Unidad Central de Procesamiento
<i>CRD</i>	Comisión Reguladora de Energía
<i>CSP</i>	<i>Common Spatial Pattern</i> (Patrón Común Espacial)
<i>CT</i>	<i>Computed Tomography</i> (Tomografía Computarizada)
<i>CUDA</i>	<i>Compute Unified Device</i> <i>Architecture</i> (Arquitectura de Dispositivo de Cómputo Unificado)
<i>DVD</i>	<i>Digital Video Disk</i> (Disco digital de Video)
<i>DBN</i>	<i>Deep-Belief Network</i> (Red de Creencia Profunda)
<i>DBN</i>	<i>Dynamic Bayesian Network</i> (Red Bayesiana Dinámica)
<i>DFT</i>	<i>Discrete Fourier Transform</i> (Transformada Discreta de Fourier)
<i>DL</i>	<i>Deep Learning</i> (Aprendizaje Profundo)

<i>DWT</i>	<i>Discrete Wavelet Transform</i> Transformada de Ondeletas Discreta
<i>EBCDIC</i>	<i>Extended Binary Coded Decimal Interchange Code</i> (Código Binario Estándar para Intercambio Decimal)
<i>ECG</i>	Electro-cardiografía
<i>EEG</i>	Electro-encefalografía
<i>ERP</i>	<i>Event-Related Potential</i> (Potencial relacionado a evento)
<i>EXCALE</i>	Exámen de la Calidad y el Logro Educativo
<i>FCM</i>	<i>Fuzzy Cognitive Map</i> (Mapa Difuso Cognitivo)
<i>FFT</i>	<i>Fast Fourier Transform</i> (Transformada Rápida de Fourier)
<i>FIR</i>	Finite-Impulse Response (Respuesta Finita al Impulso)
<i>FL</i>	<i>Fuzzy Logic</i> (Lógica difusa)
<i>FPGA</i>	Field Programmable Gate Array
<i>GA</i>	<i>Genetic Algorithm</i> (Algoritmo Genético)
<i>GPU</i>	<i>Graphic Processing Unit</i> Unidad Gráfica para Procesamiento
<i>HMM</i>	<i>Hidden Markov Models</i> (Modelos Ocultos de Markov)
<i>IA</i>	Inteligencia Artificial
<i>IC</i>	Inteligencia Computacional
<i>ICA</i>	<i>Independent Component Analysis</i> (Análisis de Componentes Principales)
<i>ICD</i>	<i>International Classification of Diseases</i>

	(Clasificación Internacional de Enfermedades)
<i>IEEE</i>	<i>The Institute of Electrical and Electronics Engineers</i> (Instituto de Ingenieros Eléctricos y Electrónicos.
<i>IIR</i>	<i>Infinite Impulse Response</i> (Respuesta Infinita al Impulso)
<i>INAOE</i>	Instituto Nacional de Astrofísica, Óptica y Electrónica
<i>INEE</i>	Instituto Nacional para la Evaluación de la Educación
<i>INEEL</i>	Instituto Nacional de Electricidad y Energías Limpias (antes Instituto de Investigaciones Eléctricas)
<i>IoT</i>	<i>Internet of Things</i> (Internet de las Cosas)
<i>IP address</i>	<i>Internet Protocol address</i> Dirección de Protocolo de Internet
<i>IPN</i>	Instituto Politécnico Nacional
<i>ITESM</i>	Instituto Tecnológico y de Estudios Superiores de Monterrey
<i>LSTM</i>	<i>Long Short-Term Memory</i> (Memoria Corta y de Largo alcance)
<i>MA</i>	<i>Moving average</i> (Promediado Móvil)
<i>MAP</i>	<i>Maximum a Posteriori Probability</i> (Probabilidad Máxima <i>a posteriori</i>)
<i>MC</i>	<i>Markov chains</i> (Cadenas de Markov)
<i>MCT</i>	<i>Memory Test Computer</i> (Computadora de Prueba de Memoria)

<i>MIT</i>	<i>Massachusetts Institute of Technology</i> (Instituto Tecnológico de Massachusetts)
<i>MODWT</i>	<i>Maximum-overlapping Discrete Wavelet transform</i> (Transformada de Ondeleta Discreta de Máximo Traslape)
<i>MRI</i>	<i>Magnetic resonance imaging</i> (Imagen de Resonancia Magnética)
<i>NB</i>	Naïve Bayes classifier (Clasificador Ingenuo de Bayes)
<i>OCT</i>	<i>Optical Coherence Tomography</i> OCT (Tomografía de Óptica Coherente)
<i>PCA</i>	<i>Principal Component Analysis</i> (Análisis de Componentes Principales)
<i>PDB</i>	<i>Protein Data Bank</i> (Banco de Datos de Proteínas)
<i>PET</i>	<i>Positron Emission Tomography</i> Tomografía de Emisiones de Positrones
<i>PPG</i>	<i>Photoplethysmography</i> (Fotopletismografía)
<i>PSCPP</i>	<i>Protein Side-Chain Packing Problem</i> (Problema de Empaquetado de las Cadenas Laterales de Proteínas)
<i>REI</i>	Redes Eléctricas Inteligentes
<i>RNA</i>	Redes Neuronales Artificiales
<i>SA</i>	<i>Simulated Annealing</i> (Recocido Simulado)
<i>SENER</i>	Secretaría de Energía
<i>SIDA</i>	Síndrome de Inmuno-deficiencia Adquirida

<i>SOM</i>	<i>Self-Organizing Map</i> (Mapa Auto-organizado)
<i>SQL</i>	<i>Structured Query Language</i> (Lenguaje de Consulta Estructurado)
<i>SVM</i>	<i>Support Vector Machine</i> (Máquina vectorial de soporte)
<i>STFT</i>	<i>Short-time Fourier Transform</i> (Transformada de Fourier de tiempo Corto)
<i>TDF</i>	Transformada Discreta de Fourier
<i>TIC</i>	Tecnologías de Información y Comunicaciones
<i>UAA</i>	Universidad Autónoma de Aguascalientes
<i>UABC</i>	Universidad Autónoma de Baja California
<i>UASLP</i>	Universidad Autónoma de San Luis Potosí
<i>UNAM</i>	Universidad Nacional Autónoma de México
<i>UP</i>	Universidad Panamericana
<i>UCI</i>	<i>Univeristy of California, Irvine</i> (Universidad de California, Irvine)
<i>VIH</i>	Virus de la Inmunodeficiencia Humana
<i>VIP</i>	<i>Visual Evoked Potential</i> (Potencial Visual Evocado)
<i>WHO</i>	<i>World Health Organization</i> (Organización Mundial de la Salud)

Índice general

Índice general	xv
1. Máquinas Capaces de Entender	I
1.1. Conceptos básicos de AS/RP	4
1.1.1. El diseño y uso de un clasificador	5
1.2. Presente y futuro de AS/RP en México	10
Referencias	12
2. Reconocimiento de Patrones	15
2.1. Introducción	15
2.2. Técnicas de Reconocimiento de Patrones	18
2.2.1. Clasificación Basada en Aprendizaje Supervisado	19
2.2.2. Clasificación No Supervisada	22
2.3. Enfoque Estadístico	24
2.3.1. Clasificador Ingenuo de Bayes (<i>Naïve Bayes</i>)	24
2.3.2. Modelos de Markov	25
2.3.3. Modelos Ocultos de Markov	27
2.4. Redes Neuronales Artificiales	29
2.4.1. Redes basadas en aprendizaje profundo	32

2.5.	Máquinas de Soporte Vectorial	33
2.6.	Conclusiones	36
	Referencias	37
3.	Señales Digitales	41
3.1.	Introducción	41
3.2.	Muestreo y Cuantización	44
3.3.	Señales y Sistemas de Tiempo Discreto	46
3.4.	Análisis Espectral de Señales	52
3.4.1.	Transformada de Fourier	52
3.4.2.	Estimación Espectral	54
3.5.	Análisis Tiempo-Frecuencia	59
3.5.1.	Espectrograma	59
3.5.2.	Transformada de Ondeletas	61
3.6.	Algunos proyectos de investigación	63
3.7.	Conclusiones	68
	Referencias	69
4.	Visión Computacional	73
4.1.	Introducción	73
4.2.	Antecedentes de la visión por computadora	74
4.3.	Procesamiento digital de imágenes	75
4.4.	Aplicaciones en Bioinformática	78
4.5.	Aportaciones internacionales	83
4.6.	Resultados recientes en México	98
	Referencias	103

5. La Minería de Datos	109
5.1. Qué es la minería de datos	109
5.1.1. Qué no es la minería de datos	111
5.1.2. Qué es la Ciencia de Datos	112
5.1.3. Riesgo: uso poco ético de la Ciencia de Datos	113
5.2. Usos de la minería de datos	113
5.2.1. Enfoques	114
5.2.2. Pasos claves en la minería de datos	115
5.2.3. Pilares de la minería de datos	117
5.3. Aplicaciones y ejemplos	118
5.4. El futuro de la minería de datos	131
5.4.1. Aplicaciones en el futuro cercano	132
5.5. Apéndice: Recursos recomendados	134
Referencias	136
 6. Administración de la Energía	 139
6.1. Diagnóstico transformadores	141
6.2. Ejemplos adicionales	149
6.2.1. Energía Eólica	149
6.2.2. Energía Solar	151
6.2.3. Energía Geotérmica	152
6.2.4. Energía Nuclear	154
6.2.5. Sistemas de Potencia	155
6.3. Tareas pendientes	158
Referencias	161

7. Patrones en Medicina y Biología	169
7.1. Introducción	169
7.1.1. Aplicaciones	172
7.1.2. Estructura y orientación del capítulo	173
7.2. Para desarrollo del conocimiento	173
7.2.1. Reconstrucción	175
7.2.2. Procesamiento	176
7.2.3. Análisis	178
7.2.4. Interpretación	179
7.3. Uso aplicado en medicina	180
7.3.1. Prevención	181
7.3.2. Apoyo al diagnóstico	182
7.3.3. Tratamiento	183
7.3.4. Seguimiento	184
7.4. Uso aplicado en biología	185
7.4.1. A nivel celular	186
7.4.2. A nivel organismo	186
7.4.3. A nivel ecosistema	187
7.4.4. Explotación industrial	187
7.5. Vanguardia del campo en México	188
7.6. Conclusiones	190
Referencias	190
 8. Interfaces Cerebro-Computadora	 193
8.1. Interfaz cerebro-computadora	194
8.1.1. Concepto y adquisición de señales cerebrales	194
8.1.2. Modelos matemáticos	198

8.2.	Aplicaciones	204
8.2.1.	Médica	204
8.2.2.	Entretenimiento	206
8.2.3.	Cognitiva	208
8.3.	Conclusiones	209
	Referencias	210

Capítulo I

Máquinas Capaces de Entender

MARÍA DEL PILAR GÓMEZ-GIL, INAOE

FELIPE ORIHUELA ESPINA, INAOE

La vida está llena de operaciones de percepción y reconocimiento. Los procesos fisiológicos, de decisión y cognitivos dependen de la sorprendente capacidad de nuestro organismo para captar lo que acontece en su exterior y reaccionar adecuadamente. Algo similar sucede con el resto de los seres vivos, desde aquellos que cuentan con las estructuras mas simples, como los virus o las bacterias, hasta los mas complejos. A nivel cognitivo, los humanos continuamente estamos realizando operaciones de aprendizaje, abstracción y generalización que nos permiten inferir patrones complejos, los cuales no fueron directamente observados por nuestro sistema sensorial. Por ejemplo, después de escuchar hablar a un amigo(a) por unos minutos, podemos intuir cuando está preocupado(a). Nuestras capacidades de reconocimiento de patrones nos han permitido desarrollar el arte, avanzar en el conocimiento científico y construir herramientas sofisticadas para vivir mejor.

Dado que la habilidad de reconocer representa una ventaja competitiva importante, nos hemos enfrascado en intentar construir máquinas capaces de percibir y entender de forma completamente autónoma, buscando incluso que realicen estas habilidades mucho mejor que nosotros; y parece que cada día estamos mas cerca de conseguirlo (Domingos, 2018). Actualmente, contamos con sistemas capaces de realizar actividades de reconocimiento que contienen sensores que les permiten “enterarse” de mediciones de todo tipo, adquiridas de forma prácticamente continua. El avance en la ciencia básica y aplicada ha permitido a los investigadores desarrollar cada día mejores algoritmos para el análisis de señales y reconocimiento de patrones (AS/RP), lo cual, aunado al desarrollo del cómputo paralelo y de alto rendimiento, ha puesto a nuestro alcance máquinas que no simplemente capturan grandes cantidades de datos y memorizan sus valores o rangos, sino que ejecutan complejos procesos de generalización y abstracción que les permiten llevar a cabo tareas de identificación o clasificación. La disminución en los costos del hardware y *firmware* ha cooperado notablemente a este auge, lo cual a su vez ha generado cantidades exorbitantes de datos que requieren ser resumidos e interpretados por algoritmos cada vez mas poderosos. El avance ha sido impresionante y ubicuo. Solo por nombrar algunos ejemplos, vemos que áreas como la visión computacional, la medicina, las ingenierías o la minería de datos utilizan actualmente en su quehacer diario aplicaciones de análisis de señales y reconocimiento de patrones que eran sueños apenas en los años 70's (Kanal, 1974).

Es interesante mirar hacia atrás en el tiempo y ver cómo es que la automatización ha llegado a este punto de sofisticación. Recordemos que algunas de las máquinas construidas en los siglos XVIII y XIX, durante la revolución industrial , mostraban los primeros intentos de automatizar la toma de decisio-

nes basada en la observación de las condiciones del medio ambiente. Sin embargo, fue hasta la aparición de las primeras computadoras digitales en los años 50, aunado a la publicación de los primeros modelos de aprendizaje automático (ver por ejemplo (Rosenblatt, 1958)), que se empezaron a tener verdaderas esperanzas de construir autómatas capaces de realizar eficientemente operaciones de reconocimiento de patrones. Oliver Selfridge, quien es considerado por algunos el padre de la Inteligencia Artificial (IA), introdujo en 1955 junto con su equipo de trabajo del *Massachusetts Institute of Technology* (MIT), el concepto de reconocimiento de patrones (Selfridge, 1955). Selfridge definió al reconocimiento de patrones como “la extracción de características significativas de un fondo de detalles irrelevantes.” Selfridge empezó su plática haciendo notar algo que sabemos muy bien: el reconocimiento de patrones es “la clase de cosas” que las personas hacen muy bien pero no así las computadoras (Selfridge, 1955). Es muy interesante observar que en aquellos años de inicio de los reconocedores, las computadoras digitales ofrecían grandes expectativas de uso. En el mismo congreso donde Selfridge presentó el concepto de “reconocimiento de patrones”, su equipo de trabajo presentó un reconocedor para distinguir entre letras y objetos básicos (Dinneen, 1955), el cual fue ejecutado en la computadora conocida como “MTC” (*Memory Test Computer*, Computadora de Prueba de Memoria), ubicada en el *Lincon Laboratory*. En la página oficial del Museo de Historia de las computadoras puede verse una foto de esta computadora. La capacidad de cómputo de la MCT era varios órdenes de magnitud menor que la que posee el teléfono celular mas básico hoy en día. En las décadas siguientes, disciplinas como la inteligencia artificial, las ciencias computacionales y la estadística continuaron desarrollando algoritmos cada vez más poderosos, a la vez que se empezaron a lanzar al mercado sistemas de procesamiento paralelo de alta

velocidad, así como dispositivos electrónicos capaces de llevar a cabo una amplia gama de mediciones.

En este libro, los investigadores miembros e invitados de la sección académica “Análisis de Señales y Reconocimiento de Patrones” (AS/RP de la Academia Mexicana de Computación (Amexcomp)) ofrecemos a los lectores una breve descripción de algunos temas de la ciencia y la tecnología que han permitido y continúan promoviendo el desarrollo de AS/RP. En este capítulo, damos una breve introducción a conceptos básicos asociados con AS/RP. En los capítulos 2 y 3, detallamos un poco más la teoría y los algoritmos más populares dentro de AS/RP. Posteriormente, incluimos varios capítulos para tratar algunos de los campos de investigación específicos en que se encuentran involucrados los investigadores de nuestra sección académica: la visión computacional, el reconocimiento de patrones en la minería de datos, el papel de AS/RP para la obtención, administración y aprovechamiento de la energía, AS/RP en medicina y biología y en el diseño de interfaces cerebro-computadora (BCI por sus siglas del inglés *Brain Computer Interface*).

1.1. Conceptos básicos de AS/RP

En el contexto de este libro, entendemos por *señal digital* a “un conjunto de magnitudes físicas medidas en un punto o intervalo de tiempo o espacio determinado” (Gómez-Gil, Guzmán-Arenas, Orihuela-Espina, Bribiesca, y Rascón, 2017). Las señales digitales pueden representar fenómenos físicos, económicos, sociales, biológicos etc. Las señales se analizan a fin de entender su contenido informativo y aprovecharlo para alguna tarea específica o para aumentar nuestro conocimiento sobre ellas.

Un *patrón* representa una relación entre dichas señales, la cual se obtiene a través de técnicas matemáticas. Normalmente un patrón no se representa con todos los valores obtenidos del dispositivo sensor, sino que se utiliza una depuración de dicho conjunto. Con esta depuración, realizada a través de técnicas de procesamiento de señales, se intenta eliminar toda la información irrelevante para la tarea que desea realizarse. A este conjunto de valores depurados, se le conoce como *vector de características* y puede estar formado de distintos tipos de datos, por ejemplo números reales, enteros o valores enumerados. Estas representaciones se pueden usar para realizar tareas discriminatorias, explicativas o predictivas, a través del uso de clasificadores (Gómez-Gil y cols., 2017). Los patrones pueden tener asociadas etiquetas que identifican una *clase* o pertenencia a un grupo. Clasificar consiste entonces en asignar una etiqueta o valor de pertenencia de un patrón a una categoría determinada. Matemáticamente hablando, clasificar consiste en dividir un espacio n -dimensional en un número de regiones, donde cada región en este espacio corresponde a una clase determinada (Tou y Gonzalez, 1977). La dimensión del espacio (n) corresponde al número de elementos que caracterizan al patrón. Cuando un reconocedor se construye a través de ajustar sus parámetros en base a información de los patrones a reconocer, se le conoce como “Clasificador adaptivo.”

1.1.1. El diseño y uso de un clasificador

Como sucede con otras áreas del conocimiento, siempre hay un debate entre si el reconocimiento de patrones es una ciencia, una ingeniería o ambas cosas. Suena razonable la posición de Duin (Duin y Pekalska, 2007), quien considera que el reconocimiento de patrones es una ingeniería cuando se aplica a la construcción de sistemas que imitan o ayudan a las personas a reconocer, y una ciencia cuando

se estudia la manera en que los seres vivos distinguen o caracterizan patrones. Según el mismo autor, ambos enfoques se benefician uno del otro, lo cual es comprobado por las miles de aplicaciones que actualmente sirven para apoyar la investigación de los procesos cognitivos y biológicos. Mas allá de la opinión de Duin, es claro que el reconocimiento de patrones también es una ciencia exacta cuando trata con la definición de teorías matemáticas y de cómputo, que dan base a la construcción de sistemas adaptivos de reconocimiento y clasificación.

Podemos distinguir dos etapas asociadas a los sistemas de reconocimiento y clasificación: diseño y uso (Gómez-Gil y cols., 2017). Diseñar un reconocedor consiste en analizar las condiciones del problema, decidir la mejor estrategia y producir las reglas de decisión que permitirán el reconocimiento. Hay tres módulos importantes en un reconocedor: pre-procesamiento, extracción de características y clasificación. El pre-procesamiento consiste en limpiar los datos de artefactos ajenos a la información en cuestión, como puede ser el ruido, información incompleta, exceso de datos en alguna de las clases, etc. En el capítulo 3 se detallan algunas de las técnicas útiles para esta tarea. La extracción de características consiste en manipular los datos de tal manera que se obtenga el mínimo de información relevante de cada uno, que permita diferenciarlo de otros. La clasificación corresponde a los algoritmos de decisión. El proceso de diseño incluye una validación del desempeño del reconocedor, es decir, una medición de su comportamiento con datos no utilizados para construirlo (capacidad de generalización). Si resulta una evaluación insatisfactoria, el diseño deberá modificarse; el proceso se repite hasta obtener un desempeño aceptable. En la figura 1.1 se muestra este proceso de diseño. En el capítulo 2 se explican con detalle algunos algoritmos de decisión y validación de reconocedores.

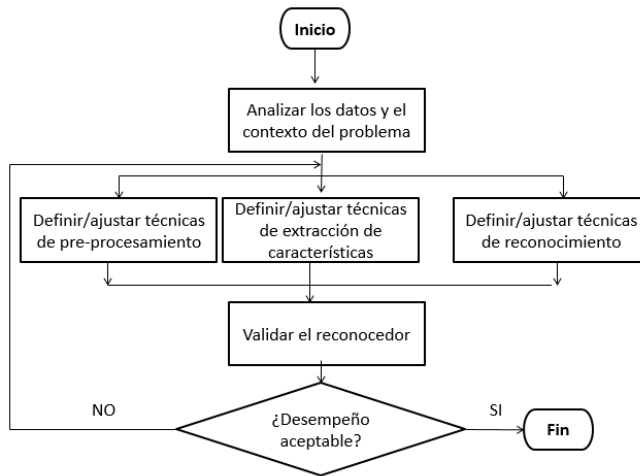


Figura 1.1: Pasos para diseñar un clasificador (Gómez-Gil y cols., 2017).

Para reconocer o clasificar un patrón se deben realizar las siguientes acciones (ver figura 1.2): primero, adquirir la señal producida por el patrón, a través de un sensor. Posteriormente, es necesario pre-procesar la señal para poder obtener las características que representan al objeto en cuestión. Enseguida se ejecuta el reconocedor utilizando el vector de características obtenido en el paso anterior, y de acuerdo a lo definido en el sistema, se decide la clase a la que pertenece.

Hay muchas maneras de diseñar reconocedores de patrones. Lo primero que viene a la mente es el hacer un programa *ad-hoc*, es decir, que se acomode exactamente al tipo de problema que se enfrenta; pero esto puede resultar demasiado limitado o servir solamente para casos particulares. Afortunadamente, actualmente existen una gran variedad de arquitecturas así como técnicas de reconocimiento de patrones que pueden utilizarse. A pesar de ello, la tarea de escribir reconocedores que puedan generalizar no es trivial, ya que no siempre es evidente cual es el tipo de clasificador que mejor se adapta al problema.

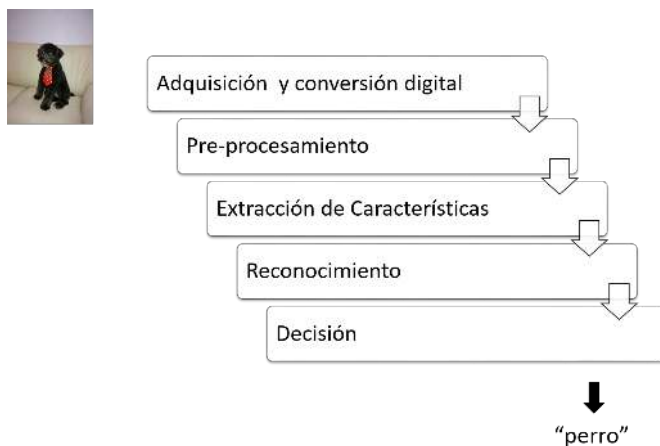


Figura 1.2: Etapas del proceso de clasificación (Gómez-Gil y cols., 2017).

Se pueden considerar 4 diferentes enfoques en las técnicas de reconocimiento de patrones (Duin y Pekalska, 2007):

Modelado de conceptos. Este enfoque incluye a los sistemas expertos, las redes de creencia y las redes probabilistas. Los sistemas expertos buscan capturar el conocimiento de los expertos humanos y colocarlo a través de reglas, en un sistema automático de toma de decisiones. En (Pineda y Brena, 2017), disponible en línea aquí, puede encontrarse una amena explicación de estos conceptos.

Modelado de objetos. Este enfoque incluye reconocimiento de patrones sintáctico, reconocimiento de patrones estructural y razonamiento basado en casos. El reconocimiento de patrones sintáctico busca relaciones estructurales de los patrones, usando gramáticas formales; generalmente se utiliza cuando los datos no están representados con un vector de características sino con grafos o cadenas de símbolos. (Ochoa y Trinidad, 2011). El razonamiento basado en casos es una técnica de IA que busca solucionar problemas re-utilizando

o modificando las soluciones aplicadas en situaciones anteriores (Mántaras, 2011).

Modelado de sistemas. Este enfoque incluye a las redes neuronales y los sistemas de visión. Las redes neuronales son sistemas de aprendizaje por ejemplos, inspirados en las neuronas biológicas, capaces de representar conocimiento por medio de números para posteriormente generalizar y abstraer de dicho conocimiento. En el capítulo 2 se discute un poco mas sobre las redes neuronales y en el capítulo 4 sobre visión computacional.

Generalización. Este enfoque incluye inferencia gramatical y reconocimiento de patrones estadístico. La inferencia gramatical busca aprender un lenguaje a partir de datos; involucra varias áreas, entre ellas aprendizaje automático, reconocimiento de patrones y teoría de lenguajes formales (Aguilar, Alonso, López, y Moreno, 2005). El reconocimiento de patrones estadístico busca estimar funciones de probabilidad y de distribución a partir de los datos. (Duda, Hart, y Stork, 2000)

Con respecto a las aplicaciones del reconocimiento de patrones, es claro que estamos hablando de una disciplina completamente ubicua, que tiene aplicación en prácticamente cualquier área de las ciencias exactas, naturales, sociales o de ingeniería. Por otro lado, su estudio es interdisciplinario, ya que toma conocimientos desarrollados en otras áreas. También es importante aclarar que el análisis de señales se aplica no solamente para apoyar el reconocimiento de patrones, sino en un también en otras áreas de ciencia e ingeniería.

1.2. Presente y futuro de AS/RP en México

La investigación de AS/RP en México ha tenido un auge importante en los últimos años. Según el sistema internacional de evaluación de publicaciones por instituciones, conocido como Scimago, en Noviembre de 2018 México ocupaba el segundo lugar en los países de América Latina en cuanto al mayor número de publicaciones asociadas a Ciencias de la Computación, en las especialidades de Visión por Computadora, Reconocimiento de Patrones e Inteligencia Artificial. El/la lector(a) interesado(a) puede consultar este valor al día de hoy en esta liga. Los principales trabajos de la comunidad mexicana dedicada a AS/RP se han enfocado en los últimos años principalmente a las siguientes áreas (Gómez-Gil y cols., 2017):

- reconocimiento de patrones en grandes bases de datos,
- reconocimiento basado en Redes Neuronales Artificiales,
- visión y etiquetado de imágenes basado en aprendizaje profundo,
- interfaces cerebro computadora,
- análisis y predicción de series de tiempo,
- visión computacional,
- audición robótica,
- vehículos autónomos,
- identificación de componentes defectuosos y en la manufactura,
- diagnóstico médico sobre imágenes y señales uni-dimensionales,

- algoritmos de clasificación y regresión para navegación y control de robots,
- sistemas de apoyo a la rehabilitación,
- biometría bimodal y multimodal.

En (Gómez-Gil y cols., 2017), disponible en línea aquí, se puede encontrar una lista detallada de citas a trabajos realizados en los últimos años por investigadores mexicanos, relacionados a estas áreas. Asimismo, en este documento producido por Amexcomp a finales de 2017, se listan los principales enfoques de investigación, proyectos y publicaciones de los miembros de la sección AS/RP de la Amexcomp.

La mayoría de las instituciones de enseñanza superior y centros de investigación del país tienen proyectos relacionados con AS/RP. En el libro “La computación en México por especialidades académicas” publicado en 2017 por Amexcomp (Pineda-Cortés, 2017), se identificaron a las siguientes instituciones que mayor aporte realizan a la investigación en esta área :

- Centro de Investigación en Matemáticas (CIMAT),
- Centro de Investigación y de Estudios Avanzados del Instituto Politécnico Nacional (CINVESTAV),
- Instituto Nacional de Astrofísica, Óptica y Electrónica (INAOE),
- Instituto Nacional de Electricidad y Energías Limpias (INEEL) (antes Instituto de Investigaciones Eléctricas),
- Instituto Politécnico Nacional (IPN),

- Instituto Tecnológico y de Estudios Superiores de Monterrey (ITESM),
- Universidad Nacional Autónoma de México (UNAM),
- Universidad Autónoma de Baja California (UABC),
- Universidad Autónoma de San Luis Potosí (UASLP)
- Universidad Panamericana (UP).

El estudio y aplicación AS/RP ofrece amplias oportunidades para el desarrollo de la industria mexicana. En (Gómez-Gil y cols., 2017) se encuentra una vasta lista de aplicaciones de AS/RP desarrolladas en México. Desde la Amexcomp, consideramos fundamental que la enseñanza de algunos conocimientos básicos asociados a AS/RP se incluya en los planes de estudio de licenciaturas, ya que muchos profesionales actualmente utilizan de forma cotidiana sistemas basados en reconocimiento de patrones. En los siguientes capítulos hemos intentado presentar estos conceptos básicos que deseamos sirvan de pauta a los lectores interesados para entender que es el reconocimiento de patrones aplicado al contexto de señales digitales, y de que manera puede beneficiar a la sociedad.

Referencias

Aguilar, R., Alonso, L., López, V., y Moreno, M. N. (2005). Descubrimiento incremental de patrones secuenciales para la inferencia de gramática. En *Actas del III taller nacional de minería de datos y aprendizaje, tamida2005* (p. 85-95). Descargado de <http://www.lsi.us.es/redmidas/CEDI/papers/409.pdf>

- Dinneen, G. P. (1955). Programming pattern recognition. En *Proceedings of the march 1-3, 1955, Western Joint Computer Conference* (pp. 94-100). New York, NY, USA: ACM. Descargado de <http://doi.acm.org/10.1145/1455292.1455311> doi: 10.1145/1455292.1455311
- Domingos, P. (2018). Our digital doubles. *Scientific American*, 88-93.
- Duda, R. O., Hart, P. E., y Stork, D. G. (2000). *Pattern classification (2nd edition)*. New York, NY, USA: Wiley-Interscience.
- Duin, R. P. W., y Pekalska, E. (2007). The science of pattern recognition. achievements and perspectives. En W. Duch y J. Mańdziuk (Eds.), *Challenges for computational intelligence* (pp. 221-259). Berlin, Heidelberg: Springer Berlin Heidelberg. Descargado de https://doi.org/10.1007/978-3-540-71984-7_10 doi: 10.1007/978-3-540-71984-7_10
- Gómez-Gil, P., Guzmán-Arenas, A., Orihuela-Espina, F., Bribiesca, E., y Rascón, C. (2017).
Análisis de señales y reconocimiento de patrones. En L. Pineda-Cortés (Ed.), *La computación en México por especialidades académicas* (pp. 233-265). Academia Mexicana de la Computación.
- Kanal, L. (1974, November). Patterns in pattern recognition: 1968-1974. *IEEE Transactions on Information Theory*, 20(6), 697-722. doi: 10.1109/TIT.1974.1055306
- Mántaras, R. L. D. (2011, Julio-Diciembre). Razonamiento basado en casos: ejemplos de aplicaciones poco convencionales. *Komputer Sapiens*, 3(2), 10-14. Descargado de http://smia.mx/komputersapiens/index.php?option=com_content&view=article&id=72&Itemid=84
- Ochoa, J. A. C., y Trinidad, J. F. M. (2011, Julio-Diciembre). Reconocimiento de patrones. *Komputer Sapiens*, 3(2), 5-9. Descarga-

- do de http://smia.mx/komputersapiens/index.php?option=com_content&view=article&id=72&Itemid=84
- Pineda, L., y Brena, R. (2017). Conocimiento y razonamiento. En L. Pineda-Cortés (Ed.), *La computación en México por especialidades académicas* (pp. 33–63). Academia Mexicana de la Computación.
- Pineda-Cortés, L. (2017). *La computación en México por especialidades académicas*. Ciudad de México, México: Academia Mexicana de la Computación.
- Rosenblatt, F. (1958). The perceptron: A probabilistic model for information storage and organization in the brain. *Psychological Review*, 65(6), 386–408. Descargado de <http://dx.doi.org/10.1037/h0042519>
- Selfridge, O. G. (1955). Pattern recognition and modern computers. En *Proceedings of the march 1-3, 1955, Western Joint Computer Conference* (pp. 91–93). New York, NY, USA: ACM. Descargado de <http://doi.acm.org/10.1145/1455292.1455310> doi: 10.1145/1455292.1455310
- Tou, J., y Gonzalez, R. C. (1977). *Pattern recognition principles* (Second ed.). Addison-Wesley.

Capítulo 2

Reconocimiento de Patrones

RIGOBERTO FONSECA-DELGADO UNIV. YACHAY TECH

MARÍA DEL PILAR GÓMEZ-GIL, INAOE

JUAN MANUEL RAMÍREZ-CORTÉS, INAOE

DUSTIN CARRIÓN-OJEDA, UNIV. YACHAY TECH

2.1. Introducción

Este capítulo cubre conceptos básicos y algunas estrategias para diseñar sistemas basados en reconocimiento de patrones. El reconocimiento de patrones tiene por objetivo la clasificación de objetos en categorías o clases; los objetos pueden ser imágenes, audio, documentos, señales de sensores u algún otro tipo de medida utilizable por una computadora. El término *patrón* normalmente hace referencia a estos objetos o a alguna representación de ellos. Actualmente, el reconocimiento de patrones es una parte sustantiva de los sistemas de toma de decisiones basados en Inteligencia Artificial (IA) (Theodoridis y Koutroumbas, 2009). La IA es el área de computación que trata con algoritmos que permi-

ten a una máquina el percibir, razonar y actuar de forma similar a un ser humano (Winston, 2010). Para los lectores interesados en profundizar en IA, se recomienda el curso en línea *Intro to IA* de la Universidad de California en Berkeley.

El *reconocimiento de patrones* se define como “la ciencia que se ocupa de los procesos de ingeniería, computación y matemáticas, relacionados con objetos físicos y/o abstractos, denominados patrones, con el propósito de extraer información que permita establecer propiedades de o entre conjuntos de dichos objetos, los cuales nos permite interpretar el mundo que nos rodea” (Carrasco-Ochoa y Martínez-Trinidad, 2011). Normalmente, los objetos se representan a través de características observadas en ellos, en la forma vectorial (Theodoridis y Koutroumbas, 2009):

$$x = [x_1, x_2, \dots, x_l]^T$$

donde el vector x está formado por las características x_i , $i = 1, 2, \dots, l$; el operador T representa la operación matricial transpuesta.

Por ejemplo, consideremos el repositorio de bases de datos comúnmente conocido como (UCI): *UCI Machine Learning Repository*¹ (Repositorio de Aprendizaje de Máquina de la Universidad de California, Irvine), creado por el “Centro para aprendizaje de máquina y sistemas inteligentes”², de la Universidad de California en Irvine, que contiene una gran variedad de problemas de reconocimiento de patrones. Un ejemplo muy popular de este repositorio es el conjunto de datos llamado “Iris” (Fisher, 1936), que contiene características correspondientes a 3 tipos o clases de una planta llamada Iris, que son: “Setosa”, “Versicolor” y “Virginica”; las características son: largo del sépalo, ancho del sé-

¹<https://archive.ics.uci.edu/ml>

²<https://cml.ics.uci.edu/>

palo, largo del pétalo y ancho del pétalo, expresados en centímetros. La tabla 2.1 muestra cuatro ejemplos de estas instancias, sus características y la clase a la que pertenecen.

Tabla 2.1: Ejemplos de instancias de la base de datos Iris, tipos de plantas y sus atributos.

Largo Sépalo [cm]	Ancho Sépalo [cm]	Largo Pétalo [cm]	Ancho Pétalo [cm]	Clase
5.1	3.5	1.4	0.2	Iris-setosa
4.9	3.0	1.4	0.2	Iris-setosa
7.0	3.2	4.7	1.4	Iris-versicolor
6.3	3.3	6.0	2.5	Iris-virginica

La figura 2.1 muestra cuatro gráficas, una por cada característica en Iris. Las clases en estas gráficas están representadas usando “cuadrados” para Setosa, “equis” para Versicolor y “círculo” para Virginica; el eje X corresponde a los valores de las características de las diferentes instancias. Allí puede apreciarse que utilizando el largo o ancho del pétalo es posible distinguir instancias de la clase Setosa del resto de las clases, pero no es posible diferenciar las clases Versicolor ni Virginica. Entonces, para poder clasificar todos los tipos de la planta es necesario considerar varias características.

La figura 2.2 muestra seis gráficas, utilizando diferentes parejas de características especificadas en los ejes horizontal y vertical. La primera gráfica utiliza el largo y ancho del sépalo y puede notarse que las instancias de Versicolor e Virginica están traslapadas, pero utilizando el ancho del pétalo y el largo del sépalo, pueden separarse mejor las instancias de cada clase.

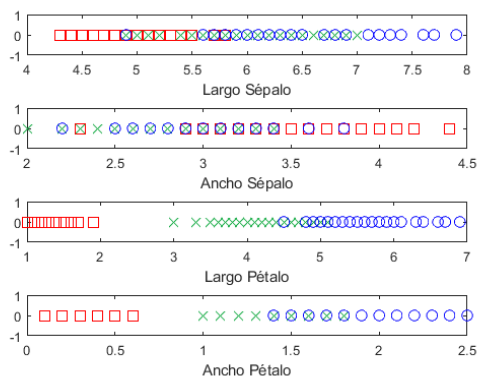


Figura 2.1: Instancias de cada tipo de plantas representadas usando una característica por gráfica.

Entonces, a partir de las características apropiadas, es posible encontrar una función que distinga entre las clases. La función que relaciona los vectores de características de una instancia con la clase respectiva se conoce como función de clasificación o clasificador.

2.2. Técnicas de Reconocimiento de Patrones

Existen una gran cantidad de técnicas para construir reconocedores que pueden agruparse en varias formas; entre las mas populares se encuentran las siguientes (Duda, Hart, y Stork, 2000):

Estimación paramétrica. Se basa en la aproximación de la probabilidad condicional de que un patrón pertenezca a una clase, asumiendo alguna forma para dicha probabilidad y suponiendo el número de clases presentes.

Estimación no paramétrica. Estos métodos realizan la clasificación sin asumir el número de clases ni la forma de las distribuciones de probabilidad.

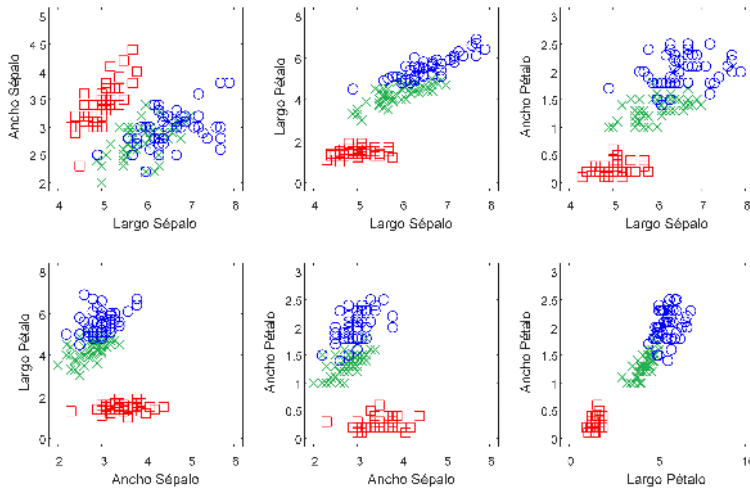


Figura 2.2: Visualizando las instancias con diferentes pares de características.

Métodos estocásticos. Usan funciones aleatorias para diseñar las funciones de clasificación.

Métodos basados en aprendizaje supervisado/no supervisado. Utilizan modelos de aprendizaje de máquina para definir a las funciones. Este aprendizaje se lleva a cabo a través de ejemplos de los patrones a clasificar, lo cuales se conocen como el conjunto de entrenamiento. Cuando se conoce de antemano el valor de clase de los datos en el conjunto de entrenamiento, se dice que el aprendizaje es supervisado; en caso contrario el aprendizaje se denomina “no supervisado”.

2.2.1. Clasificación Basada en Aprendizaje Supervisado

La clasificación supervisada recibe como entrada un conjunto de datos previamente etiquetados y busca construir una función discriminadora. Por ejemplo,

la figura 2.3 muestra una función de clasificación compuesta de la función no-lineal f_1 y la función lineal f_2 para la base de datos Iris.

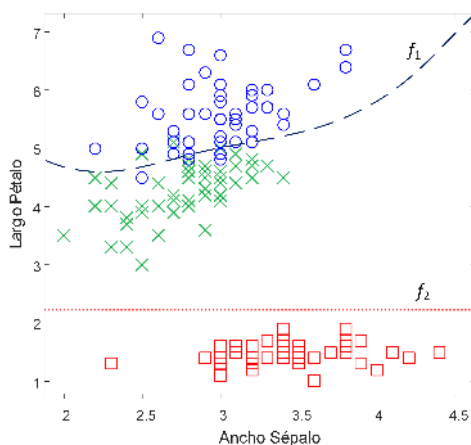


Figura 2.3: Instancias en Iris de las tres clases representadas utilizando dos características.

El desempeño de los clasificadores supervisados usualmente se mide en función del número de ejemplos que se clasifican correctamente. Si se tiene solamente una clase, el clasificador debe diferenciar si una instancia pertenece o no a dicha clase. El clasificador puede equivocarse al asignar un ejemplo a una clase cuando el ejemplo no pertenece a ésta; este error se conoce como *falso positivo* (fp). Cuando el clasificador rechaza un ejemplo de una clase, y resulta que el ejemplo si pertenece a la clase el error se denomina *falso negativo* (fn) (Manning, Raghavan, y Schutze, 2009). La tabla 2.2 muestra estos tipos de errores y así como los nombres que reciben las clasificaciones correctas.

Existen varias métricas para evaluar el desempeño de un clasificador binario basadas en este tipo de errores; entre las mas utilizadas se encuentran la precisión,

Tabla 2.2: Errores y aciertos posibles de un clasificador binario, dada la observación real.

		Observación Real	
		Verdadero	Falso
Resultado Clasificador	Verdadero	Verdadero Positivo (vp)	Falso Positivo (fp)
	Falso	Falso Negativo (fn)	Verdadero Negativo (vn)

el recuerdo, la exactitud y la medida “F”, las cuales se definen a continuación haciendo referencia a las variables definidas en la tabla 2.2:

La **precisión** es la fracción de los ejemplos correctamente clasificados sobre el total de ejemplos clasificados:

$$P = \frac{vp}{(vp + fp)} \quad (2.1)$$

El **recuerdo** es la fracción de los ejemplos correctamente clasificados que han sido identificados como tales:

$$R = \frac{vp}{(vp + fn)} \quad (2.2)$$

La **exactitud** representa a la fracción de clasificaciones correctas con respecto al total posible.

$$\text{exactitud} = \frac{(vp + vn)}{(vp + fp + fn + vn)} \quad (2.3)$$

F-Measure es una medida del balance entre precisión y recuerdo y puede ajustarse a las necesidades del problema a tratar. *F-Measure* se define como un promedio pesado de precisión y recuerdo:

$$F = \frac{1}{\alpha \frac{1}{P} + (1 - \alpha) \frac{1}{R}} = \frac{(\beta^2 + 1) PR}{\beta^2 P + R} \quad (2.4)$$

donde $\beta^2 = \frac{1-\alpha}{\alpha}$; $\alpha \in [0, 1]$, por lo que $\beta \in [0, \infty]$. Para asignar la misma importancia a la precisión y al recuerdo, $\alpha = 1/2$ o $\beta = 1$. En este caso F se denomina “balanceada” y se escribe F_1 o $F_{\beta=1}$ y:

$$F_{\beta=1} = \frac{2PR}{P + R} \quad (2.5)$$

2.2.2. Clasificación No Supervisada

La clasificación no supervisada inicia con un conjunto de ejemplos sin etiquetar y su objetivo es agruparlos en k clases, donde el número k puede establecerse previamente; se busca que los elementos dentro de un grupo sean más similares entre si y más diferentes con respecto a los elementos de otro grupo (Theodoridis y Koutroumbas, 2009). La Figura 2.4(a) muestra datos utilizando dos características, y (b) muestra el resultado de un agrupamiento de los datos en tres grupos A , B y C . En la figura (b) puede apreciarse que la distancia entre los elementos miembros de un grupo (flecha punteada) es menor comparada con la distancia entre elementos de grupos diferentes (flecha continua). Dependiendo del parámetro k y del método de agrupamiento elegido, los grupos obtenidos pueden variar significativamente.

Existe una gran cantidad de algoritmos para realizar agrupamientos. Uno de los mas utilizados es el de “k-medias” (*k-Means* en inglés), el cual consiste en escoger al azar k elementos del conjunto, y calcular la distancia de todos los demás elementos a éstos. Cada elemento es agrupado con uno de los escogidos que le es mas cercano. Posteriormente se calcula el promedio de cada grupo y con esto se genera un nuevo representante, llamado “centroide”. Esto se repite iterativamente hasta que los centroides no varían (Morales, Reyes, Morales-Reyes, y Escalante, 2107). Otros métodos populares de agrupamiento son el k-means difu-

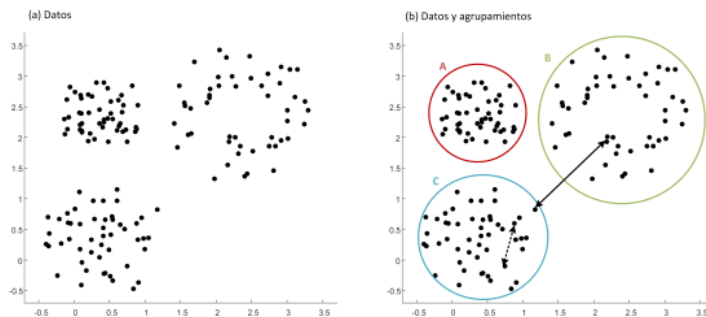


Figura 2.4: (a) Datos originales, (b) Datos originales agrupados usando circunferencias.

so, el aprendizaje no supervisado Bayesiano, el análisis de componentes principales y las redes neuronales auto-organizadas (Duda y cols., 2000; Haykin, 2009).

Existen diferentes métricas para evaluar la calidad del agrupamiento no supervisado. Algunas de estas se basan simplemente en comparar el agrupamiento obtenido con otro que se asume como verdadero. Otra opción es evaluar la calidad de los grupos en función de la distancia con el centroide o el valor promedio de la “silueta” que forma cada grupo. Una forma indirecta de evaluar agrupamientos es a través de los resultados obtenidos en la tarea para la que se requieren los grupos. Algunas definiciones de estas métricas pueden encontrarse en (Vilain, Burger, Aberdeen, Connolly, y Hirschman, 1995; Strehl, Ghosh, y Mooney, 2000).

2.3. Enfoque Estadístico para el Reconocimiento de patrones

A partir de los datos, es posible estimar probabilidades y distribuciones de los objetos de interés, utilizando estas estimaciones para construir un clasificador. En esta sección se presentan los métodos estadísticos mas populares.

2.3.1. Clasificador Ingenuo de Bayes (*Naïve Bayes*)

El clasificador de Naïve Bayes (NB), basado en el teorema de Bayes, asume una hipótesis de independendencia entre las características o variables predictoras; debido a esta hipótesis recibe el calificativo de Naïve que significa “ingenuo” (Duda y cols., 2000). El teorema de Bayes define la probabilidad de una clase C dado un conjunto de características $x_1, \dots, x_n$ como:

$$P(C|x_1, \dots, x_n) = \frac{P(C)P(x_1, \dots, x_n|C)}{P(x_1, \dots, x_n)} \quad (2.6)$$

sustituyendo el numerador por su representación como probabilidad condicional, y asumiendo que $P(x_i|C, x_j) = P(x_i|C)$, esta ecuación puede escribirse como:

$$P(C|x_1, \dots, x_n) = \frac{P(C) \prod_{i=1}^n P(x_i|C)}{P(x_1, \dots, x_n)} \quad (2.7)$$

NB utiliza una regla de clasificación en función de la clase o hipótesis más probable, conocido como “máximo a posteriori” (MAP). La función de clasificación puede expresarse como (Mitchell, 1997):

$$\hat{C} = \arg \max_{c \in C} P(c) \prod_{i=1}^n P(x_i|c) \quad (2.8)$$

Diferentes modelos de clasificación NB pueden crearse cambiando la función de distribución aplicada sobre $P(x_i|c)$; al trabajar con datos continuos se puede asumir una distribución normal, mientras que cuando los datos son discretos es más común utilizar una distribución multinomial o de Bernoulli.

La gran ventaja de NB es que puede entrenarse usando pocos datos y es capaz de tolerar características irrelevantes. Sin embargo, la suposición de independencia entre las características es algo que no siempre se presenta, puesto que en la mayoría de los problemas existe una correlación entre las mismas. Aun así, NB puede conseguir un buen desempeño.

El método NB, así como muchos otros clasificadores, están implementados en el proyecto de código abierto conocido como Weka.

2.3.2. Modelos de Markov

Los modelos de Markov son modelos estocásticos basados en la suposición de que los estados futuros de un sistema solo dependen de su estado presente, lo cual se conoce como propiedad de Markov. Los modelos de Markov más comunes son: Cadenas de Markov, Modelos Ocultos de Markov, Procesos de Decisión de Markov y Procesos de Decisión de Markov Parcialmente Observables (Rabiner, 1989).

Las cadenas de Markov (MC por sus siglas en inglés) están compuestas por una secuencia de variables aleatorias llamadas estados $(q_1, q_2, \dots, q_n)$, que poseen la propiedad de Markov. Por lo tanto, la probabilidad de transición entre estados se define como:

$$P(X_{n+1} = q_{n+1} | X_1 = q_1, X_2 = q_2, \dots, X_n = q_n) = P(X_{n+1} = q_{n+1} | X_n = q_n) \quad (2.9)$$

Cada estado de una MC cuenta con una distribución de probabilidad inicial, misma que puede representarse como grafo dirigido o tabla. La figura 2.5 y

el cuadro 2.3 representan una MC para la predicción del estado del clima; por ejemplo si hoy está soleado, la probabilidad de que mañana esté soleado es del 85 % ya que el estado futuro solo depende del estado actual (Rabiner, 1989). Conociendo la distribución inicial de probabilidades, las MC pueden realizar predicciones complejas.

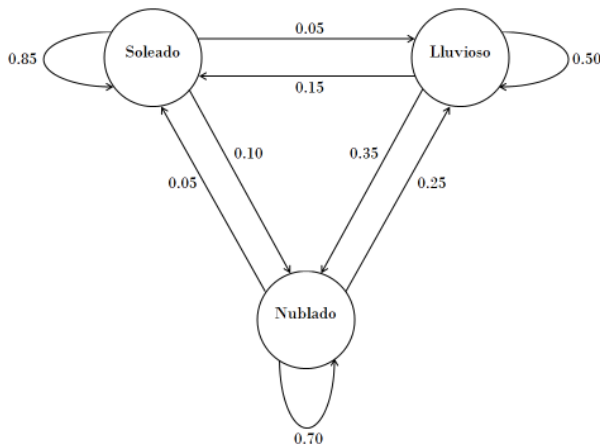


Figura 2.5: Representación gráfica de una cadena de Markov, que muestra los estados y sus posibles transiciones.

Tabla 2.3: Representación tabular de una cadena de Markov, que muestra los estados y sus posibles transiciones.

	Soleado	Lluvioso	Nublado
Soleado	0.85	0.05	0.10
Lluvioso	0.15	0.50	0.35
Nublado	0.05	0.25	0.70

2.3.3. Modelos Ocultos de Markov

Los modelos ocultos de Markov (*HMM "Hidden Markov Models"*), propuestos por L E. Baum (Baum, Petrie, Soules, y Weiss, 1970), son altamente populares debido a su simplicidad y capacidad inferencial. En los HMM se asume que el sistema a modelar es una MC, en el cual los estados están ocultos, es decir, no se pueden observar directamente, por lo que se analizan las variables visibles que son dependientes de cada estado. Los estados tienen una distribución de probabilidad sobre las variables visibles, de las cuales se extrae información sobre los estados ocultos. En el curso "Análisis de sistemas probabilísticos y probabilidad aplicada" disponible en el Sistema de Cursos Abiertos del Instituto Tecnológico de Massachusetts, se encuentra este video, que explica el funcionamiento de los modelos de Markov.

Un HMM se representa usando los siguientes parámetros (Kouemou, 2011):

- Conjunto de estados $Q = \{1, 2, \dots, N\}$.
- Posibles variables observables $V = \{v_1, v_2, \dots, V_M\}$.
- Distribución inicial de probabilidades $\pi = \{\pi_1, \pi_2, \dots, \pi_n\}$.
- Probabilidades de transición $A = \{a_{ij}\}$. Cada entrada de la matriz A se define como: $a_{ij} = P(q_{t+1} = j | q_t = i)$.
- Probabilidades de las variables observadas $B = \{b_j(v_k)\}$. Cada entrada del conjunto B se define como: $b_j(v_k) = P(o_t = v_k | q_t = j)$.

Los HMM pueden escribirse como una tupla $\lambda = (Q, V, \pi, A, B)$; la secuencia de observaciones se denota como $O = (o_1, o_2, \dots, o_T)$. La figura 2.6 ilustra los parámetros de un modelo oculto de Markov.

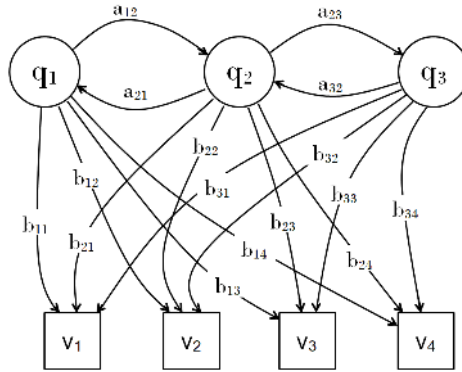


Figura 2.6: Representación gráfica de un modelo oculto de Markov con tres estados y cuatro posibles observaciones.

Los HMM generalmente se utilizan para atacar los siguientes tres tipos de problemas clásicos:

- Evaluación. Dado un modelo λ y una secuencia de observaciones O ¿Cuál es la probabilidad de que λ genere O ($P(O|\lambda)$)?
- Decodificación. Dado un modelo λ y una secuencia de observaciones O ¿Cuál es la secuencia de estados Q más propensa a producir O ?
- Aprendizaje. Dado un modelo λ y un grupo de secuencias observaciones $O' = (O_1, O_2, \dots, O_n)$ ¿Qué parámetros son los óptimos para generar dichas secuencias O' ?

Existen una vasta cantidad de algoritmos basados en modelos de Markov; referencias introductorias pueden encontrarse en (Rabiner, 1989; Sucar, Morales, y Hoey, 2011; Fink, 2014).

2.4. Redes Neuronales Artificiales

Las redes neuronales artificiales (RNA) forman parte de un área de la IA conocida como “Inteligencia Computacional” (IC), la cual ha colaborado en forma importante para el éxito del reconocimiento de patrones. IC se enfoca a estudiar, diseñar e implementar modelos que han sido inspirados biológica o lingüísticamente, incluyendo, además de RNA, a los algoritmos genéticos, la programación evolutiva, los sistemas difusos y los sistemas híbridos inteligentes. Una explicación sobre conceptos básicos de IC se puede encontrar en (Morales y cols., 2107). En la pagina <https://cis.ieee.org/> hay información detallada sobre implementaciones exitosas de clasificación usando RNA y otros modelos de IC.

Las RNA se utilizan para construir las funciones discriminadoras y de clasificación, cuando no son evidentes las reglas de decisión del sistema. Las RNA son modelos matemáticos inspirados en el comportamiento de las neuronas biológicas. Aunque dichos modelos están lejos de realmente modelar al sistema nervioso de los seres vivos, presentan características que les permiten aprender a través de ejemplos, además de abstraer y generalizar conceptos. Haykin (Haykin, 2009) define a una RNA como “un sistema masivamente paralelo y distribuido, hecho de unidades procesadoras simples, las cuales son de manera natural propensas a almacenar conocimiento adquirido de la experiencia y hacerlo útil”. Las RNA definen su comportamiento en base a ejemplos tomados de las situaciones que se desean resuelvan. Simulando a la computación “biológica” las RNA se caracterizan por ser masivamente paralelas, altamente interconectadas, tolerantes al ruido y adaptables al medio. Hay una gran variedad de modelos de RNA, pero en general todos cuentan con los siguientes componentes:

- Un conjunto de elementos de procesamiento, referidos como “neuronas.”

- Una regla de activación para cada elemento.
- Una topología de interacción entre los elementos de procesamiento.
- Un conjunto de reglas de propagación a través de las conexiones entre elementos de procesamiento.
- Una regla de aprendizaje o algoritmo de entrenamiento de los parámetros que definen al modelo.
- El medio ambiente en el que el sistema opera.

El modelo básico de la i -ésima neurona artificial de una RNA se puede definir cómo:

$$y_i = F\left(\sum_{j=1}^n x_j w_{ij} + b_i\right) \quad (2.10)$$

donde: y_i es la salida (sinapsis) de la neurona i ; n es el número de neuronas que se conectan con la neurona i ; x_j representa la salida de la j -ésima neurona conectando hacia la neurona i ; w_{ij} representa al peso (dentrita) que conecta hacia la neurona i desde la neurona j ; b_j representa un sesgo. Tanto w_{ij} como b_i , son expresados como números reales que representan la fuerza de conexión entre neuronas y son estimados a través de un algoritmo de entrenamiento. F es la regla de activación, que puede ser lineal o no lineal, dependiendo de la arquitectura de la red. Las funciones de activación mas comunes son la escalón, la hiperbólica, la exponencial, la línea y la “unidad lineal rectificadora.”

Existen una gran variedad de arquitecturas de RNA; generalmente se agrupan en 3 grandes tipos (Haykin, 2009):

- Redes de un nivel. Las neuronas reciben entradas directamente del exterior y su salida es el resultado de la RNA. Ejemplos de este tipo de re-

des son los perceptrones de un nivel (Haykin, 2009) y la red SOM (*Self-Organizing Map*) (Kohonen, 1990).

- Redes de varios niveles o alimentadas hacia adelante. Contienen varias capas o “niveles” de neuronas. Las salidas de las neuronas en el nivel k se conectan a las neuronas en el nivel $k + 1$. La salida de la red corresponde a los valores de las neuronas en el nivel de salida. Ejemplos de este tipo de redes incluyen a los “perceptrones de varios niveles” y las redes de convolución (LeCun, Bengio, y Hinton, 2015).
- Redes recurrentes. Contienen elementos donde la salida de una neurona se usa directa o indirectamente como entrada a sí misma. El valor de salida de una neurona se determina por los valores de entrada y salidas de pasos de ejecución anteriores en el tiempo, por lo que son sistemas dinámicos. Las arquitecturas mas comunes de este tipo son las redes de Hopfield (Hopfield, 1982), la red LSTM (*Long Short-Term memory*) (Hochreiter y Schmidhuber, 1997) y la red ART (*Adaptive Resonance Theory*) (Carpenter y Grossberg, 1988).

Existen dos grandes tipos de algoritmos con los que se pueden entrenar las RNA: supervisados y no supervisados, los cuales fueron descritos en las secciones 2.2.1 y 2.2.2. El algoritmo de entrenamiento supervisado más conocido es el de retro-propagación (Werbos, 1974), comúnmente utilizado para entrenar las redes de varios niveles y las redes de aprendizaje profundo. El algoritmo de entrenamiento de la red SOM es un ejemplo de algoritmo no supervisado.

Para conocer mas sobre este vasto tema, se recomienda el curso en línea del sistema Coursera llamado “Aprendizaje de Máquina para Redes Neuronales” (<https://www.class-central.com/course/coursera-neural-networks-for-machine->

learning-398) impartido por Geoffrey Hinton, uno de los investigadores más reconocidos en esta área. Por otra parte, la librería Tensor Flow y el lenguaje científico Matlab® ofrecen excelentes herramientas para diseñar clasificadores basados en RNA.

2.4.1. Redes basadas en aprendizaje profundo

En el contexto de IA, “aprendizaje profundo” (*Deep Learning*, DL) se refiere a la actividad automática de adquisición de conocimiento, a través del uso de RNA u otro sistema de aprendizaje automático, usando varios niveles de abstracción. El adjetivo “profundo” indica la manera en que las redes aprenden, mas no necesariamente el hecho de que aprendan con gran detalle. La gran ventaja de DL es que no requiere de una definición *a-priori* de los vectores de características, sino que automáticamente se generan estas representaciones. Entre las arquitecturas mas populares basadas en DL están la red profunda de convolución (*Convolutional Neural Net*, CNN)(LeCun y cols., 2015) , las redes de creencias profundas (*Deep-Belief Network*, DBN) (Lee, Grosse, Ranganath, y Ng, 2009) y las redes recurrentes de memoria corta y larga (*Long-short term memory*, LSTM) (Hochreiter y Schmidhuber, 1997). Estas redes son altamente populares actualmente y han encontrado su máximo éxito en aplicaciones de reconocimiento y etiquetado de imágenes, reconocimiento de voz y escritura manuscrita, así como en detección de fraudes, descubrimiento de componentes farmacéuticos y procesamiento de lenguaje natural. En (LeCun y cols., 2015) se presenta una amena explicación y ejemplos de estas aplicaciones. Asimismo, una excelente introducción a RNA usando DL puede encontrarse en (Schmidhuber, 2016)

2.5. Máquinas de Soporte Vectorial

Las máquinas de soporte vectorial (*Support Vector Machine, SVM*) surgen en 1992 a partir de los trabajos de Boser, Guyon, y Vapnik (Boser, Guyon, y Vapnik, 1992). SVM pueden ser consideradas como sistemas de mapeo y clasificación, dentro de un espacio de hipótesis en funciones lineales, sobre características de alta dimensionalidad, entrenadas con algoritmos de aprendizaje basados en teoría de optimización y aprendizaje estadístico. La representación inicial de los datos es llevada a un espacio de mayor dimensión a través de núcleos de transformación o kernels, entre los cuales podemos encontrar funciones polinomiales, cuadráticas, funciones de base radial, exponencial o gaussianas.

El problema de aprendizaje supervisado se puede expresar como un problema de optimización, en donde se busca la minimización de una función de error. Para tal efecto, se inicia con un conjunto de datos de entrenamiento etiquetados con la clase a la que pertenecen. En clasificadores lineales se puede visualizar el proceso de separación a través de la búsqueda de hiperplanos que lleven a cabo tal función. Para un cierto conjunto de datos, existen muchos hiperplanos que pueden realizar la clasificación requerida; sin embargo SVM se orienta a obtener el hiperplano con el máximo margen de separación, a través de la identificación de puntos cercanos a la frontera de decisión, denominados vectores de soporte.

Estos conceptos se pueden formalizar matemáticamente de la siguiente manera: supongamos que tenemos un conjunto de P muestras de entrenamiento x_i , etiquetadas con dos posibles clases y_i con posibles valores $y_i = \pm 1$, y que cada muestra está formada por un vector de características de longitud Q , expresado de la siguiente forma:

$$\{x_i, y_i\} \quad i = 1 \cdots P \quad y_i \in \{-1, 1\}, \quad x \in \mathbb{R}^Q \quad (2.11)$$

El hiperplano en cuestión podría entonces ser expresado como $\mathbf{w}\mathbf{x} + b = 0$, donde $\mathbf{w}$ es el vector normal al hiperplano, $b/\|\mathbf{w}\|$ es la distancia perpendicular del hiperplano al origen y el conjunto de muestras de entrada al sistema se denota por la familia de vectores $\mathbf{x}_i$. La minimización del margen de separación se puede expresar matemáticamente como:

$$\text{Margen} = \text{ArgMin} \left(\frac{|\mathbf{x} \cdot \mathbf{w} + b|}{\sqrt{\sum_{i=1}^P w_i^2}} \right) \quad (2.12)$$

Una vez definido dicho hiperplano, el problema de clasificación consiste en la evaluación del conjunto de muestras de entrada $\mathbf{x}_i$ en función de la clase y_i a la que pertenezcan, de acuerdo con las siguientes relaciones:

$$x_i w + b \geq +1 \text{ para } y_i = +1 \quad (2.13)$$

$$x_i w + b \leq -1 \text{ para } y_i = -1 \quad (2.14)$$

Los “vectores de soporte” corresponden a puntos que son parte del conjunto de muestras por clasificar y que se encuentran ubicados cerca del hiperplano de separación; uno por clase en el caso más simple. Tenemos entonces dos planos H_1 y H_2 definidos por las siguientes expresiones, cuya ubicación define un margen de separación. El objetivo es entonces ubicar el hiperplano de separación en forma tal que dicho margen se maximice, es decir, que quede ubicado tan lejos como sea posible de los vectores de soporte.

$$x_i w + b = +1 \text{ para } H_1 \quad y \quad x_i w + b = -1 \text{ para } H_2 \quad (2.15)$$

La maximización de dicho margen es equivalente a minimizar el valor absoluto del vector w bajo la condición de clasificación expresada como:

$$\min \frac{1}{2} \|w\|^2 \quad y_i(x_i w + b) - 1 \geq 0 \quad \forall x_i \quad (2.16)$$

En términos de multiplicadores de Lagrange α :

$$L = \frac{1}{2} \|w\|^2 - \alpha[y_i(x_i w + b) - 1] \quad \forall x_i \quad (2.17)$$

El objetivo entonces del algoritmo de SVM radica en encontrar los valores de w y b que minimicen esta expresión y el valor de α que la maximice. Los detalles de la solución matemática se pueden consultar en diversas referencias, por ejemplo (Haykin, 2009).

Los problemas de clasificación del mundo real presentan características cuyas condiciones rebasan la simplificación bidimensional descrita. Se pueden tener casos en los cuales el número de características es elevado, o bien que se requiera una clasificación en más de dos categorías, o que las muestras que se pretenden clasificar no sean linealmente separables en el espacio de características original. En estas situaciones es factible llevar el conjunto de muestras a un espacio de mayor dimensionalidad, en donde el problema pudiera ser reducido al caso linealmente separable. Esto se puede conseguir usando núcleos de transformación o kernels. La figura 2.7(a) muestra un caso hipotético en donde se hace evidente que los datos representados en un espacio bidimensional no pueden ser separados en forma lineal, mientras que en la (b), los datos han sido llevados a un espacio de tres dimensiones a través de un mapeo con algún núcleo de transformación, la clasificación es factible usando un plano de separación.

Entre las principales funciones que pueden ser utilizadas para el mapeo se encuentran: funciones polinomiales, función sigmoide, exponencial y gaussiana de base radial.

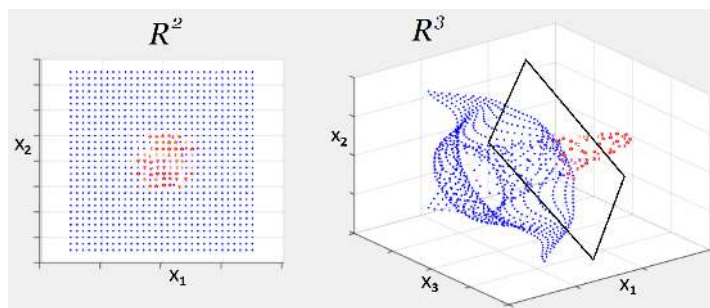


Figura 2.7: Mapeo de datos de R^2 a R^3

A la fecha existen reportados una gran cantidad de trabajos basados en máquinas de soporte vectorial (Ramírez-Cortés, Gómez-Gil, Aquino, Báez-López, y Enriquez-Caldera, 2011; J.M., P., V., J, y A., 2010). Los resultados indican que SVM representa un esquema atractivo, con muy buenas características de eficiencia tales como entrenamiento rápido y bastante robusto en las etapas de generalización.

2.6. Conclusiones

En este capítulo hemos mostrado un panorama general sobre los conceptos básicos asociados a teoría de reconocimiento de patrones y comentado brevemente sobre algunas de sus técnicas mas populares. Esta área es inmensa, tanto que pueden encontrarse especialidades, maestrías y doctorados enfocadas exclusivamente a su estudio. La importancia del reconocimiento de patrones en la actualidad es tal que resulta imposible hablar de alguna situación en la práctica que no requiera de su uso, sobre todo debido a la inseparable relación que hay entre el reconocimiento de patrones, la IA, IC y el análisis de datos.

Referencias

- Baum, L. E., Petrie, T., Soules, G., y Weiss, N. (1970, 02). A maximization technique occurring in the statistical analysis of probabilistic functions of markov chains. *Ann. Math. Statist.*, 41(1), 164–171. Descargado de <https://doi.org/10.1214/aoms/1177697196> doi: 10.1214/aoms/1177697196
- Boser, B. E., Guyon, I. M., y Vapnik, V. N. (1992, July). A training algorithm for optimal margin classifiers. En D. Haussler (Ed.), *Proceedings of the 5th annual workshop on computational learning theory (COLT'92)* (pp. 144–152). Pittsburgh, PA, USA: ACM Press. Descargado de <http://doi.acm.org/10.1145/130385.130401>
- Carpenter, G. A., y Grossberg, S. (1988, March). The art of adaptive pattern recognition by a self-organizing neural network. *Computer*, 21(3), 77–88. doi: 10.1109/2.33
- Carrasco-Ochoa, J. A., y Martínez-Trinidad, J. F. (2011). Reconocimiento de patrones. *Komputer Sapiens*.
- Duda, R. O., Hart, P. E., y Stork, D. G. (2000). *Pattern classification (2nd edition)*. New York, NY, USA: Wiley-Interscience.
- Fink, G. A. (2014). *Markov models for pattern recognition: From theory to applications* (2nd ed.). Springer Publishing Company, Incorporated.
- Fisher, R. A. (1936). The use of multiple measurements in taxonomic problems. *Annals of Eugenics*, 7(2), 179–188. Descargado de <https://onlinelibrary.wiley.com/doi/abs/10.1111/j.1469-1809.1936.tb02137.x> doi: 10.1111/j.1469-1809.1936.tb02137.x
- Haykin, S. (2009). *Neural networks and learning machines* (Third Edition ed.).

Upper Saddle River: Pearson Education Inc. (Includes Index)

- Hochreiter, S., y Schmidhuber, J. (1997, noviembre). Long short-term memory. *Neural Comput.*, 9(8), 1735–1780. Descargado de <http://dx.doi.org/10.1162/neco.1997.9.8.1735> doi: 10.1162/neco.1997.9.8.1735
- Hopfield, J. J. (1982). Neural networks and physical systems with emergent collective computational abilities. *Proceedings of the National Academy of Sciences*, 79(8), 2554–2558. Descargado de <http://www.pnas.org/content/79/8/2554> doi: 10.1073/pnas.79.8.2554
- J.M., Ramírez.-C., P., Gómez-Gil.-G., V., Aquino.-A., J, González.-B., y A., G.-P. (2010). Neural networks and SVM-based classification of leukocytes using the morphological pattern spectrum. En K. J. Melin P. y P. W. (Eds.), *Soft computing for recognition based on biometrics. Studies in computational intelligence* (Vol. 312). Berlin, Heidelberg: Springer. Descargado de https://link.springer.com/chapter/10.1007/978-3-642-15111-8_2 doi: https://doi.org/10.1007/978-3-642-15111-8_2
- Kohonen, T. (1990). The self-organizing map. *Proceedings of the IEEE*, 78(9), 1464–1480.
- Kouemou, G. L. (2011). History and theoretical basics of Hidden Markov Models. En P. Dymarski (Ed.), *Hidden markov models* (cap. 1). Rijeka: In-techOpen. Descargado de <https://doi.org/10.5772/15205> doi: 10.5772/15205
- LeCun, Y., Bengio, Y., y Hinton, G. E. (2015). Deep learning. *Nature*, 521(7553), 436–444. Descargado de <https://doi.org/10.1038/nature14539> doi: 10.1038/nature14539

- Lee, H., Grosse, R., Ranganath, R., y Ng, A. Y. (2009). Convolutional deep belief networks for scalable unsupervised learning of hierarchical representations. En *Proceedings of the 26th annual international conference on machine learning* (pp. 609–616). New York, NY, USA: ACM. Descargado de <http://doi.acm.org/10.1145/1553374.1553453> doi: 10.1145/1553374.1553453
- Manning, C. D., Raghavan, P., y Schütze, H. (2009). *An introduction to information retrieval*. Cambridge University Press. Descargado de <http://www.informationretrieval.org/>
- Mitchell, T. M. (1997). *Machine learning* (1.^a ed.). New York, NY, USA: McGraw-Hill, Inc.
- Morales, E. F., Reyes, C. A., Morales-Reyes, A., y Escalante, H. J. (2007). Aprendizaje e inteligencia computacional. En L. Pineda-Cortés (Ed.), *La computación en México por especialidades académicas* (pp. 65–90). Academia Mexicana de la Computación. Descargado de <http://amexcomp.mx/files/libro/Cap%202.pdf>
- Rabiner, L. R. (1989, Feb). A tutorial on hidden markov models and selected applications in speech recognition. *Proceedings of the IEEE*, 77(2), 257–286. doi: 10.1109/5.18626
- Ramírez-Cortés, J. M., Gómez-Gil, P., Aquino, V. A., Báez-López, D., y Enriquez-Caldera, R. A. (2011). A biometric system based on neural networks and svm using morphological feature extraction from hand-shape images. *Informatica, Lith. Acad. Sci.*, 22(2), 225–240. Descargado de <http://dblp.uni-trier.de/db/journals/informaticaLT/informaticaLT22.html#Ramirez-CortesGABE11>
- Schmidhuber, J. (2016). Deep learning. En C. Sammut y G. I. Webb (Eds.),

- Encyclopedia of machine learning and data mining* (pp. 1–11). Boston, MA: Springer US. Descargado de https://doi.org/10.1007/978-1-4899-7502-7_909-1 doi: 10.1007/978-1-4899-7502-7_909-1
- Strehl, A., Ghosh, J., y Mooney, R. (2000). Impact of similarity measures on web-page clustering. En *Workshop on artificial intelligence for web search (AAAI 2000)* (Vol. 58, p. 64).
- Sucar, L., Morales, E., y Hoey, J. (Eds.). (2011). *Decision theory models for applications in artificial intelligence: Concepts and solutions*. USA: IGI Global. Descargado de <http://www.igi-global.com/book/decision-theory-models-applications-artificial/45946>
- Theodoridis, S., y Koutroumbas, K. (2009). *Pattern recognition* (Fourth ed.). Academic Press.
- Vilain, M., Burger, J., Aberdeen, J., Connolly, D., y Hirschman, L. (1995). A model-theoretic coreference scoring scheme. En *Proceedings of the 6th Conference on Message Understanding* (pp. 45–52). Stroudsburg, PA, USA: Association for Computational Linguistics. Descargado de <https://doi.org/10.3115/1072399.1072405> doi: 10.3115/1072399.1072405
- Werbos, P. (1974). *Beyond regression : new tools for prediction and analysis in the behavioral sciences* (Tesis Doctoral no publicada). Harvard University.
- Winston, P. H. (2010). *Inteligencia artificial* (3rd ed.). Pearson.

Capítulo 3

Señales Digitales

JUAN MANUEL RAMÍREZ CORTÉS, INAOE

VICENTE ALARCÓN AQUINO, UDLAP,

RIGOBERTO FONSECA DELGADO, UNIV. YACHAI TECH

3.1. Introducción

En el inicio de un nuevo milenio, el advenimiento de impresionantes desarrollos tecnológicos que impactan positivamente nuestra vida diaria supera avasalladoramente nuestra capacidad de asombro. Un breve repaso al desarrollo que ha tenido la ciencia y la tecnología en cualquiera de sus disciplinas durante el reciente siglo nos muestra un avance exponencial que ha provocado cambios radicales en el estilo de vida sobre la faz de la tierra. Un ejemplo de tal evolución se tiene en la tecnología utilizada para el almacenamiento y distribución de contenido en forma de audio y video en la industria del entretenimiento. Haciendo un breve repaso histórico, la distribución masiva de productos musicales inicia con tecnología analógica en forma de discos de vinilo o acetato, para posteriormente dar

paso a la cinta magnética colocada en el interior de cartuchos de plástico. El salto al mundo digital ocurre con el advenimiento del disco compacto (CD), el cual utiliza tecnología óptica para leer por medio de un fino rayo laser las señales digitales previamente grabadas en la unidad. Algunos refinamientos en densidad y velocidad de lectura dan origen al video disco digital (DVD), también referido como disco versátil digital, dadas sus condiciones para almacenar video, audio o datos, en función de los requerimientos específicos. Aún cuando esta tecnología cubre exitosamente los requerimientos de capacidad, facilidad de distribución, costo y calidad a través de formatos de alta resolución como el *blue-ray*, aún se sirve para su reproducción de un lector óptico y equipo electromecánico rotativo, requerimientos tecnológicos que ante el flujo de datos en línea o *video streaming* resultan ya obsoletos. Los servicios de cómputo en la nube, o sencillamente *la nube*, conforman un paradigma de intercambio, procesamiento de datos, y servicios digitales asociados a través de la plataforma del internet, que han dado lugar al desarrollo de una infinidad de aplicaciones, muchas de las cuales han sido concebidas para su uso desde un “simple” y cotidiano teléfono inteligente. Todos estos recursos conforman un esquema tecnológico que impacta marcadamente nuestro estilo de vida en el presente y en el futuro inmediato. Tal es el caso con el llamado Internet de las Cosas (IoT), el cual congrega en ese sencillo término una avalancha de aplicaciones con base en redes inteligentes en muy diversas áreas del mundo contemporáneo tales como bioingeniería, telemedicina, seguridad, comunicaciones, sistemas de transporte, entretenimiento, o el área automotriz, entre muchas otras. De acuerdo con la empresa Cisco, Internet de las Cosas representa un mercado mundial de 35,000 millones de dólares al año 2018 en curso, con un crecimiento exponencial. En el ejemplo relacionado con aplicaciones automotrices, resulta evidente mencionar las alternativas tecnoló-

gicas que surgen cuando tanto el automóvil como el conductor son dotados de sensores de señales digitales que escalan la experiencia de la conducción y del transporte a niveles superiores. De inmediato se pueden intuir sistemas que detecten la presencia de pasajeros para señalar o accionar dispositivos de seguridad, mecanismos para climatizar automáticamente el interior, sistemas de apoyo al conductor para la navegación y detección de obstáculos o para el monitoreo de posibles estados de somnolencia, en fin, el espectro de alternativas es muy amplio. Desde luego las plantas automotrices en donde se diseñan y ensamblan todos estos sistemas no podrían permanecer al margen del mundo que se abre con la nueva tecnología, y entonces se habla de esquemas de producción en un entorno bajo la denominación de industria 4.0. Este concepto, también referido como la Revolución Industrial 4.0 es el siguiente paso a los procesos automatizados o robotizados, incorporando justamente conceptos de Internet de las Cosas, a través de servicios de cómputo ubicuo basados en sensores, señales y sistemas digitales ciberfísicos, que apoyan la conformación de nuevas formas inteligentes de organizar la producción. En palabras de Thomas Rinn, fundador de *Roland Berger Strategy Consultants* y Georg Kube, vicepresidente de *Industrial Machinery & Components Industry Business Unit at SAP*, expertos mundiales y pioneros de Industria 4.0: “La Industria 4.0 es consistente con la llamada *Cuarta Revolución Industrial*, enfatizando y acentuando la idea de una creciente y adecuada digitalización y coordinación cooperativa en todas las unidades productivas de la economía”. La palabra digital invade nuestras vidas: sitio digital, información digital, servicios digitales, procesos digitales, señales digitales, conexiones digitales, brecha digital, mundo digital, etc. Todos éstos son términos que forman parte ya, con toda naturalidad, del lenguaje cotidiano. El objetivo del presente capítulo es recapitular sobre los conceptos básicos de se-

ñales y sistemas digitales que indudablemente han dado la pauta a un mundo nuevo.

3.2. Muestreo y Cuantización

El término *señal digital* se refiere a la representación de una señal analógica que ha pasado por un proceso de discretización en magnitud conocido como cuantización y un proceso de discretización en tiempo o en espacio, según se trate de señales unidimensionales o imágenes, conocido como muestreo. Las señales analógicas tienen definido como dominio de la función que las representa a los números reales, es decir, existen en un número infinito de instantes de tiempo o espacio dentro de un cierto intervalo. En forma similar, el rango de una señal analógica también está definido como el conjunto de números reales en un determinado intervalo. Algunos ejemplos de señales analógicas que podemos encontrar en forma cotidiana son las señales de temperatura, presión, luminosidad, voz, humedad, etc. Existe una gran variedad de dispositivos transductores que permiten convertir tales señales a una representación electrónica de voltaje o corriente, lo cual permite su posterior procesamiento. Desde luego tales señales podrían ser procesadas y utilizadas directamente en su forma analógica, como sería el caso de la señal de voz de un cantante que es amplificada a través de circuitos electrónicos y convertida a una señal audible de mayor potencia, aún en forma analógica, para ser emitida a altos niveles de volumen. Esa tecnología, sin embargo, resulta obsoleta ante los sistemas actuales, digitales desde luego, que permiten llevar a cabo además de la acción referida una gran cantidad de funciones adicionales, tales como compresión de datos, almacenamiento, incorporación de efectos, etc. Existe otro conjunto de señales digitales que son amplia-

mente utilizadas en áreas tales como ciencias sociales, economía, o meteorología, referidas comúnmente como series de tiempo, las cuales admiten técnicas de procesamiento digital de señales bajo el condicionante de que presenten un período de muestreo constante. Si este no es el caso, podría ser adecuado recurrir a técnicas que permiten regularizar la serie de tiempo hasta llevarla a tal condición. Un muestreo constante permite llevar la señal discreta hacia otros dominios o representaciones a través de transformadas espectrales tales como la Transformada de Fourier o la Transformada Wavelet, entre otras. Ejemplos de series de tiempo podrían ser la temperatura ambiental medida diariamente en un observatorio meteorológico a lo largo de un año, la producción diaria en una planta de manufactura, o el costo de las acciones en la bolsa medido minuto a minuto durante un día. En estos casos, similarmente podemos interpretar la señal como un muestreo de valores que en el mejor de los casos ha sido realizado en forma uniforme o constante. Con el objetivo de iniciar el tema de las funciones que podrían ser realizadas sobre una señal en forma digital, pensemos en una señal de audio previamente digitalizada o discretizada, sobre la que ahora es factible aplicar operaciones tales como filtrados de la señal, efectos digitales, reverberación, ecualización, modulaciones, o una muy importante y de uso popular en la actualidad como es la compresión de audio en diversos formatos entre los que resalta el popular mp3. Iniciemos con la descripción cualitativa de las operaciones involucradas en filtrado digital de señales a partir de los conceptos básicos de señales y sistemas de tiempo discreto.

3.3. Señales y Sistemas de Tiempo Discreto

El tratamiento matemático de los sistemas digitales inicia subrayando el hecho de que una señal digital o discreta se puede considerar como una secuencia de muestras individuales con cierta magnitud. Formalmente una señal discreta está definida como la sumatoria de funciones delta ponderadas al valor de la señal en la posición correspondiente a cada delta, tal como se expresa en la ecuación 3.1:

$$x(n) = \sum_{m=-\infty}^{\infty} x(m)\delta(n - m) \quad (3.1)$$

La categoría básica de sistemas discretos, la cual genera todo un tema de estudio con amplísima relevancia por su ámbito de aplicación, es la de sistemas lineales. En los sistemas lineales la premisa básica es que tales sistemas mantienen una respuesta invariante bajo suma, multiplicación escalar y corrimientos en el tiempo. Tratemos de entenderlo. Si dos cantantes aplican simultáneamente su señal de voz a través de un micrófono a un sistema que amplifica mil tantos la entrada, la salida será percibida por la audiencia como mil tantos la señal del primer cantante más mil tantos la señal del segundo cantante (invariancia a la suma). Si en el mismo sistema el cantante reduce el volumen de su voz a la mitad la audiencia percibirá una señal reducida por igual a la mitad (invariancia a la multiplicación escalar). Si el cantante emite su voz con un retraso de 5 segundos la salida saldrá con un retraso igual (invarianza al corrimiento). Esta condición básica de homogeneidad de los sistemas lineales permite generalizar en forma muy interesante hacia cualquier señal discreta, es decir: dado que cualquier señal discreta se puede expresar como la suma de impulsos ponderados, es suficiente tener conocimiento sobre el efecto que tendría el sistema sobre un solo impulso unitario para tener el efecto sobre toda la señal. La función discreta

que contiene el efecto del sistema sobre un solo impulso unitario o delta se llama justamente así: respuesta al impulso, y se acostumbra utilizar la letra h para representarla. La figura 3.1 ilustra este proceso de transformación.

Unas cuantas operaciones matemáticas discretas simples conducen a la expresión que permite obtener la salida de un sistema caracterizado por su respuesta al impulso a cualquier entrada discreta. Esta operación se conoce como la *convolución* entre la señal de entrada y la respuesta al impulso. La figura 3.2 resume en forma esquemática los pasos que conducen conceptualmente a la definición del operador Convolución. Tal como se observa en la ecuación 3.1, el operador convolución consiste de sumas, multiplicaciones y corrimientos, operaciones que son fáciles de implementar a través de un sencillo programa o por medio de algún procesador digital como podría ser un microcontrolador, un FPGA (Field Programmable Gate Array), o cualquier otra unidad de microproceso. Para fines didácticos es interesante observar y entender gráficamente las operaciones involucradas en una operación de convolución.

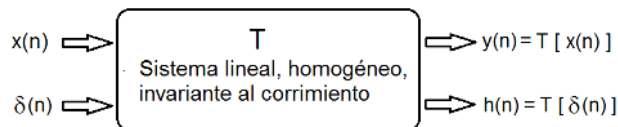


Figura 3.1: Representación del proceso de transformación de un sistema T con una señal discreta de entrada $x(n)$; respuesta al impulso.

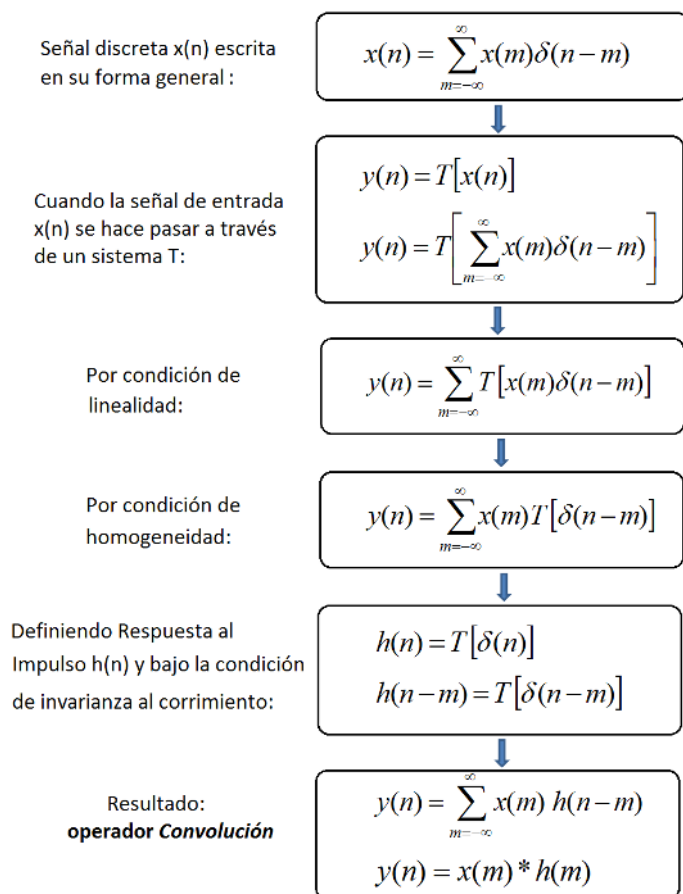


Figura 3.2: Representación del proceso matemático involucrado en el concepto de la operación “convolución”.

Algunos ejercicios simples permiten obtener cierta intuición acerca de las propiedades de la convolución, del resultado esperado a partir de una misma señal bajo diversas funciones de respuesta al impulso, y en general del concepto involucrado. Existen en la actualidad diversas páginas de internet con simuladores, animaciones, o videos en la popular plataforma *youtube* que bien pueden

servir para tal efecto. Una vez entendido el operador de convolución, la progresión natural en el estudio de señales y sistemas digitales lleva a preguntarnos sobre las posibilidades de generar e implementar diversas funciones de respuesta al impulso sobre las cuales se tenga control, a efecto de modificar las características de la señal de entrada a total conveniencia. Es decir, modificar por ejemplo su contenido armónico espectral a través de un filtrado digital. Para tal efecto, existe toda una teoría de filtros digitales bastante desarrollada, que rebasa los objetivos de este capítulo, pero que se puede encontrar fácilmente en diversos libros de texto (Proakis, 2006) (Tan y Jiang, 2018) (Semmlow, 2008). Por lo pronto ilustremos el punto de filtros digitales únicamente con un par de ejemplos: filtrado de una señal electrocardiográfica ECG y filtrado de una imagen.

Ejemplo 1: Se desea tomar la señal de ECG de la figura 3.3 y llevar a cabo un filtrado digital que permita eliminar la componente inducida de la frecuencia de línea de 60 Hz que distorsiona la señal deseada manteniendo sin variaciones al resto de componentes espectrales. Para efectos ilustrativos se restringe el tipo de filtro a los denominados FIR (Respuesta Finita al Impulso). La figura 3.4 ilustra la respuesta en frecuencia deseada. Después de utilizar alguna de las técnicas existentes en la teoría de diseño de filtros digitales, se obtiene que la respuesta al impulso que llevaría a cabo dicha tarea estaría representada por la secuencia discreta $h(n)$ de la figura 3.5.

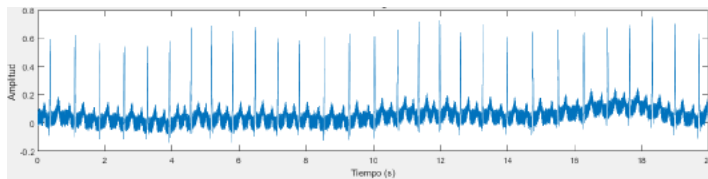


Figura 3.3: Señal de entrada $x(n)$ correspondiente a una señal EEG (de electro-encefalografía) con ruido de 60 Hz.

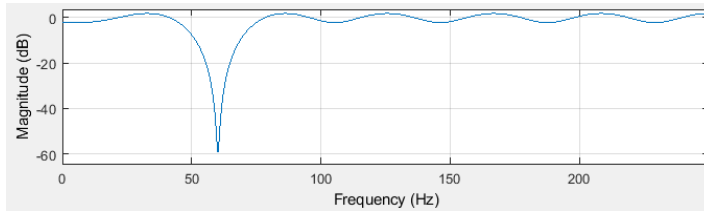


Figura 3.4: Respuesta en frecuencia del filtrado digital deseado.

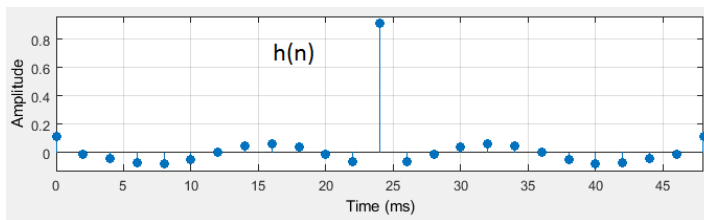


Figura 3.5: Respuesta al impulso $h(n)$ del sistema que entrega la respuesta en frecuencia deseada.

De acuerdo con los fundamentos anteriormente expuestos, la señal a la salida del sistema estará dada por la operación de convolución entre la señal de entrada $x(n)$ y la respuesta al impulso $h(n)$. La figura 3.6 muestra el resultado de dicha operación, en la cual se observa la señal EEG una vez que la componente de 60 Hz ha sido filtrada.

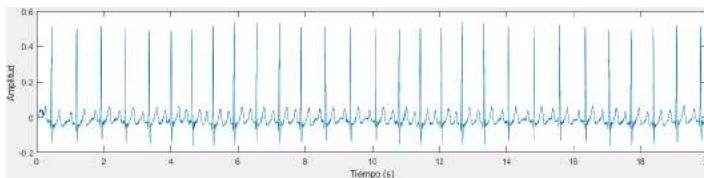


Figura 3.6: Señal $y(n)$ obtenida a la salida del filtro digital.

Ejemplo 2: Se desea aplicar el operador Laplaciano a la imagen de la figura 3.7a. La aplicación de dicho operador a una imagen en tonos de gris tiene como efecto resaltar los bordes. Cuando tal operador es aplicado a una imagen previamente binarizada, es decir, en blanco y negro, el efecto de bordes resaltados es más evidente. La teoría de procesamiento digital de imágenes – que desde luego rebasa los objetivos de este libro – indica que una de las posibilidades para la respuesta al impulso del operador mencionado está dada por una pequeña matriz de 3×3 con valores indicados en la Figura 3.7b. Esta matriz es conocida en la literatura como *Operador Sobel*. El resultado de aplicar esta pequeña matriz de 3×3 a una imagen a través de una operación de convolución en dos dimensiones se puede apreciar en la figura 3.7c, en la cual los bordes han sido resaltados. La referencia (Gonzalez y Woods, 2002) constituye una fuente clásica de tratamiento formal de esa vasta teoría contenida en procesamiento digital de imágenes.

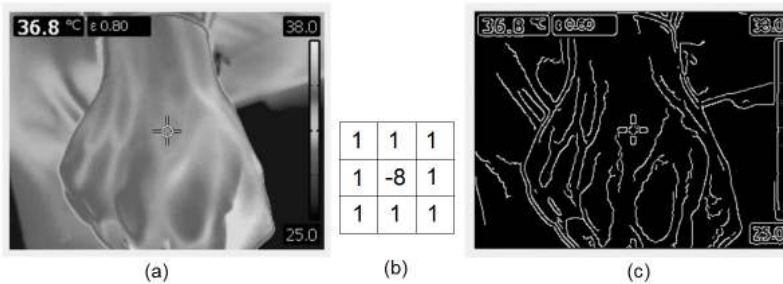


Figura 3.7: Detección de orillas (a) Imagen original en escala de grises, (b) operador Sobel, (c) Imagen obtenida al aplicar el operador en convolución 2D.

3.4. Análisis Espectral de Señales

3.4.1. Transformada de Fourier

Los desarrollos tecnológicos que han impactado la forma de vida de la humanidad durante el par de siglos recientes y que en la actualidad nos parecen tan familiares, han sido el fruto acumulado de una gran cantidad de mentes brillantes que han sentado las bases para que los conceptos teóricos abstractos se materialicen y den forma al mundo que actualmente conocemos y disfrutamos. Un tema clave dentro del área de procesamiento de señales es el Análisis Espectral, cuyos principios provienen de conceptos desarrollados por Joseph Fourier, David Hilbert y otros matemáticos brillantes a lo largo de los dos siglos recientes. El concepto fundamental estriba en la representación de una determinada función como una combinación lineal de funciones ortonormales que conforman un nuevo espacio vectorial o espacio de representación. En el caso de la Transformada de Fourier, dicho espacio está formado por el conjunto de funciones sinusoidales con frecuencias dentro de cierto rango. El análisis espectral de señales ha sido fundamental para el desarrollo de muy diversas aplicaciones y áreas tecnológicas. Tomemos como ejemplo algunos de ellos. En los sistemas de comunicación telefónica, el concepto de multicanalización de señales - transmisión simultánea de varias conversaciones telefónicas por un mismo canal - se basa en la ubicación de señales denominadas *portadoras*, cada una con diferente frecuencia. Estas frecuencias portadoras conteniendo cada una la señal de información pueden ser mezcladas y transmitidas por el mismo canal sin que las conversaciones se mezclen, y desde luego, pueden ser separadas en su momento al llegar a su destino. Esto conlleva un aprovechamiento al máximo del espacio espectral disponible. En la actualidad cualquier persona podría disponer de va-

rios números telefónicos sin preocuparse por la posibilidad de que se agoten los números disponibles. Cada una de la señales de voz, a su vez, ha sido limitada en banda a un rango de 4 Kilohertz definido y utilizado en forma global. Dicho rango tiene como fundamento el hecho de que cualquier señal temporal - y la voz no es la excepción - está constituida por la suma de una cierta cantidad de señales sinusoidales o armónicas. Un estudio básico de las señales de voz, indica que el rango de 4 Kilohertz abarca señales sinusoidales o tonos que al ser sumados entregan una señal inteligible para poder sostener una conversación. Este rango es denominado *ancho de banda del canal telefónico*. Desde luego, dicho rango no sería suficiente para transmitir un concierto musical que contendría tonos o sinusoides de mayor frecuencia, necesarios para una apreciación musical completa de la pieza. La Transformada Discreta de Fourier TDF está dada por la ecuación 3.2, y representa fundamentalmente una proyección de la señal que se pretende descomponer sobre un espacio ortonormal de funciones dadas por las señales sinusoidales en el rango bajo consideración, delimitado por la frecuencia de muestreo. A su vez, la Transformada Inversa Discreta de Fourier, dada por la ecuación 3.3 representa la reconstrucción de la señal original en forma de una combinación lineal de las mismas funciones sinusoidales. Los coeficientes complejos de Fourier estarían dando en este segundo caso, las magnitudes y fases con las cuales se modifican las funciones base para entregar una reconstrucción exacta. Existe un algoritmo muy popular que implementa la TDF de manera muy eficiente, denominado Transformada Rápida de Fourier, o FFT por sus siglas en idioma inglés. Dicho algoritmo decreta radicalmente el tiempo de ejecución, al punto de hacerlo indispensable en la mayoría de aplicaciones basadas en análisis espectral. La autoría de la FFT es tradicionalmente asignada a los científicos Tookey and Cooley, sin embargo dicho algoritmo tiene asociado

una historia muy interesante con respecto a su origen, producto de los mecanismos lentos de difusión del conocimiento característicos del siglo XX (Ramírez-Cortés, Gómez-Gil, y Baez-López, 1998). La figura 3.8 muestra una señal típica de audio y su espectro en frecuencia obtenido a través de la FFT.

$$F(u) = \sum_{n=0}^{N-1} f(n) e^{-j \frac{2\pi}{N} un} ; \text{para } u = 0, 1, \dots, N-1 \quad (3.2)$$

$$f(n) = \frac{1}{N} \sum_{u=0}^{N-1} F(u) e^{j \frac{2\pi}{N} un} ; \text{para } n = 0, 1, \dots, N-1 \quad (3.3)$$

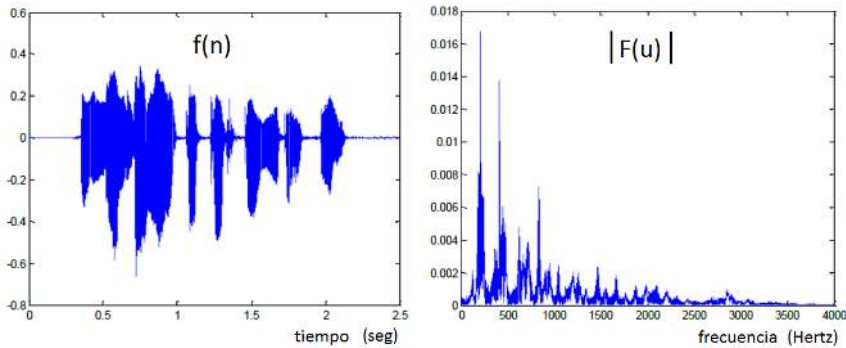


Figura 3.8: Señal de audio $f(n)$ y su espectro en frecuencia de magnitud $F(u)$.

3.4.2. Estimación Espectral

El concepto de espectro de una señal como una descomposición de ésta en funciones sinusoidales es entendible y de directa aplicación cuando se trata de señales deterministas o señales limitadas a un segmento específico, sin embargo resulta muy común y de amplia utilidad la situación de tener que *estimar* la densidad espectral de potencia o energía de una señal aleatoria estacionaria. En este

caso, dada la naturaleza estocástica de la señal y del proceso asociado, no es posible obtener con total precisión tal densidad espectral. En consecuencia, se debe hacer uso de diversas técnicas de Estimación Espectral existentes en la actualidad. Dichas técnicas son tradicionalmente clasificadas en dos grandes grupos: paramétricas y no-paramétricas. Las técnicas no-paramétricas están basadas en el uso de la TDF aplicada a través de ventanas sucesivas, usualmente traslapadas, y el promediado de los resultados con determinada estrategia, tales como el periodograma o los Metodos de Welch y Bartlett, entre otros. Un periodograma es una forma bastante eficiente, rápida, económica y con un buen nivel de resultados para llevar a cabo estimación espectral de una señal. Las técnicas paramétricas están basadas en la propuesta de cierto modelo que representa el proceso estocástico que está generando las señales por analizar, y en llevar a cabo una estimación de los parámetros que caracterizan a dicho modelo a través de diversas técnicas. Entre los principales modelos dentro de esta categoría se pueden mencionar los autoregresivos (AR), promediado móvil (MA por las siglas en inglés de Moving Average), ARMA (AutoRegresivos Moving Average). Los parámetros de tales modelos son estimados en función de las características espectrales de las señales involucradas a través de métodos como predicción lineal, mínimos cuadrados, o los métodos de Yule-Walker basados en la estimación de la secuencia de correlación o de covarianza (Semmlow, 2008).

Las figuras 3.9 a 3.12 presentan un ejemplo teórico en donde se desea llevar a cabo la estimación espectral de una señal sinusoidal inmersa en ruido aleatorio, la cual debería ser analizada como una señal estocástica, aun cuando podríamos suponer con anticipación la forma del resultado esperado. Dicha estimación es fuertemente dependiente de las condiciones establecidas por el analista, en el sentido de que en una técnica no-paramétrica se deben establecer el número,

tipo y longitud de ventanas, y similarmente, en un modelo paramétrico se debe definir con antelación el orden del sistema.

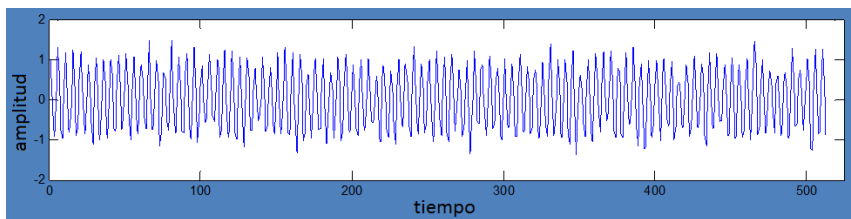


Figura 3.9: Señal sinusoidal inmersa en ruido aleatorio.

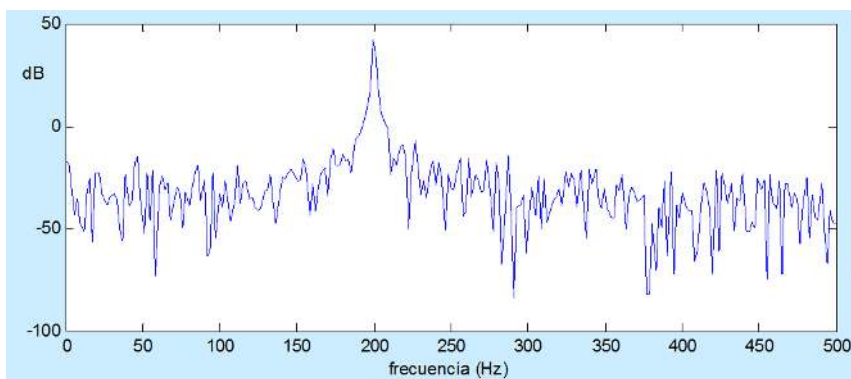


Figura 3.10: Transformada de Fourier de la señal de la figura 9; espectro de magnitud.

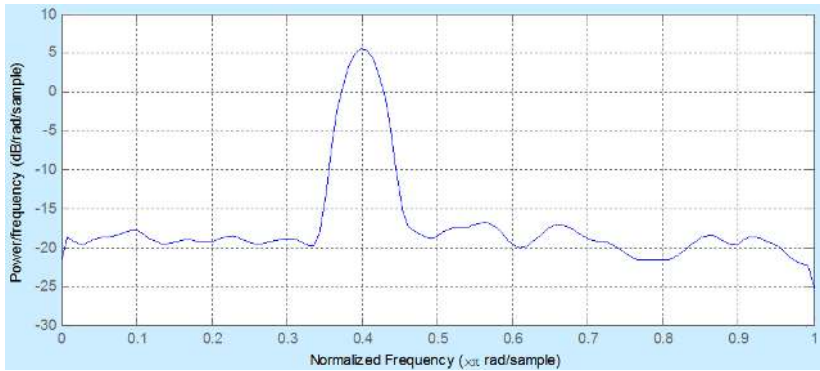


Figura 3.11: Estimación de la Densidad Espectral de Potencia a través del método de Welch de 8º. orden.

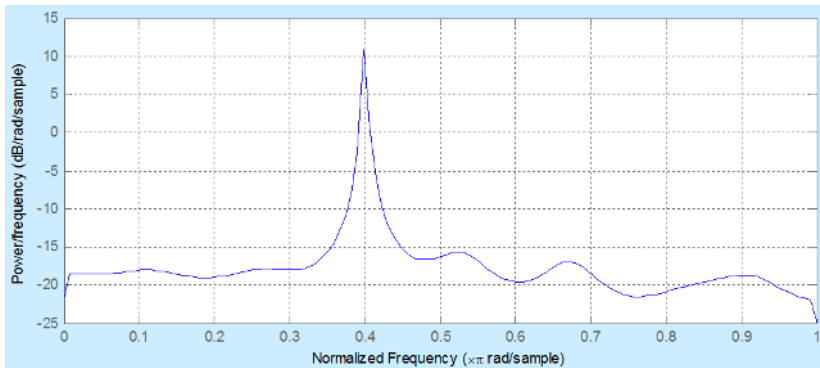


Figura 3.12: Estimación de la Densidad Espectral de Potencia a través del método de Yule Walker de 8º. Orden.

Un ejemplo de aplicación en el mundo real en donde sería requerido realizar estimación espectral con determinado objetivo científico o técnico podría ser el análisis de señales cerebrales de electroencefalografía (EEG). Las señales EEG son bioseñales que registran la actividad eléctrica cerebral obtenida por medio de electrodos superficiales colocados a nivel del cuero cabelludo o corte-

za superficial del cerebro. Estas señales han sido estudiadas a lo largo del siglo XX cada vez en mayor medida, incrementándose la intensidad de tales estudios en diversidad, calidad y cantidad, a la par del desarrollo tecnológico que ha ido poniendo a disposición de la comunidad científica equipos de EEG cada vez con mayores funcionalidades y más accesibles. Las señales de EEG han sido estudiadas con muy diversos fines, entre los cuales se pueden mencionar el estudio del cerebro asociado con la funcionalidad de los organismos, y así ha sido posible hacer mapeos de las zonas cerebrales relacionadas con la visión, con los mecanismos del habla, con el control motor, con las tareas mentales, etc. De gran interés entre los investigadores ha sido también el estudio del comportamiento cerebral como resultado de diversos estímulos tales como luminosos, visuales, auditivos, somestésicos o nociceptivos, entre otros, así como el contraste entre el comportamiento de las diversas bandas espectrales entre estados de vigilia y sueño. Sin lugar a dudas, los médicos especialistas han encontrado en las señales de EEG un excelente recurso para el estudio, diagnóstico y tratamiento de diversas patologías cerebrales, entre las cuales se pueden mencionar encefalopatías, tumores cerebrales, enfermedades degenerativas del sistema nervioso central, traumatismos craneoencefálicos, epilepsia o trastornos psiquiátricos de diversa índole. Ejemplos de esto son estudios realizados recientemente en diversas instituciones, relacionados con el diagnóstico y seguimiento de epilepsia (Juarez-Guerra, Alarcon-Aquino, y Gomez-Gil, 2015) (Gómez-Gil, Juárez-Guerra, Alarcón-Aquino, Ramírez-Cortés, y Rangel-Magdaleno, 2014). La ingeniería ha encontrado también su nicho de oportunidad en el estudio de las señales EEG para el desarrollo de tecnología, como es el caso con las llamadas Interfaces Cerebro-Computadora o BCI por sus siglas en inglés (Morales-Flores, Ramírez-Cortés, Gómez-Gil, y Alarcón-Aquino, 2013a) (Morales-Flores, Ramírez-

Cortés, Gómez-Gil, y Alarcón-Aquino, 2013b). En un sistema BCI el objetivo es obtener ventaja de las señales EEG generadas por un individuo para llevar a cabo tareas de índole práctico. Esta actividad adquiere relevancia entre personas que en condiciones de inmovilidad severas, que bien pueden ver mejorada hasta cierto grado su calidad de vida al encontrar nuevas formas de comunicación y control hacia el mundo externo. Ejemplos de aplicaciones de BCI podrían ser el control de una silla de ruedas, la manipulación de dispositivos de ambientación como persianas o climas artificiales en edificios inteligentes, poder navegar en internet o sencillamente tener el control de una unidad de televisión a partir de señales cerebrales.

3.5. Análisis Tiempo-Frecuencia

3.5.1. Espectrograma

La Transformada Discreta de Fourier es ampliamente utilizada en procesamiento digital de señales estacionarias con excelentes resultados, sin embargo, cuando se trata de señales no-estacionarias, es decir, aquellas cuyas características estadísticas varían en función del tiempo, es necesario utilizar técnicas de análisis tiempo-frecuencia. La solución inmediata consiste en aplicar la DFT en ventanas de longitud uniforme. A esta operación se le conoce como Transformada de Fourier de Tiempo Corto, o STFT por sus siglas en inglés. El resultado de la operación es una representación bidimensional conteniendo una aproximación de la distribución de la energía de la señal en los ejes de tiempo y frecuencia. A la gráfica de la magnitud de dicha información, representada típicamente en pseudocolor, se le conoce como espectrograma. Esta representación es extremadamente útil en muy diversas aplicaciones, como por ejemplo análisis de señales de voz o

en general señales audibles. La figura 3.13 muestra un ejemplo del espectrograma correspondiente al tono que se escucha durante la secuencia de marcado de un número telefónico. Es bien sabido que el teclado de un teléfono se encuentra codificado en forma de matriz, en donde cada número corresponde con un par de tonos sinusoidales asociados con el renglón y la columna correspondiente. Esa característica se aprecia claramente cuando se analiza el espectrograma, mientras que la gráfica de la señal en el tiempo solo proporciona información acerca de los momentos en que ocurren dichos tonos, pero no de su contenido armónico espectral.

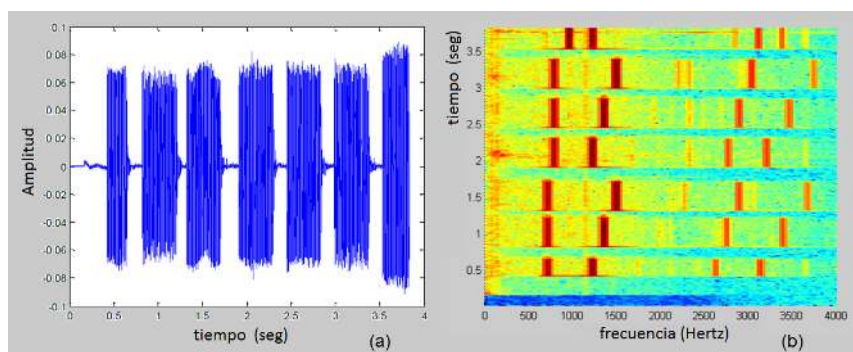


Figura 3.13: Análisis tiempo-frecuencia (a) Señal de audio obtenida con la marcación telefónica de los siete primeros dígitos (b) Espectrograma de la misma señal.

Si bien el espectrograma es altamente utilizado en el análisis de señales, su principal inconveniente radica en el hecho de ser obtenido a través de ventanas de longitud uniforme, lo cual condiciona la resolución espectral a ser uniforme también. Se presenta entonces una situación de compromiso del análisis en el cual una buena resolución temporal provoca una baja resolución espectral y viceversa. Este compromiso restringe la eficiencia del análisis especialmente en

señales no-estacionarias, y es en esos casos en los cuales resalta de manera natural la necesidad de abordar esquemas de análisis multiresolución. Una de dichas técnicas es la Transformada de Ondeletas, la cual es descrita a continuación.

3.5.2. Transformada de Ondeletas

La descomposición espectral clásica contenida en la Transformada de Fourier de Tiempo Corto podría ser considerada una Transformación básica de Ondeletas, viendo a la función faso como la señal básica de representación. Lo interesante de la Transformada de Ondeletas es que admite la definición inicial de una función denominada Ondeleta Madre, que es utilizada a diversas escalas y translaciones sobre la señal que se desea analizar. De esta manera, la función ondeleta madre da origen a una función de funciones ortonormales que constituyen la base de un nuevo espacio de representación, en forma similar a la STFT, y entregando información sobre la distribución de la energía de la señal en un espacio tiempo-frecuencia. La función ondeleta madre debe reunir algunas condiciones para calificar como tal. Debe ser una función oscilatoria que permita la identificación de cambios a diversas escalas, y debe tener un soporte compacto con energía finita. A lo largo de las décadas recientes, diversos científicos han propuesto una cantidad de funciones que pueden ser utilizadas como función ondeleta madre, pudiéndose utilizar entre las mas utilizadas las ondeletas tipo Haar, Morlett, Meyer, Mexican Hat, Daubechies y otras, con diversos órdenes. Si bien no existe una forma precisa de seleccionar alguna de ellas en determinada aplicación, se ha observado que los mejores resultados se obtienen cuando el tipo de ondeleta seleccionada se contiene características en términos de oscilaciones, variaciones y contenido espectral del estilo en cierta forma la señal por analizar. El caso particular de las señales EEG, en donde la wavelet tipo Dau-

bechies ha encontrado un uso generalizado, es un ejemplo de esto. Existen dos esquemas para implementar la Transformada Ondeleta, denominadas Transformada continua y discreta. En este trabajo únicamente nos referiremos a la versión discreta, representada por las iniciales DWT por sus siglas en idioma inglés. El algoritmo de la DWT también llamado codificación sub-banda o algoritmo de Mallat, consiste fundamentalmente en el filtrado digital en dos etapas pasa-bajas y pasa-altas frecuencias, respectivamente, seguidas por una decimación a la mitad de las muestras, tal como se observa en la Figura 3.14. Las salidas de estas dos etapas se denominan coeficientes de aproximación y detalle, respectivamente. Con esto, se enfatiza el hecho de que la DWT analiza la señal entregando simultáneamente los aspectos generales y las características finas de una señal, o como lo expresan en algunos textos, visualizando simultáneamente “*el bosque*” y “*las hojas*”. Esta operación se repite en tantos niveles como el diseñador establezca. Vale la pena resaltar el hecho de que independientemente del número de etapas utilizadas en una determinada descomposición por ondeleta, siempre es factible obtener el rango de frecuencias asociadas con cada uno de los conjuntos de coeficientes (Semmlow, 2008) (Van Fleet, 2011).

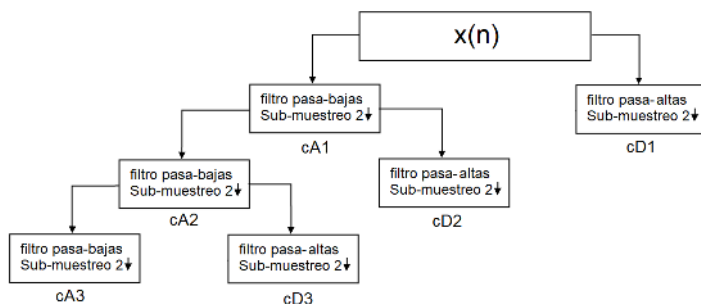


Figura 3.14: Algoritmo de codificación sub-banda utilizado en la Transformada Discreta de Ondeletas.

A partir de dichos conjuntos, es también factible obtener la representación tiempo-frecuencia, conocida como *Escalograma*, con la distribución de energía de la señal en dicho espacio. Desde luego, al igual que la gran mayoría de las transformadas espectrales, la DWT es reversible en el sentido de que existe la operación denominada Transformada Discreta Wavelet Inversa o IDWT, a través de la cual es posible recuperar íntegramente la señal original. Esta característica da lugar a diversas operaciones intermedias que pueden ser llevadas a cabo en función del tipo de análisis requerido. Por ejemplo, si es que se descartan algunos de los grupos de coeficientes en determinado nivel se podría estar haciendo una especie de filtrado digital, que bien podría ser utilizado para supresión de algún tipo de ruido. Si se descartan aquellos coeficientes que aportan muy poca información o que contienen muy poca energía, típicamente algunos coeficientes de detalle, se podría estar representando la señal con un número reducido de valores, lo cual es la operación esencial en los algoritmos de compresión de datos.

Se mencionan a continuación algunos ejemplos de proyectos de investigación recientemente desarrollados en México basados en la aplicación de DWT.

3.6. Algunos proyectos de investigación

En esta sección se describen algunos proyectos de investigación realizados en México, que tienen como fundamento la aplicación de técnicas de procesamiento digital de señales o imágenes en los dominios de tiempo o espacio, así como a través de análisis espectral. Uno de los paradigmas más utilizados en el desarrollo de Sistemas BCI es el denominado imaginación motora, el cual consiste en detectar las señales EEG generadas en determinada zona del cerebro, típicamen-

te los electrodos C₃ y C₄ del Sistema Universal de Localización de EEG 10-20, como resultado de la intención del usuario para realizar una acción de movimiento. La referencia (Carrera-León y cols., 2012) presenta un estudio en el cual las señales de imaginación motora son analizadas a través de modelos autoregresivos con el objeto de llevar a cabo la extracción de características requerida en un sistema de reconocimiento o clasificación de movimiento. Específicamente, los vectores de características se forman con los coeficientes de un modelo autoregresivo de 6º orden, el cual contiene la estimación espectral de las señales EEG bajo estudio, y es utilizado para llevar a cabo el reconocimiento de la intención de movimiento de las extremidades superiores izquierda y derecha de la persona. En dicho estudio se utilizaron dos métodos para la estimación de los parámetros del modelo autoregresivo: Método de Burg y Método de Levinson-Durbin. La figura 3.15 muestra el ejemplo de una de las señales de EEG analizada por ventanas de un segundo con un 50 % de traslape.

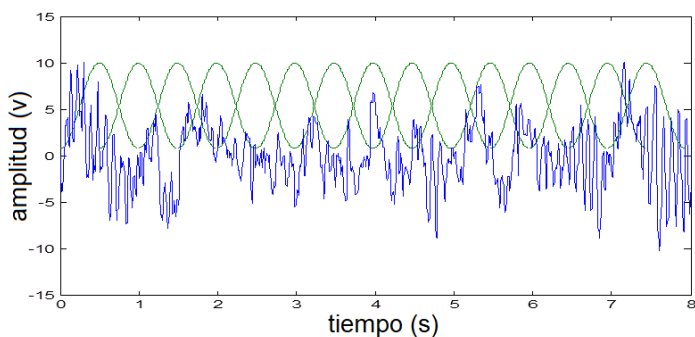


Figura 3.15: Señal EEG analizada bajo el paradigma de imaginación motora.

La Figura 3.16 muestra el desarrollo de un sistema BCI para la emulación de las operaciones de un *mouse*, diseñado para permitirle a un usuario la navegación en internet. En dicha investigación se utilizaron las señales de un giroscopio para

localizar la posición relativa del *mouse* en la pantalla, mientras que el equivalente a las acciones de selección o *clicks* se realizaron a través de las señales de electro-oculografía presentes en las señales de EEG como resultado de guiños en uno o ambos ojos. Las señales de guiño o parpadeo fueron localizadas a través de un modelo adaptivo neuro-difuso ANFIS, entrenado específicamente para tal acción (Rosas-Cholula y cols., 2013), entregando una razón de detección de un 90 % en promedio.

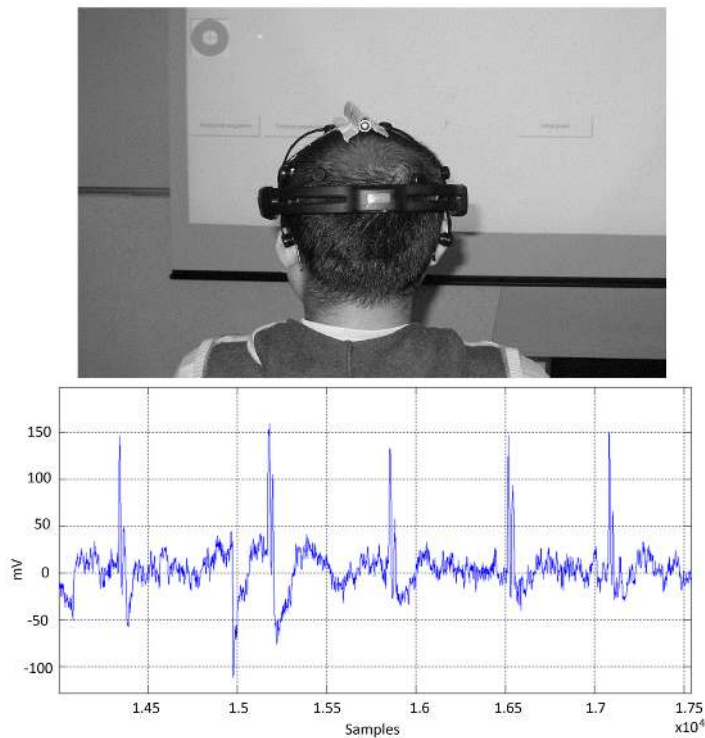


Figura 3.16: Sistema BCI para la emulación de las operaciones de un *mouse*.

En las referencias (Juarez-Guerra y cols., 2015) (Gómez-Gil y cols., 2014) se presentan estudios orientados a la detección de epilepsia sobre el análisis de seña-

les EEG utilizando la DWT y una variante de ésta conocida como DWT de máximo traslape (MODWT). Estos trabajos reportan resultados con un esquema de descomposición de 4 niveles utilizando la ondeleta madre tipo Daubechies de segundo y cuarto orden, con porcentajes de clasificación superiores al 98 %. La figura 3.17 a,b muestra un ejemplo del tipo de señales EEG en condiciones de normalidad y convulsión epiléptica, respectivamente.

Algunos otros ejemplos de proyectos de investigación desarrollados con base en técnicas de análisis tiempo-frecuencia utilizando la Transformada Discreta de Ondeletas son: análisis de mamogramas digitales para la detección temprana de microcalcificaciones asociados con patologías de cáncer (Alarcon-Aquino y cols., 2009), sistemas BCI bajo el paradigma de imaginación motora (Hernández-González y cols., 2017), clasificaciones de señales de electromiografía con potencial aplicación en el control de extremidades robóticas (Cifuentes-González, Ramírez-Cortés, Alejos-Palomares, y Conforto, 2012), aplicaciones industriales para la detección y monitoreo de fallas en motores (López-Hernandez, Rangel-Magdaleno, Peregrina-Barreto, y Ramírez-Cortés, 2018), o técnicas para compresión de imágenes y codificación de señales de video (J. C. Galan-Hernandez, Alarcón-Aquino, Ramírez-Cortes, y Gómez-Gil, 2018) (J. C. Galan-Hernandez, Alarcon-Aquino, y Ramirez-Cortes, 2010) (J. C. Galan-Hernandez, Alarcon-Aquino, y Ramirez-Cortes, 2013). En años recientes se han reportado diversos estudios y proyectos de investigación en el campo de la biometría con base en bioseñales. Existe en la actualidad una buena cantidad de modalidades de biometría. El estado de las tecnologías digitales en el mundo actual presenta una gran variedad de aplicaciones que requieren, cada vez en mayor medida, métodos de autenticación confiables y seguros para confirmar la identidad de un individuo al requerir un servicio.

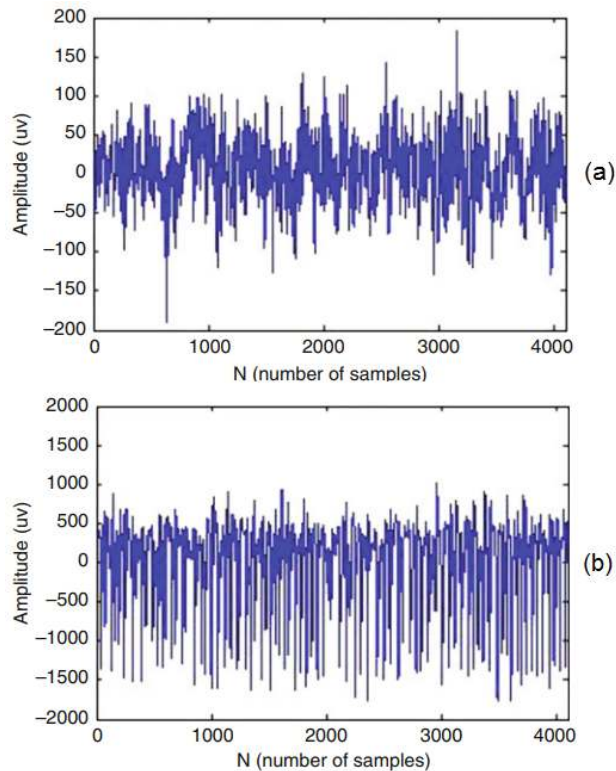


Figura 3.17: Señales cerebrales para detección de epilepsia (a) Señal EEG normal (b) Señal EEG en estado de convulsión epiléptica.

Algunos ejemplos de uso de sistemas biométricos son el acceso restringido a edificios, la entrada a centros de diversión, el uso de un teléfono celular, el acceso a sistemas de cómputo, tabletas, y muchos más. Los sistemas de seguridad basados en biometría están orientados a llevar a cabo dicha validación a través de la medición de características fisiológicas (huella dactilar, iris, mano, rostro) o comportamentales (firma, forma de andar, forma de accionar un teclado) propias de la persona. El objetivo fundamental estriba en que los sistemas de seguridad verifiquen “*quién es*” la persona que requiere el servicio, en vez de “*lo que*

trae consigo”, como implica el uso de una tarjeta de banda magnética, o “*lo que recuerda*”, como es el caso al solicitar una contraseña. Entre las más relevantes se puede mencionar el uso de la huella dactilar, rasgos del iris, forma y huella de la mano, características temporales y espectrales de la voz, reconocimiento de rostro, modo personal al caminar, cadencia en el golpeteo al teclado, rasgos estáticos y dinámicos de la firma, etc. Todas estas modalidades, sin embargo, tienen un cierto nivel de vulnerabilidad a posibles falsificaciones. Una tendencia actual es la exploración de técnicas multimodales biométricas que incorporen bioseñales del individuo tales como señal cardíaca, electroencefalográfica, presión arterial, etc., las cuales además de reunir las características requeridas para ser consideradas rasgos biométricos, resultan difíciles de falsificar. Algunos proyectos sobre biometría basada en bioseñales se encuentran actualmente en desarrollo en INAOE. Específicamente se tienen en curso los siguientes trabajos: biometría bimodal voz-EEG que incorpora técnicas de procesamiento digital de señales y análisis espectral en base a transformada de ondeletas, biometría bimodal usando señales de electrocardiografía (ECG) y fotopletismografía (PPG) con un esquema de reconocimiento basado en distorsión temporal dinámica (DTW), y biometría usando procesamiento de imágenes termográficas obtenidas de las venas de la mano. Los resultados preliminares obtenidos indican perspectivas muy interesantes para la continuación de actividades de investigación en tales temas.

3.7. Conclusiones

En este capítulo se han presentado en forma general los principales conceptos de procesamiento digital de señales en los dominios de tiempo y espacio, y de manera sucinta los conceptos básicos de análisis espectral con base en la Trans-

formada de Fourier, la cual ha abierto el panorama hacia el estudio de muy diversos espacios de representación espectral tales como la Transformada Walsh, Hilbert-Huang, o la Transformada por Ondeletas. Igualmente se ha hecho referencia a algunos cuantos proyectos de investigación actualmente en proceso. Sin lugar a dudas las técnicas de procesamiento digital de señales han sido clave en el desarrollo de una gran cantidad de aplicaciones que han surgido y seguirán apareciendo en forma de tecnologías disruptivas, las cuales modifican radicalmente nuestro estilo de vida. Seguramente el siglo XXI, cuyo inicio estamos ya viviendo plenamente, nos depara con una gran cantidad de sorpresas científicas y tecnológicas, que estarán conformando un nuevo mundo digital lleno de paradigmas que probablemente aún nos son ajenos.

Referencias

- Alarcón-Aquino, V., Starostenko, O., Ramírez-Cortés, J., Rosas-Romero, R., Rodríguez-Asomoza, J., Paz-Luna, O. J., y Vázquez-Muñoz, K. (2009). Detection of microcalcifications in digital mammograms using the dual-tree complex wavelet transform. *Engineering Intelligent Systems*, 17(1), 49.
- Carrera-Leon, O., Ramírez, J. M., Alarcón-Aquino, V., Baker, M., D'Croz-Barón, D., y Gómez-Gil, P. (2012). A motor imagery BCI experiment using wavelet analysis and spatial patterns feature extraction. En *2012 Workshop on Engineering Applications* (pp. 1–6).
- Cifuentes-González, I., Ramírez-Cortés, J. M., Alejos-Palomares, R., y Conforto, S. (2012). Lower limb EMG signal classification using discrete wavelet transform and hilbert-huang transform. En *XVII Congreso de Instrumen-*

tación (pp. 47–52).

Galán-Hernández, J., Alarcón-Aquino, S. O., Vicente, y Ramírez-Cortes, J. (2010). DWT foveation-based multiresolution compression algorithm. *Research in Computing Science*, 197–206.

Galan-Hernandez, J. C., Alarcon-Aquino, S. O., Vicente, y Ramirez-Cortes, J. (2013). Fovea window for wavelet-based compression. En *Innovations and advances in computer, information, systems sciences, and engineering* (pp. 661–672). Springer.

Galán-Hernandez, J. C., Alarcón-Aquino, S. O., Vicente, Ramírez-Cortes, J. M., y Gómez-Gil, P. (2018). Wavelet-based frame video coding algorithms using fovea and speck. *Engineering Applications of Artificial Intelligence*, 69, 127–136.

Gómez-Gil, P., Juárez-Guerra, E., Alarcón-Aquino, V., Ramírez-Cortés, M., y Rangel-Magdaleno, J. (2014). Identification of epilepsy seizures using multi-resolution analysis and artificial neural networks. En *Recent advances on hybrid approaches for designing intelligent systems* (pp. 337–351). Springer.

Gonzalez, R. C., y Woods, R. E. (2002). *Digital image processing*. Prentice hall Upper Saddle River, NJ.

Hernández-González, C. E., Ramírez-Cortés, J. M., Gómez-Gil, P., Rangel-Magdaleno, J., Peregrina-Barreto, H., y Cruz-Vega, I. (2017). EEG motor imagery signals classification using maximum overlap wavelet transform and support vector machine. En *2017 IEEE International Autumn Meeting on Power, Electronics and Computing (ROPEC)* (pp. 1–5).

Juárez-Guerra, E., Alarcón-Aquino, V., y Gómez-Gil, P. (2015). Epilepsy seizure detection in EEG signals using wavelet transforms and neural networks.

- En *New trends in networking, computing, e-learning, systems sciences, and engineering* (pp. 261–269). Springer.
- López-Hernández, M., Rangel-Magdaleno, J., Peregrina-Barreto, H., y Ramírez-Cortés, J. (2018). Detection of broken bars on induction motors using modwt. En *2018 IEEE International Instrumentation and Measurement Technology Conference (I2MTC)* (pp. 1–5).
- Morales-Flores, E., Ramírez-Cortés, J. M., Gómez-Gil, P., y Alarcón-Aquino, V. (2013a). Brain computer interface development based on recurrent neural networks and anfis systems. En *Soft computing applications in optimization, control, and recognition* (pp. 215–236). Springer.
- Morales-Flores, E., Ramírez-Cortés, J. M., Gómez-Gil, P., y Alarcón-Aquino, V. (2013b). Mental tasks temporal classification using an architecture based on anfis and recurrent neural networks. En *Recent advances on hybrid intelligent systems* (pp. 135–146). Springer.
- Proakis, J. G. (2006). *Digital signal processing: principles algorithms and applications*. Pearson Education India.
- Ramírez-Cortés, J. M., Gómez-Gil, P., y Baez-López, D. (1998). El algoritmo de la transformada rápida de Fourier y su controvertido origen. *Ciencia y Desarrollo*, 24(139), 70–76.
- Rosas-Cholula, G., Ramírez-Cortés, J., Alarcón-Aquino, V., Gómez-Gil, P., Rangel-Magdaleno, J., y Reyes-García, C. (2013). Gyroscope-driven mouse pointer with an emotiv® EEG headset and data analysis based on empirical mode decomposition. *Sensors*, 13(8), 10561–10583.
- Semmlow, J. L. (2008). *Biosignal and medical image processing*. CRC press.
- Tan, L., y Jiang, J. (2018). *Digital signal processing: fundamentals and applications*. Academic Press.

Van Fleet, P. J. (2011). *Discrete wavelet transformations: An elementary approach with applications*. John Wiley & Sons.

Capítulo 4

Visión Computacional

ELISA GÓMEZ-GÓMEZ, UNIV. A. DE AGUASCALIENTES

HIRAM H. LÓPEZ, UNIV. A. DE AGUASCALIENTES

JUAN CARLOS SÁNCHEZ UNIV. TEC. DE TECÁMAC,

HERMILO SÁNCHEZ-CRUZ, UNIV. A. DE AGUASCALIENTES

4.1. Introducción

El procesamiento digital de imágenes, su análisis, el reconocimiento de patrones y la computación gráfica (Domingo, 2004), son procesos que se trabajan en la llamada visión por computadora y se han aplicado en diversas áreas de la ciencia, desde las industrias hasta la salud humana. A continuación vamos a describir los inicios de la visión por computadora, lo que es el procesamiento de imágenes y se mencionan algunos trabajos desarrollados que han hecho del uso de la visión por computadora una herramienta clave en el descubrimiento, solución y propuesta de innovaciones para el mejoramiento de la calidad de vida del ser humano. Además, se mencionan aplicaciones en la industria de la construcción,

el tráfico vehicular y con mayor énfasis en la biología. Finalmente, se describen algunos de los trabajos que han sido desarrollados en México y que han sido publicados en revistas de prestigio internacional.

4.2. Antecedentes de la visión por computadora

Desde los años sesenta se trabajaba en el procesamiento de imágenes. En aquel entonces, se usaban algoritmos para la aplicación de teorías enfocadas a configurar una representación gráfica en una computadora, donde los puntos de una imagen fueron los elementos clave a estudiar. Se soñaba y pensaba que un día lograrían que la computadora pudiese “ver”. Ochenta años después no se ha logrado. En la siguiente década ya se trabajaba con análisis de imágenes donde se combinó la medición de las imágenes con el reconocimiento estadístico de patrones. En los ochenta, gracias a los modelos desarrollados en el análisis de imágenes y la combinación de inteligencia artificial con el modelado geométrico, dieron paso a la visión por computadora. Aunque tiene sus bases en las matemáticas y las ciencias de la computación, hay una vinculación más flexible con la física, la psicología de la percepción y de las neurociencias (Hartley y Zisserman, 2003).

Se conoce como “imagen” a algo que está grabado: ya sea un video, un dibujo digital o una fotografía (Haralick y Shapiro, 1991). Ya Solem escribió, en su libro, que la visión por computadora es “la extracción automática de información de las imágenes” (Solem, 2012). En notación matemática, el valor de una imagen en las coordenadas (r, c) se denota como $l(r, c)$. Por ejemplo, cuando la imagen es un mapa, $l(r, c)$ es un índice o símbolo asociado con algunas categorías que pueden ser tales como el color, un tipo de suelo o una característica cualquiera

asociada a la imagen en ese punto. Una imagen puede almacenarse en diversos formatos como una fotografía, un vídeo o en formato digital.

En 1973, Duda y Hart publicaron una primera edición del libro *Pattern Classification* (Duda y Hart, 1973) el cual es ahora un clásico en esta área. Ahí se tratan trabajos como las redes neuronales, el reconocimiento estadístico de patrones, la teoría del aprendizaje automático y la teoría de invariantes, pero, sobre todo, resalta el diseño de algoritmos que incluso hoy en día sirven de base y pueden ejecutarse con mayor rapidez, gracias al avance en la tecnología informática, tanto en hardware como en software. Actualmente sirven de apoyo para entrar a este mundo de la visión por computadora. Fue hasta el 2003 en que se presenta una nueva edición con ejercicios, gráficas y, en el área de la computación, se plantean proyectos de aplicaciones informáticas.

En el 2011 se usaba el procesamiento digital de imágenes, donde se daba un proceso en el que se genera una versión alterada de algún objeto. Mery (Domingo, 2004) muestra trabajos donde un elemento básico a considerar de una imagen es el pixel y de ahí parten modelos que permiten mostrar entidades optimizadas para su construcción y uso.

4.3. Procesamiento digital de imágenes

La visión por computadora es ampliamente utilizada en áreas tan posiblemente dispares o alejadas. Por ejemplo, en la industria de la construcción se obtiene información que proporcionan los escáneres de media distancia y los aplica en el desarrollo optimizado del espacio habitacional. En las ciencias de la salud, una tomografía no es otra cosa mas que la impresión de una imagen digital.

Adán y Huber realizaron un trabajo en un entorno de construcciones civiles y edificios donde se generaron de manera automática modelos BIM (Adana y Huberb, 2011) (*Building Information Models*) en espacios habitados. Aquí, se usaron sensores 3D y los obtenidos fueron manejados por aplicaciones informáticas de diseño gráfico para, finalmente, construir casas con un respaldo científico. Uno de los trabajos que propuso en el 2010 fue “identificar partes esenciales de la estructura de interiores de edificios en entornos altamente desordenados y con un alto componente de oclusión”. No es lo mismo que un decorador de interiores establezca el orden armonioso y espiritual a que lo haga un sistema computarizado donde se le provean de las variables necesarias para su posicionamiento y colocación en el espacio del edificio de cada elemento que lo conformará.

Pensar en actividades como caminar o conducir en una ciudad y su relación con un algoritmo que controle tiempos, distancias, cantidad de avenidas o carreteras secundarias, cada vez es más cotidiano. Una de las principales razones es que las grandes ciudades que sufren de un tránsito continuo y constante y que, en muchas ocasiones, el volumen rebasa la capacidad de operaciones establecidas, requieren de un modelado por computadora que les ayude a crear una visión virtual de la ciudad y un esquema eficiente de sistema de navegación. Es importante señalar que estos sistemas viales requieren de la interpretación de la información en tiempo real para agilizar la toma de decisiones eficientes. B. Mushleh propone aplicaciones avanzadas de ayuda en la conducción para la mejora de la seguridad vial (Musleh y A. De-la Escalera-Hueso, 2012). Google, o diferentes marcas de automóviles, como por ejemplo Tesla, proponen la navegación autónoma de vehículos.

No lejos de esto, los semáforos y sus tiempos implican un tráfico fluido o ralentizado con sus problemáticas de contaminación y retraso en tiempos de arribo de los conductores a sus respectivos lugares. En diferentes ciudades, como lo son Sao Paulo, Zurich, Singapur y Anda Mendoza, se propone un ordenamiento del transporte utilizando Big Data (Anda, Erath, y Fourier, 2017). En sus trabajos, Anda propone modelos de demanda de transporte para utilizarlos en la planificación vehicular de las grandes ciudades. Se han desarrollado aplicaciones informáticas para dispositivos móviles, donde se pueden mostrar la mejor vía para que un usuario se dirija a su destino o, para el establecimiento de una ruta de transporte a definir por un gobierno local. Todo a través de la aplicación de estrategias matemáticas, de ingeniería y usando la computadora como medio de pruebas y validación de posibles soluciones.

A.K. Jain revisa y resume varios métodos utilizados en los sistemas de reconocimiento de patrones e identifica tópicos de investigación y aplicaciones, que son el sustento actual de este tema. ¿‘Qué tal sería la reparación de obras de arte, donde hay huecos o partes donde la figura se ha borrado casi completamente? Con el diseño de un algoritmo (Jain, Duin, y Mao, 2000) que selecciona una sección y calcula una proyección adecuada para predecir-completar-continuar la imagen, quedan solucionados muchos recuerdos de la generación de personas que todavía conservan fotografías en papel y que ya no presentan un buen estado. Si se traslada al arte y la cultura de los pueblos del mundo, el abanico de posibilidades es mucho más amplio: pinturas del renacimiento, de arte barroco, de la época paleolítica o incluso modernas.

4.4. Aplicaciones en Bioinformática

En la biología se han desarrollado avances importantes en la detección de enfermedades que tienen un alto índice de mortalidad, como por ejemplo el cáncer. Se han desarrollado modelos computacionales que permitan conocer uno de los factores importantes en su cura o prevención: las proteínas. Para conocer su secuencia y su estructura 3D se usan diversos métodos, entre los más aplicados es la Cristalografía de Rayos X. De ahí su estructura digital obtenida se almacena en repositorios web de libre acceso. Ya desde mediados del siglo pasado se desarrollaron simulaciones de modelos de resoluciones atómicas, pero que demandaban mucho tiempo de procesador. Autores como Gō planteaban, desde 1983, trabajos en computadora donde desarrollaban procesos estocásticos (Gō, s.f.) para predecir el plegamiento de las proteínas.

T. Cover and Hart (Cover y Hart, 1967) planteaban que “para cualquier número de categorías, la probabilidad de error de la regla del vecino más cercano está limitada por encima del doble de la probabilidad de error de Bayes”. Este trabajo representó una base para realizar trabajos en temas como la biología al buscar la forma en que las proteínas se pliegan y adquieren una forma única, lo que se conoce como su estructura terciaria. Pero no solamente se aplica a esta área, también en modelos matemáticos espaciales, de física cuántica, y marcadamente en medicina.

En la biología se han desarrollado avances importantes en la detección de enfermedades que tienen un alto índice de mortalidad como el cáncer. Se han desarrollado modelos computacionales que permitan conocer uno de los factores importantes en su cura o prevención: las proteínas. Algunos de estos modelos han tenido un éxito relativo en la predicción del plegado de ellas. Los principales modelos computacionales, en este tipo de simulaciones, son la capacidad física

de las computadoras y de los algoritmos desarrollados, pues, dependiendo de la complejidad de dichos polipéptidos, será la demanda de grandes cantidades de recursos, principalmente el tiempo de uso que tarda el procesador en terminar de obtener el resultado, con base en los algoritmos establecidos. Seguramente no se tiene el código exacto y ni siquiera cercano para hacer lo que la naturaleza hace en nanosegundos, pues actualmente, una computadora tarda horas o días. De hecho, la supercomputadora Anton que en 2010 fue la supercomputadora más veloz del mundo con 512 procesadores, podía hacer una tarea 100 veces más rápido que su más cercano competidor a una velocidad de 1.000 pentafllops. La cuestión es la siguiente: se requieren 20 días de trabajo del procesador (Service, 2010) para una proteína de 20 aminoácidos y cerca de 8000 días para una de 60. La solución es mejorar los algoritmos para que la computadora pueda generar más rápidamente la solución. Y esa, ha sido una tarea que lleva décadas en desarrollo.

El primer modelo cuantitativo de interacciones electrostáticas en proteínas nativas fue propuesto por Linderstrom-Lang (Linderstrøm-Lang, 1924), basado en la teoría de Debye-Huckel. Durante la década de 1950 sus colegas encontraron un sistema para estudiar las fuerzas motrices en la formación de hélices de polipéptidos en una solución no acuosa. Aquí es donde aparece otro elemento muy importante: las fuerzas de Van der Waals. Al parecer éstas, contribuyen al estado de plegamiento y desplegamiento de las proteínas, además de su estabilidad. Otro elemento encontrado es la hidrofobicidad, donde la exposición no polar al agua podría hacer que el volumen molar disminuya, por lo que, se podría considerar que el plegado viene afectado por la fuerza que se ejerce en toda la cadena.

Ahora bien, siguiendo con los elementos clave en el plegado de las proteínas, Roland, menciona los *rotámeros* (Jr y Karplus, 1923) cuando trabaja con Brookhaven Protein Database. El modelado por homología, propuesto, alcanza de un 69 a 78 % de predicción en la configuración 3D. Otro elemento clave que ha ayudado a alcanzar una precisión muy cercana a la realidad son los algoritmos desarrollados por Tanimura donde establece tres modelos: cadena lateral-cadena lateral, teorema *dead-end elimination* (Tanimura, Kidera, y Nakamura, 1994) (DEE), interacción cadena lateral - columna vertebral, donde la primera muestra grandes porcentajes de exactitud en la predicción de plegamientos y las imágenes obtenidas de proteínas donde coinciden casi en su totalidad con las de los modelos propuestos.

Autores como Bryngelson hacen hincapié en el *panorama de energía* (Bryngelson, Onuchic, Socci, y Wolynes, 1995) donde se trabaja con tres proteínas, haciendo un análisis cuantitativo de experimentos de plegado de las mismas. Además, se ha usado para diseñar proteínas que se pueden plegar en un razonable tiempo de procesador. En este mismo sentido Guvench trabaja con un programa informático CHARMM (Guvench y MacKerell, 2008), basándose en nuevos parámetros de proteínas para la función de energía empírica de todo átomo que conforma la proteína. Algunas de ellas involucran frecuentemente en algunas de las más importantes funciones regulatorias y la falta de una estructura intrínseca queda atrás cuando se une a su molécula objetivo (Wright y Dyson, 1999). Wright comenta la importancia de conocer algunas consecuencias del mal funcionamiento y por ello la aparición de alteraciones como en la enfermedad Alzheimer, pues se origina de un mal plegamiento, aunque, se dan casos en que, incluso no estando plegadas totalmente, pueden ser funcionales bajo ciertas circunstancias. ¿Qué lo provoca? Sigue siendo una interrogante. Las cadenas poli-

peptídicas tienen miles de átomos y millones de posibles interacciones atómicas. Por cuestiones de tamaño de las proteínas que implican un uso exponencial de recursos informáticos, los trabajos han sido sobre pequeñas proteínas, del orden de 100 residuos, donde cada uno de ellos es un aminoácido (Baker, 2000).

El modelo de empaquetamiento de cadena lateral de proteínas con mayores resultados donde, autores como Bahadur, han trabajado con la aplicación de algoritmos estocásticos y ha encontrado que tienden a ser más exactos (Bahadur, Akutsu, Tomita, y Seki, s.f.), pudiendo aplicarlo en polipéptidos de más de 323 residuos. Colbes y Brizuela han desarrollado un algoritmo en el Centro de Investigación Científica y de Educación Superior de Ensenada (CICESE) llamado *empaquetamiento de la cadena lateral en proteínas* (Brizuela, Colbes, Corona, Lezcano, y Rodríguez, 2018), donde la función de energía de la proteína es la clave en el proceso de predicción del empaquetamiento y han tenido hasta un 97 % de exactitud. Las imágenes obtenidas por computadora muestran esquemas similares a los reales. Esto es un gran avance en la solución de problemáticas del sector salud. Colbes trabaja con el modelo PSCPP (*Protein Side-Chain Packing Problem*) donde alcanzan de un 77 a 87 % de éxito usando PSCPP y llega a 96.98 % usando la librería dependiente de la columna vertebral, que ya anteriormente Tanimura ha usado con gran exactitud.

Klee menciona la importancia en la salud humana de conocer las propiedades y productos de ciertas proteínas, desde la comunicación de células, diferenciación celular y desarrollo morfológico para aplicarlo desde la agricultura hasta enfermedades humanas, como detección de cánceres o desarrollo animal y la agricultura (Klee y Ellis, 2005). La fuente de datos [pdb.org](http://www.pdb.org) (([wwwPDB](http://www.pdb.org)), 1971) tiene almacenada alrededor de 130 mil proteínas y es ha sido la base para el desarrollo de trabajos que han permitido a científicos de diversas ramas del conoci-

miento desarrollar trabajos. De hecho, ahora con el desarrollo de *Big Data*, se empieza a migrar los sistemas informáticos a fin de aplicar esta nueva forma de manejar las grandes cantidades de datos para obtener la información requerida.

Un elemento importante a notar, es el formato de información que presenta la fuente de datos PDB, para poder ejecutar sus programas, pues se ha encontrado ciertos patrones irregulares en el detalle de los datos publicados. En esta plataforma, se pueden obtener datos para modelar algoritmos y conocer si son funcionales o no. Cada autor es libre de trabajar con la plataforma, pues cada estructura de cada proteína almacenada tiene una rigurosa revisión, quedando a cargo de la comunidad científica, su crecimiento y fortalecimiento. Los portales de organismos dependientes de universidades o directamente del gobierno, tales como el pdb.org ((wwPDB), 1971), donde muchos trabajos se han producido, premios Nobel han basado sus trabajos en esta fuente de datos y ofrecen sus materiales de manera libre. De igual forma, el Instituto de Ingenieros Eléctricos y Electrónicos (*The Institute of Electrical and Electronics Engineers (IEEE)*), ha sido una fuente de donde se han obtenido datos para este capítulo.

Por último, cabe mencionar que la lista de trabajos tanto a nivel biológico, de la construcción, del transporte, el arte y la cultura y más áreas que el ser humano ha desarrollado, el uso de la computadora le ha permitido desarrollar, comprender y estructurar mejor su mundo. Ciertamente hay muchas deficiencias como la contaminación, el hambre, las enfermedades, pero también es cierto que conforme pasa el tiempo, el humano va encontrando la solución y lo planteado aquí es una pequeña muestra de lo que es posible hacer, crear e innovar cuando la comunidad científica, las instituciones gubernamentales y las académicas trabajan coordinadas. Eso es tener visión, pues con el análisis intelectual y el uso adecua-

do de las herramientas informáticas, termina por convertirse en una poderosa visión por computadora.

4.5. Aportaciones internacionales

La percepción humana tiene la capacidad de adquirir, integrar e interpretar la abundante información que se nos presenta día a día, pues estamos en medio de un mundo visualmente maravilloso, que se manifiesta en muchas formas, texturas y colores (Acharya y Ray, 2005). Es hermoso capturar una puesta de sol, las nubes en el cielo, los microorganismos presentes en nuestro cuerpo, una vista aérea de la ciudad, una noche estrellada, pero resulta realmente fascinante, poder identificar los objetos presentes en cada toma e interpretar lo que se observa. Para el ser humano es una tarea que realiza inconscientemente, pues desde pequeño aprende a identificar los objetos que lo rodean, al igual que todos los seres vivos (Sossa Azuela, 2006). Cada uno identifica sus objetos de interés por sus rasgos físicos ya sea color, textura, tamaño y los clasifica según su funcionalidad (Sossa Azuela, 2006), por ejemplo, un colibrí trabaja de manera importante identificando plantas para succionar el néctar de las flores y poderse alimentar obteniendo suficientes calorías para poder volar y distribuir el polen de flor en flor. El ser humano necesita ver para poder desarrollar cualquier actividad en su vida cotidiana, identificar el alto en un semáforo, cruzar la calle, dibujar, etc. La visión es tan importante que hoy en día la ciencia a querido dotar a máquinas con ésta capacidad para crear dispositivos que identifiquen objetos, como la detección de rostros en las cámaras, sustituir al humano en tareas donde se necesita mayor precisión y rapidez o servir como auxiliar en la toma decisiones.

La visión por computadora o visión artificial es un campo de la ciencia que estudia técnicas de procesamiento de imágenes para que las computadoras puedan reconocer y localizar objetos en un medio ambiente, deduciendo de forma automática la estructura y propiedades de un mundo tridimensional a partir de imágenes bidimensionales (Gonzalez y Woods, 2002). La descripción que produce un sistema de visión por computadora debe estar relacionado con la realidad que produce la imagen procesada, debe ser una descripción numérica de la escena (Gonzalez y Woods, 2002), por ejemplo, posición del objeto, número de agujeros, tamaño, iluminación y formas. En la representación y descripción de objetos se tratan métodos que ayudan a representar la región de un objeto en términos de sus características externas es decir, su límite o su borde y en términos de sus características internas (Sossa Azuela, 2006) (Gonzalez y Woods, 2008), o sea los píxeles que componen la región. La elección de alguno de los esquemas de representación depende del enfoque que se le da a la aplicación. Pueden usarse los códigos de cadena, aproximaciones polinomiales, descriptores de Fourier, etc. para representar el borde del objeto y el color, la textura, los invariantes de momento etc. para la descripción de los píxeles que componen el objeto.

El análisis de señales y el reconocimiento de patrones son dos áreas importantes de la visión por computadora, en las que se han desarrollado diferentes algoritmos y métodos de reconocimiento de objetos en imágenes o señales, con el fin de comprender el mundo real, tomando como patrones los procedimientos que se llevan a cabo en la naturaleza para la resolución de problemas en la vida diaria mediante una computadora, como por ejemplo, la detección de huellas dactilares, la identificación de secuencias de ADN, el análisis de imágenes aéreas para la predicción del clima, la detección de cultivos, entre muchos otros. En

la modernización social y la economía en rápido desarrollo que se vive el día de hoy, la visión por computadora juega un papel importante y tiene aplicaciones en diferentes áreas como la robótica, aprendizaje e inteligencia computacional, interacción humano-computadora y tecnologías de lenguaje, entre otras.

Vehículos Autónomos

En los últimos años las ciudades han crecido a pasos agigantados al igual que la tecnología y consiguientemente han traído la necesidad de desplazamiento rápido y de forma segura en las grandes ciudades, innovaciones como la movilidad cero emisiones o el manejo autónomo. Las empresas han volcado sus esfuerzos y dinero en lo que parece ser un futuro prometedor, como los vehículos autónomos, que tienen como objetivo imitar el manejo y control del vehículo de una persona, siendo capaz de percibir el mundo que los rodea y desplazarse en una ruta trazada por el usuario, lo que representa un desafío importante para las empresas automotrices y la comunidad científica enfocada en el reconocimiento de patrones y el análisis de señales, debido a que debe procesar lo que percibe en tiempo real, con el fin de identificar los señalamientos de tránsito, peatones, obstáculos en carretera, mantener su distancia con respecto a otros vehículos, entre otras situaciones.

Los señalamientos de tráfico desempeñan un papel importante en el sistema de tráfico vial de una ciudad, sirven para regular el comportamiento de tráfico, identificar las condiciones de camino, conducir de forma segura y orientar a los peatones. El cómo identificar las señales de tráfico de forma oportuna y precisa se ha convertido en un problema urgente y un foco de investigación en los medios de transporte autónomo, debido a que la principal causa de accidentes viales son los errores humanos, la identificación oportuna de éstos señalamien-

tos puede evitar graves accidentes y garantizar la movilidad segura. Las señales de tráfico generalmente están formadas por algunos tipos de colores, como el azul, amarillo y rojo que indican su propósito y nivel de importancia (Fu, Liu, Yang, y Wang, 2014).

En la comunidad científica internacional se han desarrollado técnicas para la detección de señales, mutaciones en píxeles de fondo y segmentación de las señales de tráfico basadas en diferentes métodos de procesamiento de imágenes y reconocimiento de patrones, la mayoría de los métodos existentes se basan en la segmentación del color y la transformación morfológica, que se ven fácilmente afectados por los píxeles de fondo produciendo una interferencia muy grande para la identificación. En el artículo desarrollado en el año 2014, *Background pixels mutation detection and Hu invariant moments based traffic signs detection on autonomous vehicles* (Fu y cols., 2014), por Meng-yin Fu, Fang-yu Liu, Yi Yang y Mei-ling Wang investigadores del Instituto Tecnológico de Beijing, proponen un método donde utilizando el método de detección de mutaciones de píxeles de fondo a través de la selección de umbral global basado en la información de histograma gris, las áreas sospechosas de las señales de tráfico se extraen del fondo, y luego a través de los momentos invariantes de Hu, los valores propios se calculan para detectar las señales de tráfico. Los resultados experimentales muestran que el método propuesto tiene una rápida y buena capacidad de reconocimiento de rotación y escalado con una alta tasa de reconocimiento (Fu y cols., 2014).

Análisis de imágenes astronómicas

La observación del cielo data desde tiempos muy remotos, todas las culturas hacen referencia al cielo en sus escritos y edificaciones, pues gracias a la observa-

ción y al estudio de los astros, muchas civilizaciones construyeron calendarios precisos que les permitieron identificar temporadas para sembrar y recoger sus cosechas, predecir la migración estacional de los animales que eran su alimento ya que de eso dependía la supervivencia de cualquier grupo humano, además de tomar como puntos de orientación estrellas para su desplazamiento y descubrimientos de tierras y nuevas culturas, construyendo en sus lugares de asentamiento edificaciones orientadas a las salidas de planetas y fenómenos de la naturaleza importantes para cada cultura. La astronomía surgió cuando la humanidad dejó de ser nómada y paso a ser sedentaria. La necesidad de establecer tiempos llevo a las primeras civilizaciones a darse cuenta de los fenómenos cíclicos que se les presentaban día a día y despertar su curiosidad por la observación de los cuerpos celestes.

Por muy viejo que parezca, hoy en día la astronomía es una ciencia que estudia el comportamiento y la composición del universo, incluyendo planetas, cúmulos de galaxias, cometas, la materia oscura y todos los fenómenos ligados a ellos. Las galaxias son masas gravitacionalmente unidas de gas, polvo y miles de millones de estrellas y su clasificación e identificación según su forma, es un tema importante para los astrónomos ya que proporciona información valiosa de la evolución y el origen del universo, además de que permite a los astrofísicos probar teorías existentes o formar conclusiones nuevas que permitan entender los procesos físicos que gobiernan las galaxias y la naturaleza del universo.

Existe un esquema de clasificación que fue desarrollado en el año de 1936 por Sir Eward Hubble también conocido como la secuencia de Hubble que clasifica las fotografías de galaxias en cuatro categorías, Elípticas, Espirales, Espiral Barradas e Irregulares. El método actual de clasificación es manual y se basa en el sistema de reconocimiento visual humano, haciéndolo poco confiable debi-

do a que está sujeto a errores de clasificación especialmente cuando las imágenes contienen ruido o son borrosas (Elfattah, Elsoudaboul, y Hoon Kim, 2012). Afortunadamente, en las últimas décadas se tiene mayor acceso a grandes cantidades de datos, lo que permite el desarrollo de algoritmos de clasificación asistida por computadora con participación parcial o nula del factor humano. Recientemente, muchos algoritmos basados en aprendizaje automático se han aplicado a la clasificación automática de galaxias demostrando un éxito considerable en varios experimentos, estos algoritmos también se basan en las cuatro clases morfológicas conocidas, (Elípticas, Espirales, Espiral Barradas e Irregulares) para el análisis y la clasificación de conjuntos de datos de imágenes de galaxias.

Uno de los métodos propuestos en la comunidad científica es *Automated classification of galaxies using invariant moments* (Elfattah y cols., 2012) desarrollado por Mohamed Abd Elfattah, Mohamed A. Abu ELsoud, Aboul Ella Hassanien y Tai-hoon Kim en el año 2012, donde se propone un método de clasificación automatizado adecuado para conjuntos de datos masivos, utilizando un conjunto de imágenes obtenidas del catálogo de Zsolt Frei, que contiene 113 diferentes imágenes de galaxias proporcionado por el departamento de Ciencias Astrofísicas de la Universidad de Princeton, que se han usado como referencia para el estudio astronómico (Elfattah y cols., 2012).

Las imágenes se minimizan a dimensiones de 313×313 píxeles para facilitar la carga computacional, el primer paso es crear un vector de características de 7×1 a través del proceso de extracción de características que se realiza calculando los siete momentos invariantes de Hu (Hu, 1962) para todas las imágenes galaxia de entrenamiento y prueba. Luego, se selecciona una de las imágenes de prueba y se calcula el Error Medio Cuadrático (MSE) que se pondera de acuerdo con el vector de puntuación de Fisher. El MSE se mide entre la imagen desea-

da que es una imagen de prueba y las imágenes que constituyen el conjunto de las imágenes de entrenamiento. Y por último, decide sobre la clase de la imagen de prueba de acuerdo con la clase de la imagen correspondiente que produce la MSE mínima a través de uno de los cuatro tipos de galaxias, es decir, la espiral, el espiral barrada, elíptica y los tipos de galaxias irregulares (Elfattah y cols., 2012).

El objetivo de este trabajo fue investigar la utilidad de los invariantes de momento para la identificación automática de formas de galaxias comunes. Cuando se prueban las formas cartográficas más generalizadas, las invariantes de momento parecen funcionar. Existe una buena distinción entre las clases aunque la superposición ocurre, pero dentro de las clases la discriminación no es tan fuerte. Esto indica que los invariantes de momento por sí solos no son suficientes y para encontrar un resultado óptimo, la técnica de invariantes de momento debe compararse con otras técnicas que se están investigando actualmente. Estos incluyen descriptores de Fourier, descriptores escalares y codificación de la cadena. Además, todas las técnicas que analizan las formas del objeto de forma aislada (Elfattah y cols., 2012).

Interpretación de dibujos, planos y escritura

Los sistemas de reconocimiento de placas son una de las aplicaciones del procesamiento de imágenes y el reconocimiento de caracteres, ya que enfrentan y superan problemas de tráfico como por ejemplo, aplicar infracciones, detectar escenas de crimen, determinación de áreas de aparcamiento, estadísticas, entre otros. Esta tecnología identifica vehículos automáticamente al leer y reconocer sus placas de automóviles. Este sistema de reconocimiento está instalado en muchos lugares, como en los estacionamientos, en las entradas de edificios altamente seguros y también en las barreras de peaje. La policía está utilizando ésta apli-

cación porque puede detectar vehículos a gran velocidad desde la distancia y así poder emitir su respectiva infracción. Este tipo de sistemas resultan beneficiosos porque pueden automatizar la administración del estacionamiento y los usuarios, eliminar el uso de tarjetas de deslizamiento y multas de estacionamiento, mejorar el flujo de tráfico durante las horas pico, detectar el exceso de velocidad en las autopistas y detectar autos que no respetan los señalamiento de transito vehicular (Jusoh y Zain, 2011).

El Departamento de Carreteras y Transporte de Malasia ha aprobado una especificación para las placas de automóviles que incluye la fuente y el tamaño de los caracteres que deben seguir los propietarios de automóviles. Sin embargo, hay casos donde este estándar no se está siguiendo. Los propietarios de vehículos privados tienen a usar varios tipos de fuentes y tamaños para sus placas de automóviles. Estas fuentes y tamaños de caracteres provocan problemas durante la fase de reconocimiento, siendo incapaz de reconocer caracteres de patrones similares, como es el caso del reconocimiento de los caracteres ‘B’ o ‘3’ como ‘8’, ‘O’ y ‘D’ como ‘o’, ‘G’ y ‘5’ como ‘6’, ‘A’ como ‘4’ y ‘I’ como ‘l’ (Jusoh y Zain, 2011).

Malaysian car plates recognition using freeman chain codes and characters features (Jusoh y Zain, 2011), es una investigación desarrollada por Nor Amizam Jusoh y Jasni Mohamad Zain en el año 2011, que se centra principalmente en la realización de un experimento mediante la aplicación de códigos de cadena particularmente FCC, evaluando su precisión y eficiencia en el reconocimiento de caracteres con el apoyo de las características de caracteres y placas de automóviles malasios de una sola línea como imágenes de prueba, donde se obteniendo como resultado de la investigación que el código de cadena de Freeman se puede considerarse como una técnica de reconocimiento alternativa debido a su capa-

cidad en tiempo de reconocimiento y que al combinar el código de cadena con las características de los caracteres proporciona una precisión de reconocimiento del 95 % solo si la imagen es de alta calidad y sin ningún ruido que pueda perturbar el proceso de reconocimiento o si se combina con otra técnica (Jusoh y Zain, 2011). Otra aplicación importante es el *Análisis de Patrones Chinos Tradicionales Basado en Invariantes de Momento* (Li, Guo, Meng, y Wang, 2013), una investigación desarrollada por Xiao-Niu Li, Hai Guo, Jia-Na Meng, Bo Wang en el año 2013.

Los patrones de arte tradicional chino con su larga historia y su singularidad en el mundo, tienen características únicas. Como parte de los patrimonios culturales, suelen tener las raíces fundamentales y conservar el sabor nativo de la cultura nacional, que tiene un alto valor estético, cultural y práctico. Con el inmenso énfasis en la protección del patrimonio cultural tradicional, cada vez más las atenciones de investigación se centran en el análisis digital y el procesamiento de los patrones del arte tradicional chino. Con la ayuda de la tecnología informática, el análisis de textura de los patrones artísticos tradicionales chinos se han convertido en un problema importante (Li y cols., 2013). Ésta investigación se centra en ajustar la distribución de escala de grises de una imagen original, utilizando transformaciones de valores de grises para formar una nueva imagen con cinco niveles de grises, utilizando el histograma gris de la imagen original ya que el efecto visual de la imagen generalmente está relacionado con su histograma y cambiar la forma de éste tienen una gran influencia en las texturas de la imagen. Después, se usa el método de invariantes de geometría en el análisis y procesamiento de imágenes, se calculan los 12 momentos invariantes a la nueva imagen en gris y se usa como su vector de características de texturas. Debido a que los momentos geométricos se han utilizado ampliamente en el campo del

reconocimiento de patrones, clasificación de objetivos y que son características estáticas, se obtiene una matriz de distancia euclidiana, y por último, se lleva a cabo la clasificación utilizando el método de agrupamiento K-means en el análisis de agrupamiento utilizando la matriz de distancia euclidiana. A partir de los resultados obtenidos del cálculo de ejemplos numéricos, se puede observar que el método se puede aplicar al análisis de textura, reconocimiento y clasificación de patrones de arte tradicionales (Li y cols., 2013).

Redes Neuronales

Hay muchos tipos de especies de plantas en la tierra. Desafortunadamente, como progreso humano, cada vez más especies de plantas se encuentran al borde de la extinción. Por lo tanto, es importante reconocer correcta y rápidamente las especies de plantas para comprenderlas, gestionarlas y archivarlas antes de que sea demasiado tarde. Sin embargo, identificar correctamente las especies de plantas requiere un conocimiento experto que solo los botánicos pueden proporcionar. Debido al número limitado de botánicos, es necesario adquirir algunos de sus conocimientos y automatizar el proceso de reconocimiento (Zulkifli, Saad, y Mohtar, 2011).

La identificación de especies de plantas es un proceso en el que cada planta individual debe asignarse correctamente a series descendentes de grupos de plantas relacionadas, en función de características comunes. La clasificación de plantas no solo reconoce las diferentes plantas y nombres de la planta, sino que también proporciona las diferencias de las plantas y amplía el esquema para clasificar las plantas (Zulkifli y cols., 2011). Las plantas se pueden clasificar de acuerdo con las formas, colores, texturas y estructuras de su hoja, corteza, flor, plántula y morfología. Los métodos de taxonomía de las plantas aún adoptan el método de

clasificación tradicional, como la anatomía morfológica, la biología celular y los enfoques de biología molecular. El método tradicional consume mucho tiempo y requiere enormes esfuerzos de los botánicos. Sin embargo, debido al rápido desarrollo de las tecnologías informáticas, ahora hay oportunidades para mejorar la capacidad de identificación de especies de plantas (Zulkifli y cols., 2011). Los sistemas computarizados de clasificación de plantas se basan principalmente en imágenes bidimensionales. Esto hace que la clasificación de las plantas basada en hojas sea la elección adecuada en comparación con el uso de formas de flores, plántulas y morfología de plantas que son tridimensionalmente complejas en su estructura. La clasificación de las plantas basada en hojas implica la extracción de características foliares. La extracción de características de hojas es un proceso de identificación de características que pueden discriminar diferentes tipos de hojas (Zulkifli y cols., 2011).

La identificación de plantas vivas basada en imágenes de hojas es una tarea muy desafiante en el campo del reconocimiento de patrones y la visión por computadora. Sin embargo, la clasificación de las hojas es un componente importante del reconocimiento de la planta viva computarizada. La hoja contiene información importante para la identificación de especies de plantas a pesar de su complejidad. En el artículo, *Plant leaf identification using moment invariants & general regression neural network* (Zulkifli y cols., 2011), compara la efectividad de los momentos de Zernike (ZMI) (Khotanzad y Hong, 1990), los momentos de Legendre (LMI) (Teh y Chin, 1988) y los momentos de Tchebichef (TMI) (Mukundan, Ong, y Lee, 2001) para extraer las características de las imágenes de las hojas. Luego, las características extraídas de la técnica de momento invariante más eficaz se clasifican utilizando la red neuronal de regresión general (GRNN). Hay dos etapas principales involucradas en la identificación

de hojas de plantas. La primera etapa se conoce como proceso de extracción de características donde se aplican métodos invariantes en el momento. El resultado de este proceso es un conjunto de características vectoriales globales que representa la forma de las imágenes de las hojas. Se muestra que TMI puede extraer la característica de vector con porcentaje de error absoluto (PAE) menor que 10.38 %. Por lo tanto, la característica del vector TMI son la entrada a la segunda etapa. La segunda etapa involucra la clasificación de imágenes de hojas basadas en la característica derivada obtenida en la etapa previa. Se encuentra que los vectores de características permitieron que el clasificador GRNN alcanzara una tasa de clasificación del 100 %, pero es necesario realizar más investigaciones para confirmar esto, debido a la escasa cantidad de imágenes utilizadas pueden causar un sobreajuste del GRNN (Zulkifli y cols., 2011).

Análisis e interpretación de imágenes médicas

El reconocimiento de imágenes del diente se pueden utilizar para clasificar el tipo de diente, como molar, premolar, canino e incisivo. En el caso de desastres naturales como un tsunami, los datos dentales pueden ayudar a identificar la identidad del individuo, esto porque los dientes son fuertes y únicos para cada persona. Además, a diferencia de otros órganos humanos, los dientes no se disuelven después de la muerte. Generalmente, el humano tiene dos juegos de dientes. El primero se llama dientes deciduos, y el segundo, dientes permanentes. Las raíces de los dientes están incrustadas en el hueso de la mandíbula o en el hueso maxilar y están cubiertas por encías. Un conjunto de dientes permanentes generalmente consta de treinta y dos dientes que se pueden separar en cada diente individual en cuatro cuadrantes. Son el cuadrante superior (maxilar) derecho, el cuadrante superior izquierdo, el cuadrante inferior izquierdo

(mandibular) y el cuadrante inferior derecho. Cada cuadrante consta de ocho dientes que son dos incisivos, un canino, dos premolares y tres molares. Por lo tanto, hay dieciséis dientes en el maxilar y otros dieciséis en la mandíbula, para un total de treinta y dos dientes (Pattanachai, Covavisaruch, y Sinthanayothin, 2012).

Las radiografías dentales o las radiografías son imágenes de dientes, que se compone de huesos y tejidos blandos circundantes. Hay dos tipos de radiografías dentales: intraoral y extraoral, cada tipo con sus diferentes clases (Pattanachai y cols., 2012). Internacionalmente se presenta un artículo llamado *Tooth recognition in dental radiographs via Hu's moment invariants* (Pattanachai y cols., 2012) propuesto por Nakintorn Pattanachai, Nongluk Covavisaruch y Chanjira Sinthanayothin, este trabajo tiene el alcance del reconocimiento de los dientes molares y premolares. Propone un método para reconocer un diente en imágenes de rayos X dentales utilizando las invariantes de momento de Hu.

El método propuesto comienza con la digitalización de películas radiográficas dentales. Las radiografías dentales utilizadas en esta investigación son en forma de películas de rayos X panorámicas de tamaño 10×12 . Las películas de rayos X se escanean con un escáner de película (EPSON Perfection V700 Photo) para transformar los rayos X en formularios digitales y guardarlos en formato de archivo de mapa de bits (BMP). El proceso continúa con la mejora de la imagen de rayos X con ecualización de histograma que es una técnica para ajustar las intensidades de la imagen, para mejorar el contraste al aumentar el rango dinámico de intensidad dado a los píxeles, con los valores de intensidad más probables. La ecualización del histograma reorganiza los valores de intensidad del píxel mediante el uso de un histograma acumulativo. La intensidad del píxel de entrada de la ecualización del histograma se transforma en un nuevo nivel de in-

tensidad. Posteriormente, los dientes se segmentan con la ayuda del método de Otsu y luego se calculan los invariantes de momento de Hu que se usan como las características de los dientes para finalmente, reconocer los dientes mediante la coincidencia de características con la distancia euclidiana. Los resultados experimentales del reconocimiento dental por suma de la diferencia mínima de las invariantes de momento, demuestra que algunas imágenes son difíciles de aplicar con el método propuesto porque las imágenes son de muy mala calidad y están borrosas (Pattanachai y cols., 2012).

Análisis de las imágenes para la compresión

En los últimos años, se transfirió un volumen incomparable de información textual a través de Internet a través de correo electrónico, chat, blogs, Twitter, bibliotecas digitales y sistemas de recuperación de información y como el volumen de datos de texto ha superado el 40 % del volumen total de tráfico en Internet, la compresión de datos textuales se vuelve imprescindible. Se introdujeron y emplearon muchos algoritmos para este propósito, incluida la codificación Huffman, la codificación aritmética, la familia Ziv-Lempel, la Compresión Dinámica de Markov y la Transformación Burrow-Wheeler (Alzahir, 2011).

Técnicamente, los códigos de longitud fija como ASCII, EBCDIC y Unicode se han utilizado para codificar caracteres de idiomas naturales, como el árabe o el inglés, para manipular y procesar computadoras. De manera similar, muchos esquemas de compresión representan símbolos de lenguaje a través de códigos de longitud fijos o variables basados en algunos modelos estadísticos, tales como diagramas o análisis de triagramas de muestras de texto, también se presentan en la literatura. Algunas de esas técnicas dieron como resultado un tamaño mayor que la imagen textual original u ofrecieron una relación de compresión mode-

rada. Estos métodos fueron específicos del idioma. En principio, cada idioma tiene sus especificidades y características. Algunos idiomas tienen una gran cantidad de signos diacríticos y una morfología compleja, como el árabe, el persa y el hebreo, y algunos idiomas no, como el inglés. Muchas de estas características abarcan impedimentos a la compresión de texto. La mayoría de las técnicas de compresión de texto no tomaron tales características y especificidades con seriedad y algunas de ellas, como el algoritmo de Burrows - Wheeler, son insensibles a las conjugaciones de palabras que provienen de la misma raíz y, por lo tanto, no pueden usar la morfología del lenguaje. Estas técnicas de compresión se diseñaron para la compresión de textos en inglés, que es un lenguaje más estructurado. Por ejemplo, muchas de las secuencias, como zx, bv o qe, son gramaticalmente ilegales o inexistentes. En el caso de los lenguajes estructurados, la negligencia morfológica puede ser un error significativo en idiomas como el árabe, el persa y el hebreo, donde las reglas gramaticales permiten mucha más libertad en las secuencias de letras. Por lo tanto, comprimir sin tener en cuenta tales reglas dará como resultado una relación de compresión más baja (Alzahir, 2011). El modelado del lenguaje árabe plantea un desafío para los lingüistas informáticos debido a su rica morfología y su compleja formación de palabras. En consecuencia, la construcción de modelos estadísticos y el estudio de la entropía del idioma árabe para determinar un límite superior para la compresión, por ejemplo, se convierte en un proceso complejo ya que se consideran n-gramas más grandes. Debido a que la entropía impone un límite teórico a la compresión.

La compresión del texto árabe se vuelve más difícil ya que requiere un análisis lingüístico de los caracteres y los tallos. Tal manifestación llevó a algunas técnicas de compresión de texto árabe en la literatura (Alzahir, 2011). En el artículo, *A fast lossless compression algorithm for Arabic textual images* (Alzahir, 2011) se

presenta un nuevo algoritmo para comprimir imágenes textuales. El algoritmo propuesto contiene dos componentes: el primer componente es el modelo de libro de código que está diseñado para proporcionar un código (un número de bits) para cada bloque de 8×8 de la imagen textual o para reemplazar el bloque por sus bits comprimidos utilizando codificación aritmética basada en el probabilidades disponibles en el libro de códigos. El segundo componente es el esquema RCEC, que se emplea en aquellos bloques que no se encuentran en el libro de códigos. Los dos componentes comprimen casi el 98.5 % del total de bloques en cualquier imagen textual.

Los resultados experimentales del algoritmo propuesto en la base de datos de imágenes textuales en árabe muestran que tiene una relación de compresión constante del 87 % para los conjuntos donde JBIG2 tiene una mayor variabilidad en la relación de compresión. Debido a la forma en que está diseñado, el algoritmo propuesto se puede utilizar para comprimir imágenes textuales, independientemente de sus tipos de fuente, tamaño de fuente y si la página contiene una o dos columnas. Además, este algoritmo se puede generalizar y usar para comprimir otras imágenes textuales de diferentes idiomas, manteniendo una alta tasa de compresión (Alzahir, 2011).

4.6. Resultados recientes en México

Ahora vamos a platicar sobre algunos de los resultados que han sido desarrollados en México usando la idea de códigos de cadena.

Equivalencia de Códigos

Un código de cadena nos sirve para codificar una imagen en 2D o un objeto en 3D. En otras palabras, convierte una imagen o un objeto en una cadena de símbolos, de tal forma que si queremos estudiar el objeto, sólo tenemos que estudiar dicha cadena. Los códigos de cadena mas conocidos para codificar imágenes en 2D son F4 (Freeman, 1974), F8 (Freeman, 1961), AF8 (Liu y Žalik, 2005), 3OT (Sánchez-Cruz y Rodríguez-Dagnino, 2005) y VCC (Bribiesca, 1999). Los métodos mas conocidos para codificar objetos en 3D son 3DRC (Sánchez-Cruz, López-Valdez, y Cuevas, 2014) y ODCCC (Bribiesca, 2000).

A un arreglo rectangular de números se le conoce como **matriz**. Por ejemplo, el siguiente arreglo rectangular de 6 números

$$\begin{pmatrix} 0 & 2 & 1 \\ 3 & 5 & 10 \end{pmatrix}$$

es una matriz, la cual tiene dos filas y tres columnas. Este concepto de matriz ha sido usado desde hace varios cientos de años, y tiene aplicaciones en prácticamente todas las áreas de las ciencias que uno se pueda imaginar. Por ejemplo, si queremos que un robot mueva su brazo de un punto a otro punto, tenemos que indicar cuantos grados tiene que rotar su brazo, y para lograr ese objetivo, se usan las llamadas matrices de rotación, las cuales son matrices que están relacionadas con el movimiento de rotación.

Dada la importancia de las matrices en el mundo de las ciencias, no es de sorprender que también aparezcan en los códigos de cadena. En el trabajo “*Equivalence of chain codes*” (Sánchez-Cruz y López-Valdez, 2014), el cual fue desarrollado por los doctores Hermilo Sánchez-Cruz e Hiram H. López, los autores encuentran una nueva y elegante manera de pasar de un código en 2D a otro código en 2D. Antes de continuar, contestemos la pregunta ¿qué significa pasar

de un código en 2D a otro código en 2D?. Dada una imagen binaria en 2D, ya hemos visto que existen diferentes códigos de cadena con los cuales podemos encapsular la información de la imagen. Por ejemplo, podemos hacerlo con F4 (Freeman, 1974), F8 (Freeman, 1961), AF8 (Liu y Žalik, 2005), 3OT (Sánchez-Cruz y Rodríguez-Dagnino, 2005) o VCC (Bribiesca, 1999). Ahora bien, pasar de un código en 2D a otro código en 2D significa pasar de un código de cadena de una imagen dada a otro código de cadena de la misma imagen dada. Antes de la publicación del trabajo *Equivalence of chain codes* (Sánchez-Cruz y López-Valdez, 2014), la única manera que existía para pasar de un código en 2D a otro código en 2D, era a través del objeto dado. En otras palabras, si queríamos ir de un código a otro, teníamos que codificar la imagen dos veces, cada una de ellas usando los códigos deseados. Pero en el trabajo *Equivalence of chain codes* (Sánchez-Cruz y López-Valdez, 2014) se demostró que existen matrices las cuales se pueden usar para pasar de un código a otro. La idea es que sólo se necesita codificar la imagen usando un código, y después, usando matrices, es posible encontrar el resto de los códigos.

Cadenas simples para representar grupos de objetos

Un código de cadena sirve para encapsular la información de una imagen en 2D. Pero si se tienen varias imágenes en 2D, ¿Cómo podemos encapsular la información de todas las imágenes usando códigos de cadena?. Antes del trabajo *Single chains to represent groups of objects* (López-Valdez, Sánchez-Cruz, y Mascorro-Pantoja, 2016), la creencia era que se tenían que encadenar cada una de las imágenes y entonces, guardar cada una de las cadenas en archivos diferentes, para que la información no se mezclara. Pero en el trabajo *Single chains to represent groups of objects* (López-Valdez y cols., 2016) se demostró que es po-

sible guardar toda la información en una sola cadena. La idea es muy sencilla, simplemente se tienen que concatenar las cadenas de los objetos. En este trabajo, los autores demuestran que gracias a las propiedades de los códigos de cadena, la información no se mezcla, y dada la cadena concatenada, es posible recuperar cada una de las cadenas de los objetos.

Un nuevo código de cadena relativo en 3D

La idea de este trabajo es extender la codificación 3OT (Sánchez-Cruz y Rodríguez-Dagnino, 2005), que es el que se encarga de codificar una imagen binaria en 2D usando tres vectores consecutivos, al caso de objetos en 3D. Lo que se pretende es claro, codificar objetos en 3D usando tres vectores consecutivos, pero la respuesta no es tan sencilla. La razón es porque si vamos en un automóvil que no tiene reversa, entonces, al llegar a un crucero sólo tenemos tres opciones: continuar derecho, vuelta a la derecha ó vuelta a la izquierda. Por esta razón el código 3OT, a pesar de que codifica usando tres vectores consecutivos, necesita solamente tres símbolos para crear la cadena. Pero en el caso de tres dimensiones se tienen mas opciones, pues existen 25^2 posibles combinaciones de tres vectores consecutivos. Pero a pesar de este número tan grande de combinaciones, en el trabajo (Sánchez-Cruz y cols., 2014) se demuestra que se necesitan solamente 25 símbolos para crear la cadena de un objeto en 3D usando combinaciones de tres vectores consecutivos.

Compresión

En la Universidad Autónoma de Aguascalientes (UAA), el equipo de investigación en Visión Artificial y Reconocimiento de Patrones (VARP), en conjunto con el ITESM-Monterrey, la UNAM y el CINVESTAV, han desarrollado dife-

rentes algoritmos para codificar objetos binarios. En 2005 propusieron un método innovador para representar objetos binarios en imágenes. Aprovechando dicho código, propusieron un método de compresión sin pérdida de información. De acuerdo con sus resultados, encontraron que este método era adecuado para la representación de imágenes bi-nivel, obteniendo un 25 % más eficiencia en compresión que el método de código de cadena de Freeman, y un 29 % mejor que el estándar internacional para imágenes bi-nivel: *Joint Bilevel Image Experts Group (JBIG)* (Sánchez-Cruz y Rodríguez-Dagnino, 2005). Posteriormente, en 2007, con el fin de encontrar hasta dónde su método era bueno, aplicaron códigos de cadena para representar objetos con diferentes formas y con diferente cantidad de hoyos (genus) en ellas. Compararon su eficiencia en la compresión con seis códigos de cadena diferentes, así como con JBIG, hasta entonces desarrollados todos estos métodos de manera independiente y por diferentes investigadores. (Sánchez-Cruz, Bribiesca, y Rodríguez-Dagnino, 2007). El equipo VARP ha participado en la búsqueda de descriptores utilizando los códigos de cadena, así como la equivalencia entre los mismos, descubriendo que existen diferentes familias, donde todas ellas, desarrolladas a lo largo del tiempo, se descubrieron que son equivalentes (Sánchez-Cruz y López-Valdez, 2014). También han sacado ventaja en sus propiedades de concatenación (López-Valdez y cols., 2016). Otra las ventajas que el equipo VARP ha obtenido con el uso de los códigos de cadena es en la compresión de documentos en imágenes, compitiendo y superando estándares internacionales, tales como JBIG1, JBIG2 y DjVu (Sánchez-Cruz y Rodríguez-Díaz, 2012) (Rodríguez-Díaz y Sánchez-Cruz, 2014).

Referencias

- Acharya, T., y Ray, A. K. (2005). *Image processing: principles and applications*. New Jersey: John Wiley & Sons.
- Adana, A., y Huberb, D. (2011). Análisis de Datos 3D Para Generación Automática de Modelos BIM de Interiores Habitados,. *Revista Iberoamericana de Automática e Informática Industrial RIAI*, 8(4), 357-370.
- Alzahir, S. (2011). A fast lossless compression algorithm for arabic textual images. En *ICSIPA* (pp. 595-598).
- Anda, C., Erath, A., y Fourier, P. J. (2017). Transport modelling in the age of big data, (2017). *International Journal of Urban Sciences*, 21(sup1), 19-42.
- Bahadur, K. C., Akutsu, T., Tomita, E., y Seki, T. (s.f.).
En 04, no., 29 no. APBC.
- Baker, D. (2000). A surprising simplicity to protein folding. *Nature*, 405(6782), 39.
- Bribiesca, E. (1999). A new chain code. *Pattern Recognition*, 32(2), 235-251.
- Bribiesca, E. (2000). A chain code for representing 3d curves. *Pattern Recognition*, 33(5), 755-765.
- Brizuela, C. A., Colbes, J., Corona, R. I., Lezcano, C., y Rodríguez, D. (2018). Protein side-chain packing problem: is there still room for improvement. *Briefings in Bioinformatics*, 18(6), 1033-1043.
- Bryngelson, J. D., Onuchic, J. N., Socci, N. D., y Wolynes, P. G. (1995). Funnel pathways, and the energy landscape of protein folding: A synthesis. *Proteins: Structure, Function, and Bioinformatics*, 21(3), 167-195.
- Cover, T., y Hart, P. (1967). Nearest neighbor pattern classification. *IEEE Transactions on Information Theory*, 13(1), 21-27.

- Domingo, M. (2004). *Visión por computador*. Santiago de Chile, Chile: Universidad Católica de Chile, Departamento de Ciencia de la Computación. Descargado de <http://dmery.sitios.ing.uc.cl/Prints/Books/2004-ApuntesVision.pdf>
- Duda, R. O., y Hart, P. E. (1973). *Pattern recognition and scene analysis*. Wiley, New York.
- Elfattah, M. A., Elsoudaboul, M. A. A., y Hoon Kim, E. H. (2012). Automated classification of galaxies using invariant moments. *Lecture Notes in Computer Science*, 7709, 103-111.
- Freeman, H. (1961). On the encoding of arbitrary geometric configurations. *IRE Transactions on Electronic Computers*, EC-10(2), 260-268.
- Freeman, H. (1974). Computer processing of line-drawing images, (CSUR). *ACM Computing Surveys*, 6(1), 57-97.
- Fu, M. Y., Liu, F. Y., Yang, Y., y Wang, M. L. (2014). Background pixels mutation detection and hu invariant moments based traffic signs detection on autonomous vehicles. En *Control conference (ccc)* (p. 670-674). IEEE.
- Gō, N. (s.f.). Protein folding as a stochastic process journal of statistical physics, 30 (1983). no., 2, 413-423.
- Gonzalez, R. C., y Woods, R. E. (2002). *Digital image processing* (second ed.). Beijing: Publishing House of Electronics Industry.
- Gonzalez, R. C., y Woods, R. E. (2008). *Digital image processing*. Upper Saddle River, NJ: Pearson Prentice Hall.
- Guvench, O., y MacKerell, A. D. (2008). Comparison of protein force fields for molecular dynamics simulations. En *Molecular modeling of proteins* (pp. 63-88). Springer.

- Haralick, R. M., y Shapiro, L. G. (1991). Glossary of computer vision terms. *Pattern recognition*, 24(1), 69–93.
- Hartley, R., y Zisserman, A. (2003). *Multiple view geometry in computer vision*. Cambridge university press.
- Hu, M.-K. (1962). Visual pattern recognition by moment invariants. *IRE Transactions on information theory*, 8(2), 179–187.
- Jain, A. K., Duin, R. P. W., y Mao, J. (2000). Statistical pattern recognition: A review. *IEEE Transactions on pattern analysis and machine intelligence*, 22(1), 4–37.
- Jr, R. L. D., y Karplus, M. (1923). Backbone-dependent rotamer library for proteins application to side-chain prediction. *Journal of Molecular Biology*, 230(2), 543–574.
- Jusoh, N. A., y Zain, J. M. (2011). Malaysian car plates recognition using Freeman chain codes and characters' features. En *International conference on software engineering and computer systems* (pp. 581–591).
- Khotanzad, A., y Hong, Y. H. (1990). Invariant image recognition by Zernike moments. *IEEE Transactions on pattern analysis and machine intelligence*, 12(5), 489–497.
- Klee, E. W., y Ellis, L. B. (2005). Evaluating eukaryotic secreted protein prediction. *BMC bioinformatics*, 6(1), 256.
- Li, X.-N., Guo, H., Meng, J.-N., y Wang, B. (2013). Traditional chinese patterns analysis based on moment invariants. En *2013 ieee 9th international conference on mobile ad-hoc and sensor networks* (pp. 582–585).
- Linderstrøm-Lang, K. (1924). On the ionization of proteins. *CR Trav. Lab. Carlsberg*, 15(7), 1–29.

- Liu, Y. K., y Žalik, B. (2005). An efficient chain code with Huffman coding. *Pattern Recognition*, 38(4), 553–557.
- López-Valdez, H. H., Sánchez-Cruz, H., y Mascorro-Pantoja, M. C. (2016). Single chains to represent groups of objects. *Digital Signal Processing*, 51, 73–81.
- Mukundan, R., Ong, S., y Lee, P. A. (2001). Image analysis by Tchebichef moments. *IEEE Transactions on image Processing*, 10(9), 1357–1364.
- Musleh, B., y A. De-la Escalera-Hueso, J. M. A. (2012). Detección de obstáculos y espacios transitables en entornos urbanos para sistemas de ayuda a la conducción basados en algoritmos de visión estéreo implementados en GPU. *Revista iberoamericana de automática e informática industrial*, 9(4), 462–473.
- Pattanachai, N., Covavisaruch, N., y Sinthanayothin, C. (2012). Tooth recognition in dental radiographs via Hu's moment invariants. En *2012 9th international conference on electrical engineering/electronics, computer, telecommunications and information technology* (pp. 1–4).
- Rodríguez-Díaz, M. A., y Sánchez-Cruz, H. (2014). Refined fixed double pass binary object classification for document image compression. *Digital Signal Processing*, 30, 114–130.
- Sánchez-Cruz, H., Bribiesca, E., y Rodríguez-Dagnino, R. M. (2007). Efficiency of chain codes to represent binary objects. *Pattern Recognition*, 40(6), 1660–1674.
- Sánchez-Cruz, H., y López-Valdez, H. H. (2014). Equivalence of chain codes. *Journal of Electronic Imaging*, 23(1), 013031.
- Sánchez-Cruz, H., López-Valdez, H. H., y Cuevas, F. J. (2014). A new relative chain code in 3D. *Pattern Recognition*, 47(2), 769–788.

- Sanchez-Cruz, H., y Rodriguez-Dagnino, R. M. (2005). Compressing bilevel images by means of a three-bit chain code. *Optical engineering*, 44(9), 097004.
- Sanchez-Cruz, H., y Rodríguez-Díaz, M. A. (2012). Binary document image compression using a three-symbol grouped code dictionary. *Journal of Electronic Imaging*, 21(2), 023013.
- Service, R. F. (2010). Carbon-linking catalysts get nobel nod. *Science*, 330(6002), 308–309. Descargado de <https://science.sciencemag.org/content/330/6002/308.2> doi: 10.1126/science.330.6002.308-b
- Solem, J. E. (2012). *Programming computer vision with python: Tools and algorithms for analyzing images*. O'Reilly Media, Inc.
- Sossa Azuela, H. (2006). Rasgos descriptores para el reconocimiento de objetos. *SEP. Instituto Politécnico Nacional. Primera edición. México*.
- Tanimura, R., Kidera, A., y Nakamura, H. (1994). Determinants of protein side-chain packing. *Protein Science*, 3(12), 2358–2365.
- Teh, C.-H., y Chin, R. T. (1988). On image analysis by the methods of moments. *IEEE Transactions on pattern analysis and machine intelligence*, 10(4), 496–513.
- Wright, P. E., y Dyson, H. J. (1999). Intrinsically unstructured proteins: reassessing the protein structure-function paradigm. *Journal of molecular biology*, 293(2), 321–331.
- (wwPDB), T. W. P. (1971). *The protein data bank archive (PDB)*. <https://www.wwpdb.org/>.
- Zulkifli, Z., Saad, P., y Mohtar, I. A. (2011). Plant leaf identification using moment invariants & general regression neural network. En *2011 11th inter-*

national conference on hybrid intelligent systems (HIS) (pp. 430–435).

Capítulo 5

La Minería de Datos

ADOLFO GUZMÁN-ARENAS, CIC-IPN

La información digital disponible en la actualidad ha crecido en grandes volúmenes. Esta información es fácil ahora adquirirla o generarla, almacenarla, transmitirla y procesarla. Este asunto da lugar a un fenómeno curioso: se tiene a nuestra disposición una gran cantidad de datos de interés, pero no es fácil analizarla a simple vista, o con métodos sencillos. Se quiere extraer información relevante de estos datos, útil para tomar decisiones. ¿Qué nos quieren decir los datos? ¿Qué tendencias hay en ellos? ¿Qué secretos esconden? Muchas de las técnicas de este libro abordan este problema, y en este capítulo se presenta una herramienta más: la minería de datos.

5.1. Qué es la minería de datos

La minería de datos es la detección y extracción de tendencias, patrones, desviaciones y anomalías encontradas en un conjunto de datos. Precisando, estas

tendencias, patrones, desviaciones y anomalías, deben ser no triviales, implícitas (derivables de los datos), previamente desconocidas y potencialmente útiles. Cuando se trabaja en minería de datos, se piensa en la fastidiosa labor de un minero humano: buscar entre toneladas de tierra, rocas y mineral, las pequeñas gemas y diamantes que se esconden entre tanto lodo.

Otros nombres de la minería de datos: descubrimiento del conocimiento en bases de datos, minado del conocimiento, análisis de datos y hallazgo de patrones, filtrado selectivo de datos, analítica predictiva, inteligencia de negocios.

Normalmente, los datos a analizar se organizan en bases de datos relacionales. Es información estructurada. Esto significa que cada renglón (un estudiante con sus calificaciones, carrera, edad, género, y otros atributos) de una tabla está organizado de la misma forma. No es válido que unas fechas sean descritas como 20 de julio de 2018, y otras como 2018/VII/20. Si se usan tablas de distinto origen (del área de inscripciones y del área de becas, por ejemplo) los nombres y otros atributos deben coincidir –habrá que hacer una conversión o equivalencia si en una la estudiante es Gpe. Pérez Trujillo, y en otra la becaria es Pérez Trujillo Guadalupe. Obvio decir que los datos deben de estar limpios, libres de errores de ortografía, de errores de dedo, meses que no existen, lugares de nacimiento imaginarios, etc. Todo esto da origen a un proceso de *limpieza y generalización* de la información fuente, antes de estar lista para ser analizada por el minero. Desafortunadamente, este proceso es por lo común lento y fastidioso; aproximadamente la mitad del tiempo que consume una labor de minería de datos, se destina a la limpieza, generalización y estandarización de los datos. A menudo los libros y cursos de la minería de datos omiten o soslayan este trabajo, con resultados funestos: basura entra, basura sale.

También es común hacer minería de datos sobre documentos de texto (en español, por ejemplo), y se habla de minería de texto. Análisis de opiniones, de sentimientos, de encuestas. ¿Está satisfecho el cliente en esta carta que me envía, o se queja, o simplemente pregunta precios? Aquí son útiles técnicas de lingüística computacional: lematización (suprimir declinaciones), sinónimos, jerarquías de conceptos (hiperónimos o superclases, hipónimos o especializaciones). También hay minería de datos en imágenes, en señales (sonidos, electrocardiogramas), en videos.

Resumen: La minería de datos es un campo interdisciplinario. Involucra métodos científicos, procesos y sistemas para extraer conocimiento o un mejor entendimiento de datos en sus diferentes formas, ya sean estructurados o no estructurados (Han, Kamber, y Pei, 2011).

5.1.1. Qué no es la minería de datos

- Una búsqueda simple (con un comando *select* del Lenguaje de Consulta estructurado *Structured Query Language*, SQL, por ejemplo).
- Procesamiento de preguntas en una base de datos.
- Cuando los datos extraídos provienen de un fenómeno cuyo modelo se conoce, es a menudo más útil usar este modelo y ajustar sus parámetros a los datos obtenidos, para que el fenómeno así modelado arroje la información útil que se busca. Por ejemplo, si se tienen datos del calentamiento y deformación de una placa de hierro expuesta a una fuente puntual de calor (una llama), es mejor usar la ecuación de Laplace para predecir, por ejemplo, los desplazamientos de la esquina inferior de ella. Por otra parte, si se conocen los síntomas de la diabetes y de la hepatitis, es mejor usar-

los (en vez de emplear un minero de datos) para determinar cuáles de los pacientes de mi base de datos tienen qué enfermedad. Además, se puede usar un modelo y también un minero, para ver “qué más hallo”.

- Un sistema experto que hace deducciones.

5.1.2. Qué es la Ciencia de Datos

Cuando los datos son muchos, la minería de datos adquiere el nombre de «*Big Data*» (en inglés). ¿Cuándo ya son muchos, para cambiar el nombre? Respuesta subjetiva, a menudo influenciada por el deseo de exagerar. Se considera que cuando los datos no caben en un disco (de un terabyte, por ejemplo), ya son muchos. O cuando caben en un disco, pero su procesamiento por una sola máquina sería demasiado lento. Se usan entonces varias máquinas que entre sí llevan a cabo el minado. Si los procesos que corren en máquinas distintas están débilmente acoplados (los cambios de estado en uno influyen sólo de vez en cuando en los cambios de estado en otro –en otras palabras, si el intercambio de información entre dos procesos es escaso o esporádico), se usa un arreglo de PCs bajo Hadoop y MapReduce, normalmente. Si los procesos están fuertemente acoplados (interactúan a menudo, intercambian información frecuentemente), entonces se usa un grupo («*cluster*») de procesadores, entre los cuales pueden fluir varios mensajes simultáneamente. Afortunadamente, casi siempre el minado de muchos datos puede repartirse entre procesos que intercambian poca o ninguna información entre sí. Es creciente el uso del procesador gráfico (GPU, «*Graphic Processing Unit*») además de la unidad central de proceso (CPU, «*Central Processing Unit*») de cada PC.

5.1.3. Riesgo: uso poco ético de la Ciencia de Datos

La acumulación sistemática de datos por parte de los gobiernos y algunas empresas (Google, Facebook) mediante la vigilancia electrónica, se ha convertido en un fin en sí mismo y no en una forma de alcanzar otros objetivos comerciales complementarios. Esto plantea cuestiones de ética, invasión a la privacidad y abuso de poder, aun en presencia de “accidentes normales” (Según la Teoría de Accidentes Normales, los accidentes normales son normales en el sentido de que estos eventos negativos son inevitables y ocurren cuando los sistemas organizacionales son complejos y están estrechamente relacionados).

5.2. Usos de la minería de datos

La información sistemáticamente organizada sugiere su análisis con minería de datos, en vez de un enfoque tradicional. Por ejemplo:

- Cuando los datos son tantos que ya no caben en memoria principal.
- Cuando los datos tienen gran número de dimensiones (atributos).
- Datos provenientes de torrentes de datos (los clics en una página web, llamadas telefónicas, solicitud de acceso a ciertas direcciones IP) y de sensores (cámaras, micrófonos). Los datos se procesan conforme llegan, en vez de almacenarlos y procesarlos después.
- Series de tiempo (señales), datos temporales, datos secuenciales.
- Datos estructurados, gráficas, redes sociales, o con muchas ligas entre sí (páginas web).

- Bases de datos heterogéneas; bases de datos históricas o heredadas.
- Datos espaciales, espacio-temporales, multimedia; textos; datos obtenidos de la Web.
- Bitácoras del sistema operativo; simulaciones científicas. Electroencefalogramas.

5.2.1. Enfoques

¿Qué se obtiene de la minería de datos, qué resultados puede arrojar?

- Caracterización y discriminación. Generaliza, resume, y contrasta las características de diferentes grupos. Entender cómo están distribuidos los datos, qué enfermedades son las más frecuentes, cuáles son las edades típicas del estudiante de preparatoria. Regiones secas contra regiones húmedas.
- Asociación. Patrones frecuentes. Aquéllos que compran pan también compran leche. La mayoría de los fumadores fuma cigarrillos con filtro. ¿Hay correlación, o es casualidad? Encuentra asociaciones y correlaciones entre las ventas de ciertos productos, y predice basado en tales asociaciones. Identifica los mejores productos para distintos tipos de consumidores.
- Clasificación y predicción. Agrupa los datos en conjuntos con características similares a los de ciertas clases dadas de antemano: buenos pagadores, regulares, malos. Predice, dado un nuevo solicitante de crédito, qué clase de pagador será. Construye modelos (funciones) que describen clases o conceptos para predicción futura. Clasifica países según su clima, o automóviles según su rendimiento de gasolina. Pronostica algún valor

desconocido de un cierto objeto (basado en su parecido a otros objetos similares). ¿Qué tipo de clientes compran ciertos productos?

- Agrupamiento («*clustering*»). Agrupa los objetos en conjuntos (clases) con características similares, pero distintas de las características de los objetos de otra clase. Aquí, no se conocen los nombres de las clases, ni cuántas son. Agrupa los estudiantes de secundaria que obtuvieron resultados por debajo del nivel básico en el Examen de Calidad y Logro Educativo (EXCALE), en varios grupos de “estudiantes con bajos resultados”. Maximiza la similitud dentro de cada clase, y minimiza la similitud entre clases. Encuentra grupos de “clientes modelo” que comparten las mismas características: interés, nivel de ingreso, hábitos de consumo, etc. Determina el patrón temporal de compra de clientes.
- Tendencias y desviaciones. Regresión. Análisis de periodicidad. Patrones y reglas de asociación: Automóvil_BMW → Propietario_rico. Análisis basado en casos similares.
- Hallar objetos anómalos. Detecta objetos que no concuerdan con el comportamiento general de los datos. ¿Es ruido o excepción? Descubrimiento de fraude en tarjetas de crédito, análisis de eventos raros. Aparece una estrella supernova en el firmamento nocturno. Detección de computadoras que están infectadas y envían correo no deseado.

5.2.2. Pasos claves en la minería de datos

1. Entender el dominio de la aplicación. ¿Qué se quiere lograr? ¿Cuál es la meta de este trabajo? ¿Cuáles son los objetos a analizar (estudiantes, ventas, viajes, discursos, opiniones)? ¿Cuáles son sus características princi-

pales (longitud típica de un discurso, rango usual de calificaciones en la prueba EXCALE, gasto promedio de un comprador en un supermercado)? ¿Qué tan a menudo una señora que ingresa a un hospital para parto, engendra gemelos?. Familiarizarse con los datos.

2. Decidir cuáles datos se deberán analizar. En las bases de datos operacionales hay datos irrelevantes para la minería de datos (número de teléfono, dirección del padre).
3. Diseñar la estructura de tablas que albergará a los datos limpios y preprocesados (Paso (4)). A menudo esto requiere construir una *bodega de datos*. Definición: una bodega de datos es un conjunto de datos integrados u orientados a un objetivo específico, que varían con el tiempo (datos históricos) y que no son transitorios. Los datos provienen de distintas fuentes, están *limpios* (libres de errores) y estandarizados. Los datos en una bodega no cambian: cada semana o mes (por ejemplo) los datos se extraen de las bases de datos operacionales y se agregan (fuera de línea, no en forma transaccional) en una nueva capa a la bodega: datos de julio de 2018, datos de agosto de 2018, etc. La estructura de la bodega no es adecuada para el manejo de transacciones (proceso transaccional), sino para estudios estratégicos, de corto, mediano y largo plazo. Es adecuada para manejar consultas y facilitar la minería de datos. Estos datos soportan la toma de decisiones; la bodega está orientada a recibir y manejar grandes volúmenes de datos provenientes de diversas fuentes, que originalmente estaban en diversos formatos. Contiene información histórica, actualizada, de la empresa, útil para hacer análisis de tendencias, patrones frecuentes, y en general para entender la evolución de la empresa.

4. Limpieza y preprocesamiento de datos. Convertir los datos seleccionados en el Paso (2) a la estructura diseñada en el Paso (3). A menudo este proceso consume entre el 50 y 60 % del esfuerzo (tiempo) dedicado a la minería.
5. Reducción y transformación de datos. Homogeneizar la representación de los datos. Encontrar variables (atributos) útiles. Eliminar detalles innecesarios. Agrupar las edades en rangos (niño, adolescente, adulto, etc.). Generalizar atributos. Hallar representaciones invariantes.
6. Escoger las funciones que se usarán para el minado. Funciones de resumen, clasificación, regresión, asociación, agrupamiento.
7. Escoger los algoritmos de minería. Minado de patrones de interés. Conjuntos de ítems (objetos) que ocurren frecuentemente. Cambios en el tiempo. Hallar objetos anómalos.
8. Evaluar los patrones y resultados obtenidos, así como representar este conocimiento de forma entendible al usuario final. Uso de visualización, gráficas de barras, de pastel, etc. Remover patrones redundantes.
9. El usuario final usa y explota el conocimiento y resultados obtenidos en el Paso (8).

5.2.3. Pilares de la minería de datos

La minería de datos descansa en tres pilares. El primer pilar es el manejo de bases de datos. Es necesario que quien se dedique a minería disponga de un manejo adecuado, fluido, de bases de datos, y que conozca técnicas para manejar grandes conjuntos de datos sin perder tiempo. Indexación o «*hashing*», creación de

diccionarios invertidos, por ejemplo; manejo de llave única, llave primaria y llave foránea. El segundo pilar es la estadística. Esta ciencia estudia las propiedades de un subconjunto o muestra de los datos, para predecir las propiedades de todos los datos. La minería de datos aplica la estadística, excepto que no usa una muestra de datos, sino usa *todos* los datos disponibles para hallar un comportamiento confiable, y discriminar el ruido o la casualidad. El tercer pilar es el Reconocimiento de Patrones, o más generalmente, la Inteligencia Artificial. Clasificación supervisada, clasificación no supervisada, medidas de similitud, aprendizaje máquina, reconocimiento de patrones, redes bayesianas, redes neuronales, algoritmos genéticos, análisis de texto, generalización, extrapolación, entre otras. Todas estas herramientas y métodos se usan en minería de datos. Poco se usan la planeación y los juegos de mesa (ajedrez).

Un poderoso aliado de la minería de datos es la *Visualización de la información*, porque los hallazgos del minero deben ser entendibles por el usuario final, que normalmente no maneja la estadística o la computación. Una representación adecuada facilita entender lo que los datos nos quieren decir. Una buena imagen vale más que cien palabras.

5.3. Aplicaciones y ejemplos

En general, la minería de datos es una herramienta empleada por la gerencia de nivel medio (Gerente regional o de producto, Director de un hospital o de una escuela) y alto (Rector de una universidad, Secretario de Turismo, Director de un Centro de investigación o de una empresa), y no por los departamentos operacionales (los que llevan el quehacer cotidiano de la organización, su *labor sustantiva*). Es muy útil en planeación, toma de decisiones, cambios de estrate-

gia, análisis de mercados. La minería de datos se aplica en diversos ámbitos, por ejemplo: cadenas de supermercados, cadenas de hospitales, escuelas, análisis de estudiantes y tendencias, comercio minorista, publicidad dirigida, telecomunicaciones, diseño de servicios y productos especiales para ciertos tipos de clientes, diseño de productos bancarios, análisis de fraude, de epidemias, de opiniones, análisis bursátil, minería de textos, de datos biológicos, tendencias de las redes sociales, minería web, entre otros.

A continuación se lista una breve descripción de algunas aplicaciones:

- No siempre mientras más datos es mejor. Angel Kuri (Kuri-Morales y Rodríguez-Erazo, 2009) describe cómo reducir el número de objetos a analizar, en un problema de grandes cantidades de datos.
- El trabajo en (Kuri-Morales, 2018) detalla el reconocimiento de patrones en bases cuyos datos son tanto numéricos como categóricos.
- Los autores en (Rodríguez-González y cols., 2018) explican cómo hallar únicamente patrones interesantes.
- La aplicación de árboles de clasificación en el clima se presenta en (Coria y cols., 2016).
- Descubra cómo hallar prototipos o “ejemplares típicos” rápidamente leyendo a (Olvera-López, Carrasco-Ochoa, y Martínez-Trinidad, 2018).
- Para transformar problemas de regresión en problemas de clasificación, es útil estudiar a (Molina-Félix, Oliveira-Rezende, Monard, y Caulkins, 2000).
- En (Morales-González, García-Reyes, y Sucar, 2018) se explica para qué sirve combinar reconocimiento y segmentación.

Algunas aplicaciones realizadas en colaboración con el Laboratorio de Ciencia de Datos y Tecnología de Software (LCDTS) del Centro de Investigación en Computación (CIC), son:

Cómo el contexto familiar, escolar y personal afecta el rendimiento del estudiante.

Arturo Heredia dice en su tesis (Heredia-Márquez, 2017) (ver Figura 5.1): “Se usa primordialmente la Minería de Datos y la Visualización. Un caso con un alto número de dimensiones es el Examen de Calidad y Logro Educativo (EXCALE), en 3° de Secundaria en Matemáticas, que realizaba el Instituto Nacional para la Evaluación de la Educación (INEE). Lo acompaña un cuestionario de contexto que mide las condiciones de los estudiantes en tres entornos: 1) entorno personal, 2) entorno familiar, y 3) entorno escolar. Se analizaron los casos de cerca de 55,000 estudiantes que presentaron el examen. En algunos casos son más de 100 variables entre evaluación y contexto. Se observaron patrones y tendencias en variables de interés como la modalidad, sus aspiraciones académicas e historial académico (Promedio General, Promedio anterior en Matemáticas), la atención y frecuencia con la que los estudiantes entienden al profesor y que encuentran correlación con el desempeño académico de los estudiantes. Con la reducción de dimensiones se logró aumentar las clasificaciones correctas de los algoritmos utilizados, se disminuyeron los tiempos de generación de los modelos de clasificación y del tamaño del árbol, así como del espacio de análisis.”

Indicadores y tendencias delictivas en la ciudad de México. Trayectorias seguras en zonas peligrosas.

De la tesis de Manuel Gutierrez (Gutiérrez-Ceballos, 2017) (ver Figura 5.2) se extrae: “En este proyecto se propone una plataforma de software para el análisis en su parte de descripción de fenómenos como el indicado. Este fenómeno es de interés para México, ya que cuenta con

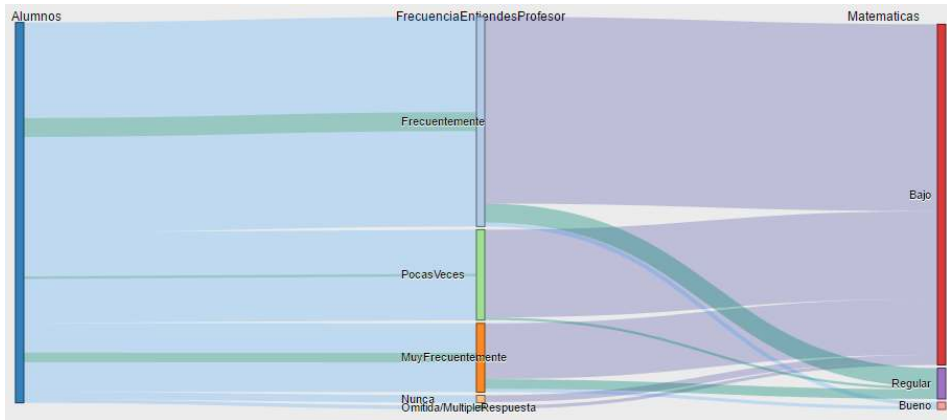


Figura 5.1: Ejemplo del diagrama Sankey con la lectura de derecha a izquierda. Esta figura muestra la relación de los estudiantes (alumnos) con un desempeño regular en Matemáticas, con sus respuestas para la pregunta “Frecuencia Entiendes Profesor”, identificando que la mayoría responde “Frecuentemente” y “Muy Frecuentemente”. La columna de la derecha “Matemáticas, bajo, regular, bueno” indica visualmente los resultados obtenidos por los alumnos.

10 ciudades en el Ranking de las 50 ciudades más peligrosas del mundo, utilizando los delitos de alto impacto de la Ciudad de México en el lapso de tiempo del 2013 al 2016. El análisis de descripción se realiza a diferentes niveles:

- El análisis geográfico en diferentes regiones de la ciudad (colonias, delegaciones, cuadrantes y no convencionales)
- El mapeo de información geo-referenciada de delitos con otros puntos de interés
- Así como funcionalidades que puedan aportar a un análisis temporal.

Para las funcionalidades se destacan como aportaciones, varios algoritmos propios como:

- El que localiza las rutas entre dos diferentes puntos
- y el que califica a las regiones en diferentes niveles de peligrosidad.

Además de las estructuras de almacenamiento que permiten responder las funcionalidades desarrolladas (base de datos con información delictiva, base de datos de polígonos de las regiones de análisis y de puntos de interés, cubos de datos y matrices de adyacencia).”

Ecobicis. Mejor redistribución de las bicicletas públicas en la ciudad de México por los camiones repartidores.

Carlos Esquivel habla de la demanda y el reacomodo de bicicletas (ver Figura 5.3). Reza su tesis (Esquivel-Briceño, 2018): “Como primera tarea se realizó la construcción de una base de datos que contiene la información de los viajes realizados en ECOBICI desde el mes de febrero de 2010, con una cantidad de más de cuarenta y cinco millones de viajes realizados hasta la actualidad, datos del clima por hora en cada colonia de la Ciudad de México, datos del estado de cada Cicloestación con la información de cuantas bicicletas están disponibles minuto a minuto. Con la información anterior fue posible plantear una serie de alternativas de solución basada en dos técnicas de aprendizaje automático para solucionar el problema del cálculo del inventario necesario por estación; estas técnicas fueron un Modelo Autorregresivo de Medias Múltiples y un Bosque Aleatorio, en el primer método se aprovecharon los datos históricos de los viajes para tratarlos como una serie de tiempo; en cuanto al bosque aleatorio, contempla la información histórica de los viajes, e incluye dentro de su vector de características información «online» y «offline» del clima de la ciudad de México, con la intención de robustecer el modelo.

Finalmente, el autor de este trabajo plantea un método propio para realizar la estimación de inventarios necesarios, pero a diferencia de los dos anteriores, que están basados en el historial de viajes realizados, el nuevo método aprovecha la información histórica del estado de las Cicloestaciones. Finalmente se plantea un algoritmo de fuerza bruta para el establecimiento de las rutas de rebalanceo basado en la construcción de un grafo de reorden que incorpora la información obtenida por la predicción del inventario necesario por estación; este algoritmo de fuerza bruta pudiera parecer intratable en un sistema con muchos nodos y aristas, sin embargo también se plantean heurísticas para realizar podas espaciales y temporales en el grafo de reorden, con la finalidad de realizar tales podas aprovechando la información a priori con la que se cuenta del sistema. Esta información fue conocida gracias al apoyo del Sistema de Análisis de Movilidad (SAM), desarrollado en el LCDTS del CIC por un equipo de estudiantes y profesores dentro de los cuales se encuentra el autor de este trabajo.”

Visualizando muchas propiedades (datos multivariados) de modo entenable. Dice Rodolfo Vilchis en el resumen de su tesis (Vilchis-Mompala, 2013): “La visualización de la información es una técnica muy usada para analizar las relaciones entre las variables de un conjunto de datos. Éstas pueden ser tanto numéricas como simbólicas y al generar una visualización la mayoría de las veces se muestran tres variables, y seis como máximo, usando color, forma y tamaño. Después de todo, el ojo humano puede percibir con claridad datos en papel o en pantalla con dos dimensiones (ejes), máximo tres. Sin embargo, en conjuntos multivariados es común tener más de cinco dimensiones, por lo cual, al visualizarlos, el usuario no detecta cómo varían las variables a través de los datos, ni las relaciones entre éstas. Para resolver este problema, se suelen usar varias gráficas y tablas. Esto da una visión fragmentada de cuáles objetos tienen qué valores

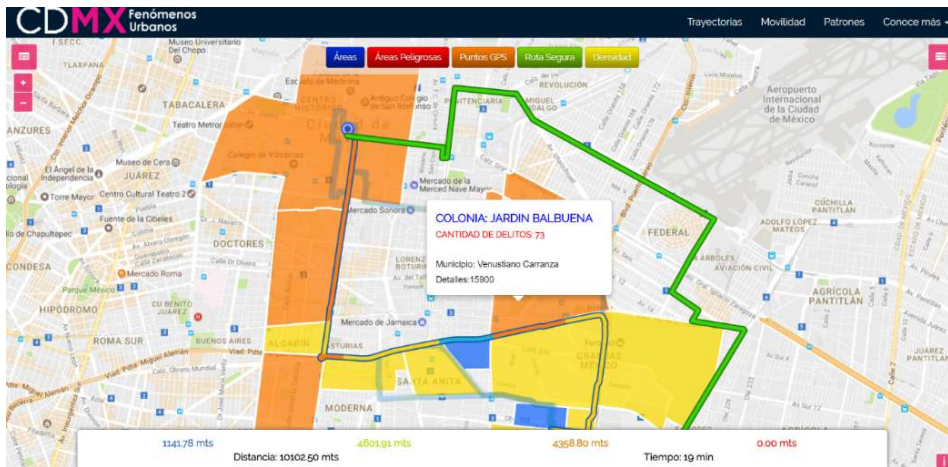


Figura 5.2: Trayectoria segura entre un origen y un destino, según el medio de transporte empleado. El software provee rutas entre dos puntos, como lo hace Waze o Google Maps, según el medio de transporte. Pero toma en cuenta la peligrosidad de la ruta (y la altera para ir por zonas menos conflictivas) según el tipo de delito, hora del día, entre otras.

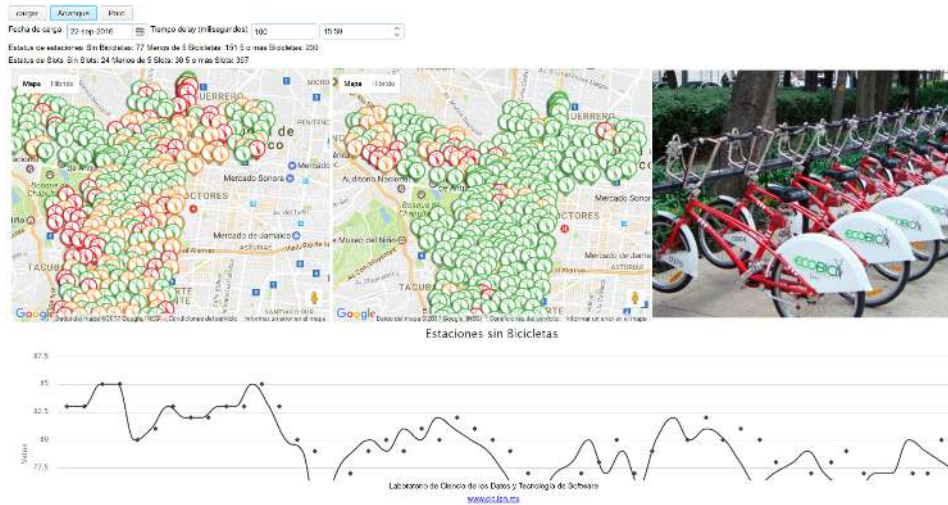


Figura 5.3: Estaciones con bicicletas (verdes), con pocas (amarillas) y sin bicicletas (rojas), en dos tiempos distintos, en la ciudad de México. La gráfica inferior muestra la variación de la demanda a corto plazo.

en cuáles atributos. El propósito de este trabajo es transmitir la mayor cantidad de información posible presente en un conjunto de datos de tal forma que sea fácilmente comprensible por el ser humano, es decir, agilizar la detección de las relaciones entre las variables numéricas y simbólicas. El trabajo presenta un nuevo método para mostrar en una sola gráfica tantas variables como sea posible, de modo que el usuario tiene una visión más integral de los datos. El sistema desarrollado, automáticamente escoge el mayor número posible de variables a mostrar (dado unos parámetros) y las agrupa para que la comprensión se efectúe sobre un mayor número de variables. Para ello, se hace uso de la regresión lineal, utilizando dos métodos, los mínimos cuadrados y *Multivariate Adaptive Regression Splines* (MARS). La idea es encontrar comportamientos monótonos (crecientes y decrecientes) entre las variables, para poder graficarlas en un mismo

eje cartesiano, cada variable con una escala diferente. Si hay variables constantes o casi constantes (varían muy poco) éstas se muestran en la visualización con una etiqueta. Aquellas variables que no poseen un comportamiento monótono con otras se tratan de ajustar mediante un particionado (reduciendo la precisión) sobre alguno de los ejes cartesianos.

Para las variables simbólicas, se busca una partición de dos o tres conjuntos de tal forma que encajen (se particionen) sobre algún eje, lo que permitirá graficarlas. Si existen variables simbólicas sobrantes, es decir, no se ajustaron mediante un particionado, se seleccionan dos de ellas y se muestran mediante el color y forma, siempre y cuando cumplan con algunas restricciones.

Con las técnicas empleadas, un conjunto de 3194 registros con 52 variables fue posible mostrarlo con nueve de sus 52 variables, otro conjunto de 4898 registros y 12 atributos se mostró con ocho de sus 12 atributos y otros conjuntos han mostrado buena visualización. En general, ambos métodos dan buenos resultados; bajo ciertas condiciones es mejor usar mínimos cuadrados y en otras MARS. Para las variables simbólicas en algunos casos se logró encontrar una partición dando buenos resultados en la visualización.”

Generación del cubo de datos en una GPU. Paralelización de consultas SQL en la GPU. El resumen de la tesis de Mario Torres (Torres-Rivera, 2013) explica las ventajas del uso de la unidad de procesamiento gráfica (GPU): “La generación de cubos de datos de manera eficiente es un problema central en los almacenes de datos y el procesamiento analítico en línea. Es un proceso que puede implicar la ejecución de gran número de operaciones aritméticas, además de consumir bastante tiempo cuando se realiza a partir de datos de gran volumen. Para una relación R con n atributos o dimensiones más un atributo de medida M , $R(A_1; A_2; \dots; A_n; M)$, el problema básico del cálculo del cubo de datos

implica la agregación de R para construir 2^n grupos de tuplas respecto a toda posible combinación de las n dimensiones (i.e., el conjunto potencia de las n dimensiones de R). A cada uno de estos grupos de tuplas se le llama cuboide. Dicho problema ha sido investigado y se han propuesto estrategias para resolverlo, sin embargo, hasta ahora la mayoría de los algoritmos no consideran las ventajas del paralelismo y las recientes arquitecturas de CPUs y GPUs.

En este trabajo se presenta el diseño de un conjunto de operaciones paralelas llamadas primitivas que aprovechan el paralelismo proporcionado por los modelos recientes de GPUs y CPUs multi-núcleo. Las primitivas facilitan la generación de cubos de datos llevando a cabo rutinas de ordenamiento, partición y agregación. La implementación del software para GPU de este trabajo se realizó mediante la plataforma de cómputo en paralelo conocida como CUDA del fabricante de procesadores gráficos NVIDIA, y para implementar el paralelismo en procesadores multi-núcleo se utilizaron hilos POSIX.

Posteriormente, se introducen tres métodos paralelos para generación de cubos de datos completos y de tipo iceberg. Además de las primitivas previamente diseñadas, estos métodos utilizan hilos POSIX con el fin de explotar el paralelismo de CPUs multi-núcleo en la construcción simultánea de varios cuboides, i.e., todos los cuboides distribuyen en grupos y posteriormente los cuboides de cada grupo se generan en paralelo. Se utiliza almacenamiento en memoria lineal a través de arreglos de una dimensión para almacenar tuplas en memoria principal, evitando costos relacionados con la construcción de estructuras de datos más complejas. Así mismo, se utilizan algunas estrategias conocidas en la literatura a fin de agilizar la generación del cubo de datos. Sin embargo, a diferencia de los trabajos previos, los métodos presentados en esta tesis constan de un paralelismo de grano fino que se obtiene a través del uso de las primitivas paralelas.”

Identificación y seguimiento a personas con una red de cámaras. Sistema de segmentación y seguimiento de objetos en movimiento.

Claudia Tavera describe en el resumen de su tesis (Tavera-Olivares, 2017): “Este trabajo desarrolla un algoritmo para la detección y seguimiento de una persona por medio de una cámara fija utilizando técnicas de procesamiento de imágenes como la substracción de imágenes, cambio de imagen de color a tonos de grises, umbralado automático de imagen e identificación de región de encuadre. El sistema consta de un módulo de detección de objetos en movimiento y un módulo de seguimiento de un objeto en movimiento con base al color.”

¿De qué hablan las noticias en México?.

El resumen de la tesis de Víctor Zaragoza (Zaragoza-Luna, 2016) señala: “El trabajo presente describe un modelo que permite hallar la configuración de cúmulos naturales ‘al sentido humano’, con mayor relación semántica en una clasificación no supervisada; es decir, convertir un clasificador no supervisado, como lo es LDA («*Latent Dirichlet Allocation*», en inglés) en un clasificador que clasifica un conjunto de documentos en cierto número de grupos, que concuerdan bastante bien con la clasificación que una persona normalmente haría.

Los documentos (información no estructurada) empleados son noticias recolectadas de los sitios web de ciertas instituciones (prensa digital mexicana); las noticias reciben tratamiento de lenguaje natural para hacer coherente el proceder del algoritmo de clasificación no supervisado, además de proveer un formato apropiado para la etapa del cálculo de similitud entre palabras; el tratamiento consistió en: remover «*stop-words*»; lematizar y aplicar sinónimos.

Los tópicos principales inmersos en los documentos se encuentran empleando LDA, que es un algoritmo de clasificación no supervisado y que es capaz de

identificar cuáles son las palabras que representan a cada uno de los k tópicos que se soliciten.

La aportación del trabajo es la propuesta de calcular las distancias que existen entre las palabras de un grupo o cúmulo (intra-distancias) y las distancias entre cúmulos (inter-distancias) para encontrar su compacidad y lejanía natural; con ello se puede identificar el número de cúmulos naturales que se hallan en un conjunto de datos. Para el cálculo de la distancia se emplea una jerarquía o taxonomía que es capaz de trabajar con las relaciones semánticas de las palabras; la taxonomía es *Wordnet*. La compacidad de los cúmulos y la distancia natural entre los mismos se calcula empleando una función de similitud llamada *Path*. También, se propone una forma de etiquetar los cúmulos con múltiples palabras empleando nuevamente la relación semántica de los términos. Para apreciar las propiedades del agrupamiento de documentos se desarrolló una visualización que muestra los resultados de la investigación; por ejemplo, la etiquetación, los cúmulos, la cantidad de documentos asignados a cada cúmulo, y relaciones entre cúmulos.” Ver las Figuras 5.4 y 5.5.

Búsqueda de patrones en cadenas de ADN. Este trabajo (Ortíz-Chan, 2016) analiza vastos volúmenes de texto (secuencias de ADN) para hallar estructuras llamadas *motifs*. El resumen de su tesis reza “Uno de los problemas clásicos en la bioinformática es la búsqueda de patrones frecuentes en secuencias de ADN.

En este trabajo se desarrolla un algoritmo alternativo llamado KTreeMotif para la búsqueda de *motifs* en secuencias de ADN, con rendimiento similar a los algoritmos más usados hoy en día como son el Gibbs Sampler, el Motif Sampler, el MEME y el SP-Star. KTreeMotif busca aprovechar las metodologías y cualidades de un conjunto de estos algoritmos con la finalidad de rearmar los resultados obtenidos por otros medios. KTreeMotif, además, implementa una

nueva estructura de datos para almacenar y recorrer las sub-cadenas de una manera más sistemática y rápida que el recorrido secuencial que suele usarse; también se hace una simplificación de la función de distancia entre una PWM y una secuencia sin perder información logrando el mismo resultado, y finalmente una mejora en exactitud de la función de distancia entre dos secuencias para el caso específico de la búsqueda de patrones frecuentes. Todas estas novedades son integrables a los otros algoritmos que atacan este problema.”

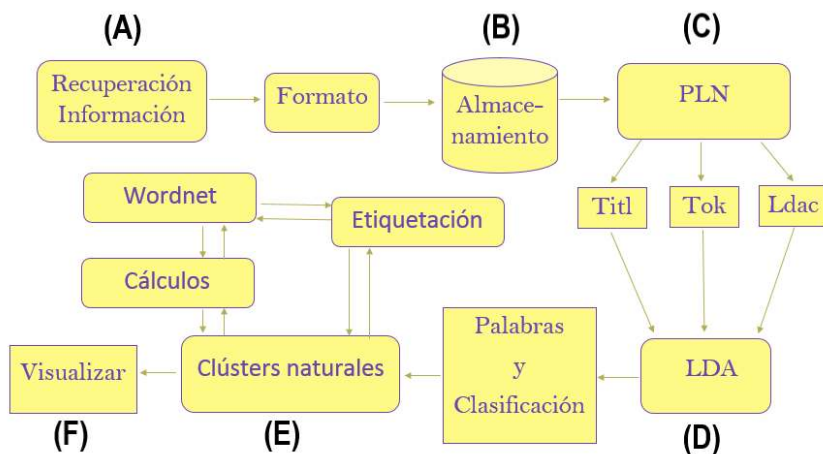


Figura 5.4: Diagrama general del proyecto: ‘Visualización de tópicos importantes usando un enfoque semántico’. Los documentos recuperados en (A) se guardan en (B), se procesan bajo tres técnicas en (C) y se someten al clasificador no supervisado (D). Luego viene la etapa de hallar los cúmulos naturales (E) usando *Wordnet*, para finalmente visualizarlos (F).

Detección de objetos variables (candidatos a supernovas) en imágenes astronómicas, a través del tiempo. Ana Cruz analizó en su tesis (Cruz-Martínez, 2016) las imágenes nocturnas del firmamento del hemisferio sur, to-

madras desde un telescopio en Chile durante dos años. Mediante análisis de señales, supresión de ruido, filtrado, normalización, y clasificación supervisada, logró la detección de objetos cuya luz varía en el tiempo. Concretamente, identificó diferentes objetos que son supernovas de clase I y II (formalmente, son “candidatos a supernova”, ya que no han sido verificados aún en la forma rutinaria).

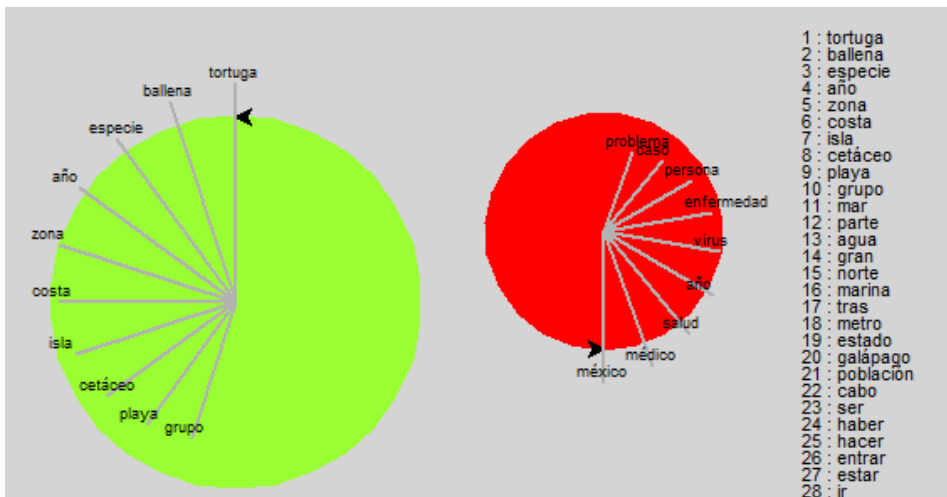


Figura 5.5: Visualización de las agrupaciones. Se muestra la salida del módulo de visualización (G) de la Figura 5.4. Se le ha pedido al sistema que agrupe las noticias en cúmulos («clusters») naturales, y que exhiba dos de ellos.

5.4. El futuro de la minería de datos

Cuando hay una acumulación histórica, ordenada, de datos en forma estructurada o semi-estructurada (texto, documentos), es apetecible usar minería de

datos para extraer de ellos información útil. Si una organización desea usar minería de datos de manera continua, es útil construir primero en una *bodega de datos* (ver sección 5.2.2).

En México, la minería de datos está empezando, y cobra cada vez más fuerza. Una de las razones es la existencia de fuentes de datos abundantes en los negocios (internet, correos, comercio electrónico, transacciones, documentos), la ciencia (percepción remota, bioinformática, simulación científica), el gobierno (datos abiertos, transparencia, INEGI), los dispositivos (cámaras digitales, cámaras en vías públicas y en empresas, sensores, internet de las cosas), la sociedad (noticias, discursos, videos, Twitter, YouTube). Además, en la actualidad, los datos son sencillos de obtener, transmitir, almacenar y procesar. Por otra parte, existe la tecnología, como por ejemplo el reconocimiento de patrones, el análisis de señales, la inteligencia artificial, la estadística, la ciencia de datos y la estructura de la información en redes. Una razón más es la necesidad existente debido a un mundo globalizado, mercados abiertos, una competencia feroz, uso generalizado de dispositivos móviles; ¿Qué opinan sobre lo que hago, digo o vendo? Se disponen de los datos, pero ¿qué significan? ¿Cómo debo conducir mi negocio el próximo año?.

5.4.1. Aplicaciones en el futuro cercano

Empezamos a ver, y cobrarán fuerza en un futuro cercano (5 años) aplicaciones en las siguientes áreas:

- Ciberseguridad. Patrones anómalos en el comportamiento de mi máquina; de mi red.

- Asistencia médica. Conocimiento de las dolencias de cada individuo y de las enfermedades y problemas más comunes de la región. Acciones preventivas; epidemias; seguros colectivos.
- Internet de las cosas. Datos de muchos sensores para que la computadora supervise o controle la fábrica, la irrigación de cultivos, el transporte, el tráfico citadino y en autopistas, producción eléctrica, iluminación, calefacción.
- Compromiso y experiencia del cliente. Dónde compra, qué compra, cuándo, qué ve, qué anuncios le llaman la atención, ofertas instantáneas (vía su celular), cupones, publicidad dirigida. Nuevas formas de atraer clientes, de retenerlos. Cupones instantáneos.
- Todo inteligente. Ventas, universidades, deportes, viajes, que se adaptan conforme cambian sus usuarios (o clientes).
- Capital humano. Conocimiento de las preferencias, fobias, aptitudes y actuación de cada empleado, estudiante, cliente, miembro de mi distrito electoral, o simpatizante, contra las necesidades, problemas y ofertas más comunes de mi empresa, organización o partido político. Publicidad dirigida. Campañas electorales individualizadas (a cada quien le digo lo que quiere oír).
- Sociedad y ecosistema. Mejor entendimiento y predicción del clima, cosechas, epidemias, hambrunas, descontento social, migración, tendencias de precios, de la moda. Uso de grandes volúmenes de datos para simular leyes o reglamentos antes de ponerlas en vigor.

5.5. Apéndice: Recursos recomendados

Algunas fuentes de consulta acerca de minería de datos:

- Han, Jiawei & Kamber, Micheline. (2000). *Data Mining: Concepts and Techniques*. Third edition. ISBN 1-55860-489-8. El libro de Han & Kamber es una buena referencia para profundizar en la minería de datos. Es un libro teórico, con un capítulo (esencial) de limpieza y generalización de datos, con ejemplos prácticos y con lecturas de artículos recientes, que profundizan en las técnicas de cómputo usadas. Varios de los conceptos de este capítulo han sido tomados, con cambios, de él. También es útil consultar la página del autor: <http://hanj.cs.illinois.edu>.
- Jure Leskovec, Anand Rajaraman, Jeffrey D. Ullman. *Mining of massive data sets*. Stanford University, <http://infolab.stanford.edu/~ullman/mmds/book.pdf>. Libro apropiado para manejar grandes cantidades de datos (ver sección 5.1.2).
- <https://www.kdnuggets.com/>. La página de kdnuggets contiene mucha información sobre Ciencia de Datos, tendencias, cursos, entrevistas, herramientas, etc.

Algunas herramientas de uso libre para minería de datos:

- Para una sola computadora: WEKA (de amplio uso, www.cs.waikato.ac.nz/ml/weka), KNIME (de amplio uso, www.knime.com), TALEND (para integración de datos, es.talend.com), EXCEL, TABLEAU (versión gratuita, www.tableau.com), PENTAHO (www.pentaho.com/marketplace). También son útiles Python (Lenguaje de programación con librerías para minería, www.python.org), R (lenguaje de

programación para cómputo estadístico, www.r-project.org), RapidMiner (rapidminer.com), SQL, Anaconda (lenguaje de programación para ciencia de datos, anaconda.org/anaconda/python, www.anaconda.com/distribution/), Tensorflow (librerías para aprendizaje automático, www.tensorflow.org), Scikit-learn (aprendizaje automático en Python, scikit-learn.org/stable/), Keras (aprendizaje profundo en Python, keras.io), Apache Spark (plataforma para análisis de muchos datos, spark.apache.org)

- Para un arreglo de computadoras. Para análisis de grandes conjuntos de datos, que no caben en un disco o que no se procesan convenientemente por un solo procesador. Hadoop y MapReduce (instalada en el LBDTS del CIC, hadoop.apache.org), Flink (instalada en el mismo laboratorio, para cálculos sobre torrentes de datos, ver sección “Uso de la minería de datos”), flink.apache.org), Spark (spark.apache.org).
- Para limpiar datos: OpenRefine (también sirve para explorar y transformar datos, openrefine.org/), Trifacta Wrangler (www.trifacta.com/products/wrangler/), DataKleenr (chi2innovations.com/datakleenr/), Google Refine (ahora se llama OpenRefine, q.v.).
- Para visualizar datos: Tableau (versión gratuita, www.tableau.com), Infogram (para crear infografías, gráficos, mapas, paneles infogram.com/es), cartoDB (carto.com), Piktochart (piktochart.com), Socrata (socrata.com), PowerBI (powerbi.microsoft.com/es-es/), Google fusion tables (support.google.com/fusiontables/answer/2527132?hl=en). Otras menos comunes: Gephi (gephi.org), Axiis (www.axiis.org).

Herramientas comerciales (con costo):

- TABLEAU (versión comercial, www.tableau.com), SAS (estadística, http://www.sas.com/en_us/software/stat.html), SPSS (estadística, www.ibm.com/analytics/mx/es/technology/spss/index.html), Endeca (ORACLE, www.oracle.com/technetwork/middleware/endeca/overview/index.html).

Sitio dedicado exclusivamente a *Latent Semantic Analysis* (LSA):

- <https://lsa.colorado.edu/>. LSA es una técnica matemática y estadística para extraer y representar la similitud de significado de palabras y frases, mediante el análisis de grandes conjuntos de texto. Usa la descomposición en valores singulares (*eigenvalues*), una forma general de análisis factorial, para condensar una matriz muy grande de palabras-en-contexto a otra mucho más pequeña, pero aún grande, del rango de 100 a 500 dimensiones. El número correcto de dimensiones es crucial.

Referencias

- Coria, S. R., Gay-García, C., Villers-Ruiz, L., Guzmán-Arenas, A., Sánchez-Meneses, O., Ávila-Barrón, O. R., ... Martínez-Luna, G. L. (2016). Climate patterns of political division units obtained using automatic classification trees. *Atmósfera*, 29, 359 - 377.
- Cruz-Martínez, A. B. (2016). *Detección de objetos variables en imágenes astronómicas a través del tiempo* (Tesis de Maestría). CIC-IPN.
- Esquivel-Briceño, C. R. (2018). *Modelo para la predicción de la demanda y el rebalanceo dinámico del sistema de bicicletas públicas de la ciudad de México* (Tesis de Maestría). CIC-IPN.

- Gutiérrez-Ceballos, M. (2017). *Análisis regional de delitos de alto impacto en la ciudad de México con mapeo de puntos de interés* (Tesis de Maestría). CIC-IPN.
- Han, J., Kamber, M., y Pei, J. (2011). *Data mining: Concepts and techniques* (3rd ed.). San Francisco, CA, USA: Morgan Kaufmann Publishers Inc.
- Heredia-Márquez, A. (2017). *Técnicas de visualización para asociar condiciones de estudio con los resultados en matemáticas de la prueba EXCALE* (Tesis de Maestría). CIC-IPN.
- Kuri-Morales, A. (2018). Pattern discovery in mixed data bases. En *Mexican Conference on Pattern Recognition* (pp. 178–188).
- Kuri-Morales, A., y Rodríguez-Erazo, F. (2009). A search space reduction methodology for data mining in large databases. *Engineering Applications of Artificial Intelligence*, 22(1), 57 - 65.
- Molina-Félix, L. C., Oliveira-Rezende, S., Monard, M. C., y Caulkins, C. W. (2000). Transforming a regression problem into a classification problem using hybrid discretization. *Computación y Sistemas*.
- Morales-González, A., García-Reyes, E. B., y Sucar, L. E. (2018). Image annotation by a hierarchical and iterative combination of recognition and segmentation. *IJPRAI*, 32(1), 1–31.
- Olvera-López, J., Carrasco-Ochoa, J., y Martínez-Trinidad, J. (2018, 1 de 1). Accurate and fast prototype selection based on the notion of relevant and border prototypes. *Journal of Intelligent and Fuzzy Systems*, 2923–2934.
- Ortíz-Chan, L. A. (2016). *Búsqueda de patrones en cadenas de ADN* (Tesis de Maestría). CIC-IPN.
- Rodríguez-González, A. Y., Lezama, F., Iglesias-Alvarez, C. A., Martínez-Trinidad, J. F., Carrasco-Ochoa, J. A., y de Cote, E. M. (2018). Closed

frequent similar pattern mining: Reducing the number of frequent similar patterns without information loss. *Expert Systems with Applications*, 96, 271 - 283.

Tavera-Olivares, C. (2017). *Sistema de segmentación y seguimiento de objetos en movimiento* (Tesis de Maestría). CIC-IPN.

Torres-Rivera, M. A. (2013). *Generación del cubo de datos empleando paralelismo de GPU'S y CPU'S multinúcleo* (Tesis de Maestría). CIC-IPN.

Vilchis-Mompala, R. A. (2013). *Visualización de objetos multivariados utilizando agrupación de variables* (Tesis de Maestría). CIC-IPN.

Zaragoza-Luna, V. U. (2016). *Modelo de detección de cúmulos naturales basado en una taxonomía semántica* (Tesis de Maestría). CIC-IPN.

Capítulo 6

Administración de la Energía

PABLO H. IBARGÜENGOYTIA, INEEL

GUILLERMO SANTAMARÍA BONFIL, INEEL

ALBERTO REYES, INEEL

La energía es uno de los campos del quehacer y del conocimiento humano más importantes del siglo pasado y más fuertemente en el presente. Entre los temas importantes a considerar en el sector energético se pueden mencionar los siguientes:

- Generación de la energía, tradicionalmente se lleva a cabo en grandes centrales termoeléctricas, hidroeléctricas o nucleares. últimamente se están considerando fuentes renovables de energía como eólica y fotovoltaica.
- Transmisión, se refiere al envío de energía desde las centrales generadoras a las subestaciones en las zonas urbanas.
- Distribución, se refiere al envío de energía desde las subestaciones urbanas a los consumidores como casas habitación, fábricas u hospitales.

- Sustentabilidad, esto es, la tendencia actual a hacer un uso eficiente de la energía así como disminuir los efectos del calentamiento global del planeta.

Para cada uno de los temas importantes del sector energía, se pueden identificar las siguientes funciones o problemas que se pueden resolver con técnicas inteligentes:

- Monitoreo y diagnóstico, identificar el estado de funcionamiento de un sistema en base a las señales que provea.
- Planificación, seleccionar un conjunto de acciones a seguir para conseguir una meta específica.
- Optimización, seleccionar una solución a un problema entre un conjunto de posibles otras soluciones con uno o varios objetivos y restricciones a cumplirse.
- Toma de decisiones, selección de una acción generalmente en presencia de incertidumbre, dado el estado de un sistema y las metas perseguidas.
- Pronóstico, proponer un posible estado futuro de algún parámetro dada la historia de ese y otros parámetros relacionados.

En estos temas importantes del sector de energía y tratando de resolver los tipos de problemas, la computación y en especial la Inteligencia Artificial (IA) han jugado un papel importante (Dabbaghchi, Christie, Rosenwald, y Liu, 1997). El Análisis de señales y el reconocimiento de patrones (AR/RP) representan técnicas de IA muy utilizadas en aplicaciones reales.

6.1. Diagnóstico de transformadores con señales de vibración

Un claro ejemplo de un proyecto nacional utilizando AS/RP fue el diagnóstico de transformadores de potencia utilizando señales de vibración (P. H. Ibargüen-goytia, nan, y Betancourt, 2010). Los transformadores de potencia son equipos eléctricos sin partes móviles usados para transferir energía de un circuito a otro a través de un campo magnético y sin conexión eléctrica directa entre los dos circuitos. La idea del proyecto fue aprender un modelo con las vibraciones de transformadores cuando operan normal y adecuadamente. Las fallas que identifica el sistema de diagnóstico son mecánicas. A diferencia de otras pruebas como análisis de gases en el aceite o la detección de descargas parciales, las vibraciones identifican fallas mecánicas como aflojamiento de devanados o del núcleo. El proyecto fue parcialmente financiado por la compañía Prolec-General Electric que además aportó transformadores en su última etapa de pruebas de fabricación (ver Fig. 6.1) para las pruebas de modelos de vibración.

Los sensores de vibración producen mediciones de aceleración en el dominio del tiempo. Esas mediciones no producen mucha información por lo que se procedió al análisis de las señales en el dominio de la frecuencia. Utilizando la transformada discreta de Fourier (DFT por sus siglas en inglés) se obtuvieron las componentes de la señal en las diferentes frecuencias que componen las señales. Por supuesto que las frecuencias donde se encontraron esas componentes son múltiplos de 60 Hertz, la frecuencia de la corriente alterna en el continente americano.

El reconocimiento de patrones de vibraciones anormales se logra a través de la detección de desviaciones al comportamiento normal de vibración. Este



Figura 6.1: Parte interna de un transformador. Se muestra el ensamblado del núcleo y los devanados. (Cortesía de Prolec GE ®).

comportamiento normal deberá ser modelado en todo el rango de condiciones de operación. Las condiciones de operación incluyen las temperaturas y los porcentajes de corriente y voltaje aplicados al transformador. Para conocer las condiciones normales de operación se realizaron experimentos en la fábrica de transformadores de Prolec General Electric en Monterrey, México. La figura 6.2 muestra las pruebas que se hicieron a un transformador operando normalmente. Nótese los sensores de vibración pegados a las paredes del tanque del transformador.

La tabla 6.1 describe el conjunto de experimentos realizados en la fábrica.

El procedimiento de adquisición de datos para aprender los modelos de vibraciones normales consistió en lo siguiente:

- Se colocaron sensores en las paredes del transformador como muestra la Fig. 6.2.

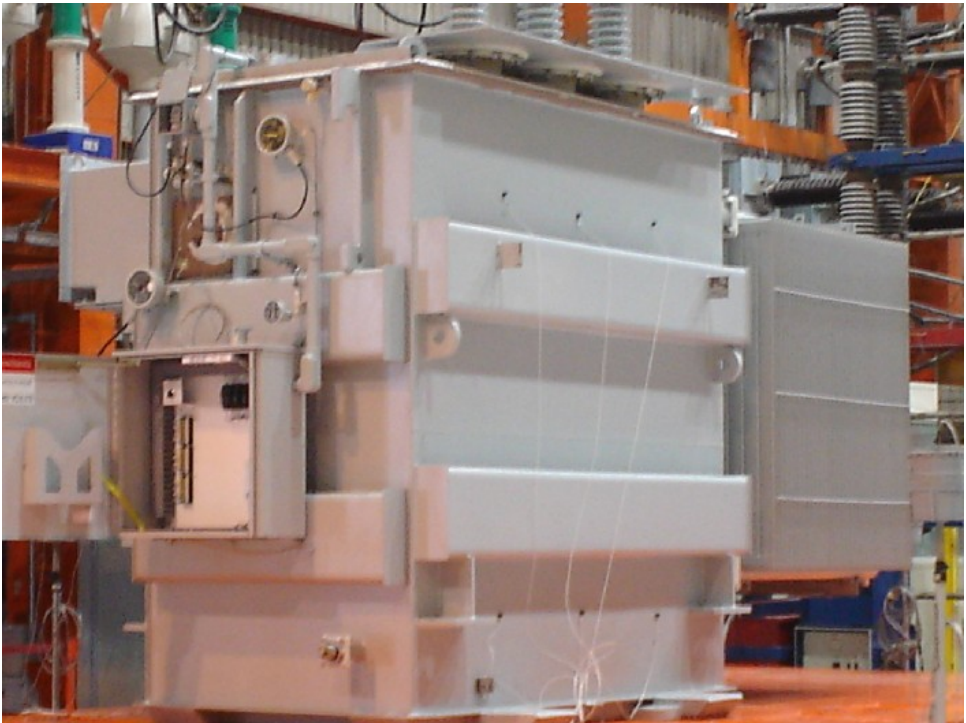


Figura 6.2: Sensores de vibración colocados en un transformador operando correctamente.

Tabla 6.1: Experimentos realizados en fábrica.

	Temperatura		
Excitación	Frío	Caliente	Condición
Voltaje	Efecto del voltaje en vibraciones del núcleo	Efecto del voltaje y temperatura en vibraciones del núcleo	70 %, 80 %, 90 %, 100 %, 110 %
Corriente	Efecto de corriente en vibraciones del paquete de devanado	Efecto de corriente y temperatura en vibraciones del paquete de devanado	30 %, 60 %, 100 %, 120 %

- Se puso el transformador a operar en las combinaciones marcadas en la Tabla 6.1.
- Para cada condición de operación, se adquieren las señales de vibración a una velocidad de 5 K muestras por segundo durante 2 segundos para los 8 sensores.
- Se repite la adquisición 10 veces de cada experimento.
- Por último, se aplica la DFT para adquirir el contenido de frecuencia de cada señal. Las figuras 6.3 y 6.4 muestran un ejemplo de las gráficas logradas.
- Se adquiere la base de datos para alimentar a los programas de aprendizaje automático de modelos.

La figura 6.3 muestra la amplitud de las señales de vibración en la frecuencia de 120 Hz en todos los sensores durante las diferentes condiciones de excitación de

corriente a diferentes porcentajes de operación nominal. Nótese que los sensores 6 y 7 muestran las mayores intensidades de vibración a estos 120 Hz.

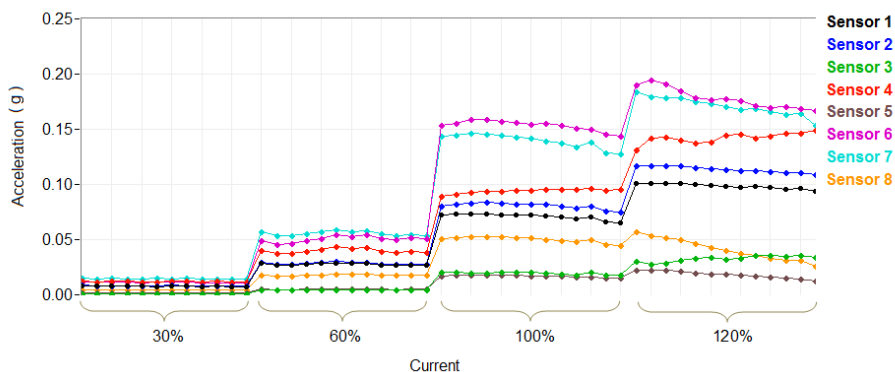


Figura 6.3: Señales de vibración a 120 Hz del transformador excitado con corriente en todos los sensores

La figura 6.4 muestra la amplitud de las señales de vibración captadas por un solo sensor mostrando todas las frecuencias en las mismas condiciones de operación de la figura 6.3. Nótese que las frecuencias de 120 Hz y 240 Hz proveen la información útil. Según los expertos, es de esperarse que los múltiplos pares de la frecuencia fundamental de 60 Hz son las más perceptivos de vibración.

En suma, las variables disponibles para este proyecto son sensores, frecuencias, temperaturas y la excitación del transformador (voltaje o corriente). Siguiendo el consejo de los expertos, consideramos dos posibles conjuntos de modelos. El primero es un modelo que relaciona las condiciones y frecuencias operacionales. Un modelo para cada sensor. El segundo conjunto posible de modelos relaciona condiciones operacionales y sensores. Un modelo para cada frecuencia. Se decidió probar un conjunto de modelos que relaciona condiciones y sensores, es decir, condiciones operacionales y vibraciones detectadas en ciertas

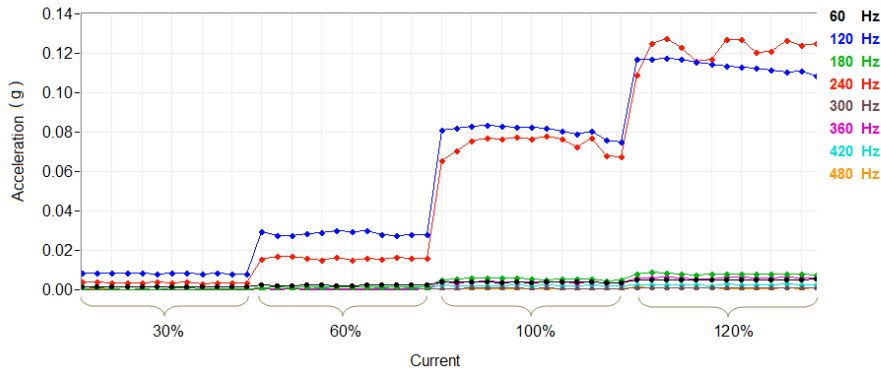


Figura 6.4: Señales de vibración de un sensor del transformador excitado con corriente en todas las frecuencias

partes del transformador. La figura 6.5 muestra una instancia del modelo resultante.

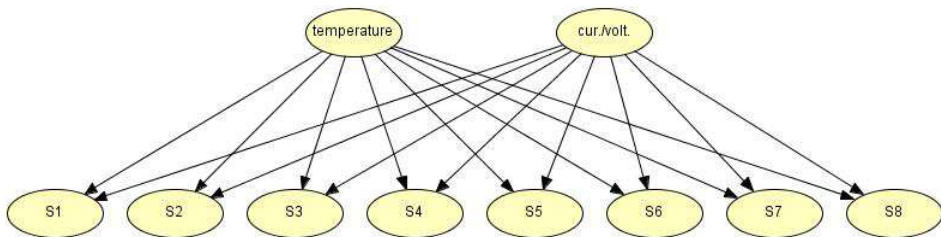


Figura 6.5: Modelo que relaciona condiciones operacionales con la amplitud de vibración medida en cada sensor.

En realidad, el modelo completo está formado por dos Redes Bayesianas (*Bayesian Network*, BN) como la que se muestra en la figura 6.5. Uno correspondiente al componente de 120 Hz y el segundo correspondiente a 240 Hz. Una vez definida la estructura, se utiliza el algoritmo EM (Estimation-Maximization) (Lauritzen, 1995) para obtener las tablas de probabilidad condicionales. Usamos

10 experimentos de cada tipo como se indica en la Tabla 6.1 y aplicado en 5 transformadores. La estructura y los parámetros aprendidos completa los modelos para el diagnóstico.

Se diseñó un sistema computacional que utiliza las mediciones obtenidas en los experimentos descritos en la Tabla 6.1. Se corrieron los experimentos y se identificó si hay una falla. Un experimento consiste en establecer las condiciones operacionales de excitación y temperatura. A continuación, el sistema obtiene las medidas de los sensores y ejecuta el algoritmo para la detección de vibraciones anormales. Como tenemos 16 indicadores de falla (ocho sensores en dos frecuencias), se requiere una técnica de fusión del sensor. Decidimos hacer una suma ponderada de la conclusión de cada sensor. Si un sensor es sensible y detecta desviaciones fácilmente, se le asigna un peso bajo. Otros sensores menos sensibles pueden tener pesos más altos. Dada la suma de cada frecuencia, podemos configurar rangos para declarar un transformador como correcto, sospechoso, defectuoso. La figura 6.6 muestra los resultados de un experimento. En la esquina superior izquierda de la ventana se indican las condiciones de operación. Primero, carga en % con el 100 % de la corriente, y temperatura fría del transformador (frio). En el lado izquierdo de la ventana, hay dos luces. Uno corresponde a un modelo para 120 Hz. y el otro corresponde a 240 Hz. Como se ha mencionado anteriormente, se está utilizando un modelo para cada frecuencia. Estas luces se vuelven verdes si el transformador es correcto, amarillo si el transformador es sospechoso y rojo si hay una falla definitiva. A continuación, en la esquina inferior izquierda de la ventana, hay una pequeña caja para cada sensor en el transformador. Los primeros 8 corresponden a 120Hz y el último correspondiente a 240 Hz. Si se detecta una falla se marca NO-OK y OK en caso contrario. Nótese que el sexto sensor detectado una desviación en el modelo de

240 Hz. En la esquina superior derecha de la ventana aparecen cuatro filas de datos. Las dos primeras corresponden a la amplitud de vibración actual medida en 8 sensores en el transformador. Las siguientes dos filas corresponden a la información normalizada. En realidad, son las entradas a los modelos de redes Bayesianas. La parte inferior derecha de la ventana muestra otra información del prototipo.

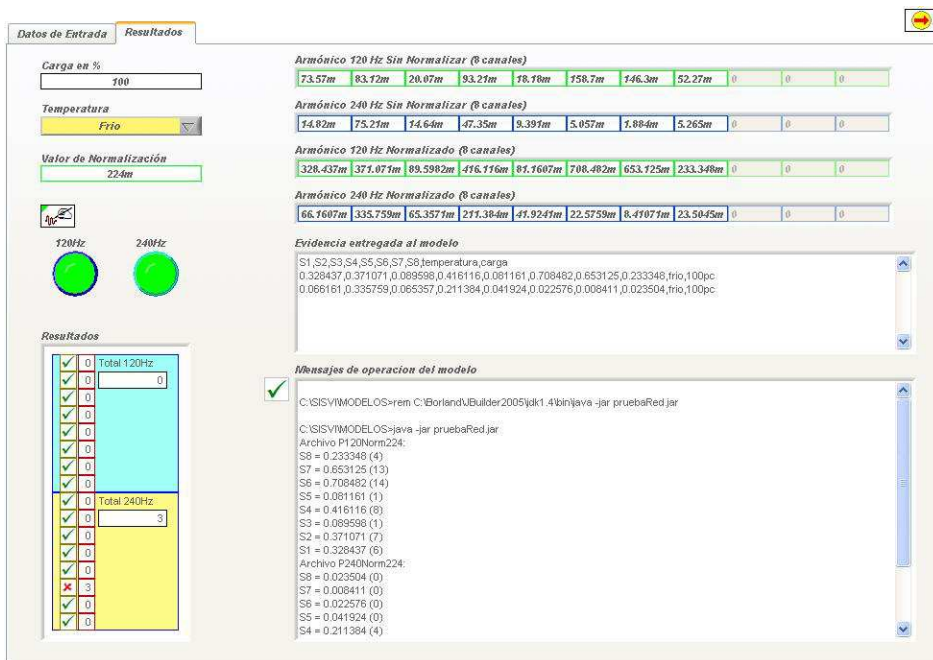


Figura 6.6: Interfaz de usuario del sistema computacional de diagnóstico.

6.2. Ejemplos adicionales de proyectos AS/RP en México

En México el uso de la Inteligencia Artificial (IA) y Reconocimiento de Patrones (RP) ha sido empleado abundantemente en diferentes campos relacionados a la energía y los sistemas de potencia. Entre las aplicaciones de IA/RP para el pronóstico de recursos y generación de energías renovables se encuentran la energía eólica, solar, geotérmica, energía nuclear, además de aplicaciones en los sistemas de potencia.

A continuación se mencionan los diferentes grupos que han trabajado y/o trabajan en dichas áreas. La presentación de los trabajos será por tipo de recurso energético, y aplicaciones a sistemas de potencia. Dentro de cada sección, los trabajos serán agrupados de acuerdo a las técnicas de IA y RP.

6.2.1. Energía Eólica

El pronóstico de la velocidad del viento y de producción de energía eólica es una de las áreas con más aplicaciones de IA y RP en México, desde el 2005 hasta la fecha. El primero consiste en predecir la velocidad del viento a una altura determinada y para un horizonte de pronóstico dado. Mientras que en el segundo se debe predecir la energía generada de acuerdo a las especificaciones de un parque o turbina de viento.

Típicamente los Horizontes de Pronóstico (H_p) empleados para la velocidad del viento y generación de energía son a corto plazo (minutos-horas), mediano plazo (días-semanas), y largo plazo (meses-años). La mayoría de los modelos propuestos se han enfocado en problemas de pronóstico a corto plazo. Por ejemplo, en (Cadenas, Jaramillo, y Rivera, 2010) el $H_p = 10$ minutos, en (Cadenas

y Rivera, 2010) el $H_p = 1$ hora, en (Graff, Peña, y Medina, 2013) $H_p = 1, \dots, 6$ horas, mientras que en (Graff, Peña, Medina, y Escalante, 2014; Santamaría-Bonfil, Reyes-Ballesteros, y Gershenson, 2016) el $H_p = 1, \dots, 24$ horas. Sin embargo, existen algunos trabajos que realizan pronósticos a mediano y largo plazo. Por ejemplo (Flores, Loaeza, Rodríguez, y Cadenas, 2009; Flores, Graff, y Rodríguez, 2012) realiza pronósticos a mediano plazo con un $H_p = 1$ día y de un $H_p = 1$ mes.

Entre las herramientas más utilizadas de IA en esta área se encuentran los Algoritmos Genéticos (GA) (Flores y cols., 2012; Santamaría-Bonfil y cols., 2016; Rodríguez y cols., 2017) y Programación Genética (GP) (Flores, Graff, y Cadenas, 2005; Graff y cols., 2013, 2014). El uso que se le ha dado a GA ha sido el de selección de variables, sintonización de parámetros de modelos de regresión y, en particular, GP como herramienta de pronóstico (Graff y cols., 2013, 2014).

Por otra parte, las herramientas más utilizadas como modelos de regresión son las Redes Neuronales Artificiales (ANN) (Flores y cols., 2009; Cadenas y Rivera, 2010; Cadenas y cols., 2010; Flores y cols., 2012), así como modelos estadísticos clásicos para series de tiempo como modelos Autorregresivos (AR) (Santamaría-Bonfil y cols., 2016), ARMA y ARIMA (Cadenas y Rivera, 2010; Flores y cols., 2012; Graff y cols., 2013, 2014; Santamaría-Bonfil y cols., 2016), además modelos de suavizamiento exponencial (Cadenas y cols., 2010). Otras técnicas de RP menos utilizadas son las Máquinas de Vectores de Soporte (SVM) (Santamaría-Bonfil y cols., 2016). Cabe mencionar, que todas estas propuestas emplean modelos univariantes, es decir, solo emplean una serie de tiempo para realizar tanto el entrenamiento de los modelos como el pronóstico. Adicionalmente, modelos basados en Redes Bayesianas Dinámicas (DBN) (Reyes y cols., 2016; P. Ibar-

güengoytia, Reyes, Borunda, y García, 2018) han sido desarrollados para el pronóstico de la velocidad del viento.

Además de las contribuciones antes mencionadas, existen otro tipo de aplicaciones en este sector. Por ejemplo, en (Rodriguez y cols., 2017) se propone una técnica para la reconstrucción (esto es, imputación de valores faltantes) de series de tiempo del viento empleando GA y ANN. Mientras que en (Sánchez-Pérez, Robles, y Jaramillo, 2016) se desarrolló un método para la caracterización del recurso eólico en regiones geográficas utilizando vectores (i.e. velocidad y dirección del viento) agrupados en *estados de viento* mediante K-medias; una vez agrupados en estados, se utilizaron Cadenas de Markov para determinar las probabilidades de transición entre los estados.

Otras aplicaciones menos relacionadas al pronóstico, son la planeación y diseño de parques eólicos (Herbert-Acero J. and Franco-Acevedo J. and Valenzuela-Rendón M. and Probst-Oleszewski, O., 2009), y el control de turbinas de viento (Ricalde y Cruz, 2010). En el primero, se desarrolló una aplicación para maximizar la producción de energía de un parque eólico hipotético determinando el mejor arreglo de aerogeneradores posible utilizando algoritmos de optimización metaheurística como Recocido Simulado (SA) y GA. En el segundo caso, se desarrolló un modelo de control adaptable basado en ANN para un aerogenerador de pequeña escala con el fin de maximizar la energía obtenida por el mismo.

6.2.2. Energía Solar

En lo referente a energía solar (i.e. fotovoltaica y fototérmica) y su aprovechamiento mediante IA/RP, únicamente se encontraron dos trabajos. En el primero se realiza un modelo de optimización del desempeño térmico de un Con-

centrador Solar Parabólico (CSP) usando ANN y GA (May-Tzuc y cols., 2017). Un CSP es un sistema fototérmico el cual convierten la radiación solar en calor el cual es aprovechado usando un fluido de trabajo. Mediante un ANN se relacionan mediciones de la eficiencia térmica (i.e. y) del CSP con sus características (i.e. x) como el flujo del agua, temperatura ambiente, radiación solar, entre otras. Posteriormente, la optimización de la eficiencia térmica se lleva a cabo optimizando el valor del flujo del agua a través de un GA. El segundo trabajo está relacionado a la predicción de generación de energía fotovoltaica en un parque fotovoltaico al norte de México (Alejos-Moo, Tziu-Dzib, Canto-Esquivel, y Basam, 2018). Más precisamente, este trabajo emplea Redes Neuronales Profundas (DNN) y Redes Neuronales Profundas Recurrentes (R-DNN), aplicadas al pronóstico de generación de energía fotovoltaica. El pronóstico de generación de energía fotovoltaica, al igual que la energía eólica, se divide en horizontes de pronóstico. El trabajo (Alejos-Moo y cols., 2018) pertenece a los pronósticos a corto plazo, con un $HP = 15$ minutos en el futuro.

6.2.3. Energía Geotérmica

La energía geotérmica es el calor contenido en el interior de la Tierra que genera fenómenos geológicos a escala planetaria. Este término es a menudo utilizado para indicar aquella porción del calor de la Tierra que puede o podría ser recuperado y explotado. Sin embargo, realizar perforaciones para explotar estos recursos geotérmicos es altamente costoso, por lo tanto es necesario determinar aquellas zonas con mayor potencial geotérmico. Para ello es necesario llevar a cabo la Estimación de la Temperatura Estática (SFT) de los yacimientos. En este sentido, la mayoría de las contribuciones de IA/RP a esta área están abocadas a la SFT en yacimientos geotérmicos (A. L. Laureano-Cruces y

Espinosa-Paredes, 2005; Olea-Gonzalez, Vázquez-Rodríguez, García-Gutiérrez, y Anguiano-Rojas, 2007; Espinosa-Paredes y cols., 2010; Bassam, Santoyo, Andaverde, Hernández, y Espinoza-Ojeda, 2010; Santoyo y Díaz-González, 2010; Álvarez del Castillo, Santoyo, y García-Valladares, 2012; Díaz-González, Hidalgo-Dávila, Santoyo, y Hermosillo-Valadez, 2013; Pérez-Zárate y cols., 2015).

Para realizar SFT básicamente existen dos clases de modelos: el primero simula la evolución de la temperatura en la columna completa del pozo, mientras que la segunda se enfoca en emplear mediciones de la Temperatura de Fondo del Pozo (BHT). Para el primer caso se han utilizado modelos inversos para la calibración de simuladores físicos del comportamiento de la temperatura en la columna del pozo (A. L. Laureano-Cruces y Espinosa-Paredes, 2005; Olea-Gonzalez y cols., 2007; Espinosa-Paredes y cols., 2010; Bassam y cols., 2010; Álvarez del Castillo y cols., 2012). Para ello, el modelo físico del simulador considera la geometría cilíndrica y flujo radial, transferencia de calor, temperatura del lodo, así como las condiciones iniciales y de frontera. La idea, es determinar los valores de los parámetros del simulador que permitan reproducir valores similares a los medidos. Para ello se han utilizado, agentes computacionales reactivos, modelos cognitivos basados en reglas de expertos humanos modelados con Lógica Difusa (FL), el algoritmo de optimización de Marquardt, y ANN. Para el segundo caso, las herramientas más económicas y usadas para la estimación de la BTH son los geotermómetros químicos, los cuales pueden emplear solutos, gases e isótopos. Los geotermómetros de solutos permiten relacionar la BTH con la composición química de los fluidos geotérmicos que emergen en manantiales hidrotermales, sin la necesidad de perforar. En este sentido, se ha utilizado principalmente ANN para la construcción de un geotermómetro virtual que permita relacionar la BTH con los solutos como el sodio, potasio, magnesio, li-

tio, entre otros (Santoyo y Díaz-González, 2010; Díaz-González y cols., 2013). Para el caso de geotermómetros de gases también se han utilizado ANN empleando como variables explicativas dióxido de carbono, sulfuro de hidrógeno, argón, nitrógeno, entre otros (Pérez-Zárate y cols., 2015).

No obstante, tanto los simuladores de evolución de temperatura en la columna, como los geotermómetros, necesitan mediciones reales de las temperaturas en los pozos. En particular, en el caso de los geotermómetros de solutos estos tienden a sobre estimar las BHT debido en gran parte al hecho de que la cantidad de datos disponibles para el desarrollo es pequeña o su origen no es confiable. Estas bases de datos son pequeñas y poco confiables a causa de mediciones incompletas las cuales son eliminadas, lo cual sesga los modelos estimados. Para resolver esto, los autores de (Román-Flores, Santamaría-Bonfil, Arroyo-Figueroa, y Díaz-González, 2019) propusieron modelos de imputación clásicos y SVM para completar una base de solutos de pozos geotérmicos.

Otras aplicaciones de IA/RP en geotermia son la estimación de caídas de presión en pozos geotérmicos usando ANN y herramientas de simulación de pozos (Bassam, Álvarez Del Castillo, García-Valladares, y Santoyo, 2015).

6.2.4. Energía Nuclear

La energía nuclear proviene de la división de los átomos en un reactor, principalmente de elementos capaces de realizar la fisión nuclear como el uranio o el plutonio, para calentar y evaporar agua haciendo girar una turbina y generar electricidad. Las plantas o centrales nucleares aprovechan este principio para generar energía eléctrica. Sin embargo, debido a las características de alto peligro de los combustibles nucleares, es crítico para las plantas el manejo apropiado de las fallas en los sistemas de control. Mediante la toma de deci-

siones, las fallas mecánicas y/o eléctricas de las diferentes variables que mantienen el reactor en funcionamiento deben atenuarse para restablecer el equilibrio. Por ejemplo, en (A. Laureano-Cruces, Ramírez-Rodríguez, Mora-Torres, y Espinosa-Paredes, 2006) se emplearon Mapas Cognitivos Difusos (FCM) para representar el proceso de toma de decisiones requerido para accidentes por pérdida de refrigerantes, obteniendo resultados congruentes con la opinión de los expertos. En (Espinosa-Paredes, Nuñez Carrera, Laureano-Cruces, Vázquez-Rodríguez, y Espinosa-Martínez, 2008) también se utilizaron FCM para representar la toma de decisiones para procedimientos de operación de emergencia durante situaciones anormales como la pérdida de refrigerante en un reactor de agua en ebullición. En particular, los reactores de agua en ebullición suelen presentar un comportamiento anómalo del flujo de recirculación en la bomba de reacción conocido como *flujo biestable* y, según estudios preliminares, es producto de un comportamiento caótico. Aunque no este tipo de evento es considerado poco riesgoso es importante poder caracterizarlo. Por lo cual en (Nuñez Carrera, Prieto-Guerrero, Espinosa-Martínez, y Espinosa-Paredes, 2009) se emplearon transformadas de ondeleta (*Wavelet* en inglés) continuas y discretas en un análisis multiresolución, para caracterizar las series de tiempo de este fenómeno.

6.2.5. Sistemas de Potencia

Los sistemas de potencia son redes de componentes eléctricos los cuales suministran, transfieren y usan la energía eléctrica. Un ejemplo de estos son las redes eléctricas, las cuales proporcionan energía en un área extendida. Existen varios tipos de aplicaciones en sistemas de potencia, de las cuales, los grupos mexicanos de IA/RP han trabajado principalmente en Sistemas de Aislamiento Controla-

do (AC) (Gamez, Sanchez, y Ricalde, 2011; Gamez-Urias, Sanchez, y Ricalde, 2015; Cruz-May, Ricalde, Atoche, Bassam, y Sanchez, 2018; Arrieta-Paternina, Zamora-Mendez, Ortiz-Bejar, Chow, y Ramirez, 2018), disturbios en la calidad de la energía (Valtierra-Rodriguez, Romero-Troncoso, Osornio-Rios, y Garcia-Perez, 2014) (Idarraga, Guillen, Ramirez, Zamora, y Arrieta-Paternina, 2015), (Guillen y cols., 2018), y planeación energética (Zuñiga, Santamaría-Bonfil, Arroyo-Figueroa, y Batres, s.f.; Palacios, Valdes, y Batres, 2018).

Los métodos de AC son utilizados para mejorar significativamente el manejo de perturbaciones del sistema eléctrico. Estas perturbaciones pueden causar apagones y dañar equipos eléctricos a diferentes niveles en la red. Por lo tanto, la función de un modelo AC es definir los puntos de separación entre una isla (i.e. sección de la red) y el resto del sistema para reducir cargas, aumentar o reducir la generación en un área, y apagar los interruptores necesarios para aislar la perturbación. Por ejemplo, cuando una perturbación sucede en un sistema de potencia estresado, los generadores síncronos interconectados del sistema experimentan oscilaciones electromecánicas locales o entre áreas. Es tarea de los sistemas de AC separar los generadores del sistema en grupos *coherentes* para evaluar la seguridad de la red y diseñar medidas apropiadas para mantener su estabilidad. En este sentido, el trabajo de (Arrieta-Paternina y cols., 2018) permite agrupar de forma temporal los generadores coherentes, extrayendo los modos oscilatorios inter-áreas usando una técnica de identificación modal y el uso de agrupamiento jerárquico aglomerativo. Otro concepto, altamente relacionado a los sistemas AC son las Micro Redes (μR). Una μR es una red de distribución activa, que conecta diferentes fuentes de energías renovables/no-convencionales y cargas, equipada con esquemas de control para garantizar su funcionamiento como un sistema aislado de la red. Varios proyectos de μR se han implementado,

por ejemplo en (Gamez y cols., 2011; Gamez-Urias y cols., 2015) se describe un enfoque de optimización de una μR usando ANN recurrentes combinado con un Sistema Multiagente (MAS), para determinar la cantidad óptima de energía eólica, solar y baterías (incluida la de un vehículo eléctrico), a fin de minimizar la cantidad de energía requerida de la red pública. En (Cruz-May y cols., 2018) se aplicaron ANN recurrentes de alto orden para predecir las variables energéticas presentes en la μR con el objetivo de determinar las cantidades óptimas de energía para los sistemas eólicos, solares, bancos de baterías y redes públicas.

Como se mencionó anteriormente, los disturbios en la calidad de la energía se han convertido en una preocupación de alta prioridad debido a sus consecuencias. Fuentes, transitorios, cortes de suministro y ruido son algunos de los disturbios más comunes. Por lo cual es necesario desarrollar herramientas que permitan la detección y clasificación de dichas perturbaciones. Por ejemplo, en (Valtierra-Rodriguez y cols., 2014) una metodología basada una Red Adaptativa Lineal (ADALINE) y ANN fue propuesta para clasificar 8 tipos de disturbios y sus combinaciones. Otros trabajos como (Idarraga y cols., 2015; Guillen y cols., 2018) proponen utilizar una descomposición en ondículas para caracterizar cada disturbio para posteriormente aplicar la norma Euclidiana (Idarraga y cols., 2015) o modelos de aprendizaje automático como K-Vecinos Cercanos, ANN, SVM, y Redes Bayesianas (Guillen y cols., 2018).

La planificación en el sector energético implica múltiples objetivos. Entre los requerimientos necesarios para la creación de políticas de planificación y funcionamiento del sistema eléctrico se encuentra el pronóstico de la demanda eléctrica. Este consiste en la predicción de la energía demandada en una región específica dada un horizonte de pronóstico determinado. La falta de precisión en la estimación de la demanda eléctrica puede causar una sobre o sub estimación de

la misma lo que puede causar disturbios en el balance de energía y sobre costos. En (Zuñiga y cols., s.f.) los autores propone un método basado en reglas de asociación y el método *a priori* para el pronóstico a corto plazo de la demanda eléctrica en el Norte de México con $HP = 2$ horas en el futuro. Respecto a la planeación a largo plazo, los autores de (Palacios y cols., 2018) presentan el problema de Expansión de Generación (GEP) el cual consiste en determinar la capacidad de generación requerida para satisfacer la demanda estimada típicamente con 10–30 años de antelación. Para resolver este problema, se plantea el uso de Optimización Multi-Objetivo el cual considera tres criterios opuestos: impacto económico, impacto ambiental e impacto ambiental. El planteamiento del problema genera 37,800 variables y 36,078 restricciones y es resuelto usando NSGA-II.

6.3. Tareas pendientes de AS/RP en el campo de la energía

Para concluir este capítulo, se incluye una reflexión sobre los pendientes por hacer en el campo del AS/RP. Los pendientes van de la mano con otras áreas de la computación y de los avances tecnológicos del sector eléctrico. Específicamente existen áreas de oportunidad para AS/RP dada la creciente tendencia a la adquisición masiva de datos y por consiguiente, la acumulación de grandes cantidades de información. El ejemplo más importante de esta tendencia es el concepto de redes eléctricas inteligentes (REI) (Smart grid en inglés). Es una tendencia que todos los países están adoptando a diferentes velocidades y con diferentes intensidades (Bush, Goel, y Simard, 2013). Se trata de integrar los sistemas eléctricos con tecnologías de información y comunicaciones (TICs). Las REI están for-

mados con la integración de los componentes del sistema eléctrico tradicional: generación, transmisión, distribución y comercialización de la electricidad, más un sistema de informática y comunicaciones. La arquitectura de las REI está integrada por la generación, transmisión y distribución, además de los clientes y los proveedores de servicios quienes ofertan la electricidad a los clientes, la instalación y el mantenimiento. La implantación de una REI implica la instalación de instrumentación nueva que cuente con sensores para la adquisición de información en todo el espectro de la REI. Se requiere por lo tanto de sistemas informáticos que almacenen y analicen información geográfica y estadística, así como el estado de la red eléctrica. Sistemas informáticos de adquisición de datos para control y monitoreo de equipos de campo, además de otros sistemas informáticos que otorguen la energía de manera segura, económica y confiable. Se requieren medidores inteligentes que permitan la comunicación entre los medidores y la empresa proveedora, así como un sistema de administración de datos que recopile la información de los distintos servidores y la procese. En suma, la implantación de las REI implica básicamente la adquisición de grandes volúmenes de información nunca antes considerado en el sector de electricidad. Datos sobre el consumo en el tiempo de cada cliente, su localización geográfica. Datos del estado de cada componente que participa en la generación, la transmisión, la distribución. Datos sobre la disposición momentánea de generación fotovoltaica o eólica. Datos sobre segmentos de la red en estado de falla. Datos sobre el estado de los activos que permita programar mantenimientos óptimos para alargar la vida de los equipos. Y muchas aplicaciones más.

Con esta gran cantidad de datos, se extrae conocimiento para resolver problemas como los siguientes:

- Monitoreo y diagnóstico de cada elemento de la REI.

- Planificación del despacho de energía considerando fuentes tradicionales y renovables, además de la expansión de la red.
- Optimización de la operación de la REI en presencia de fallas, es decir, limitar la zona sin servicio a la menor cantidad de clientes y en el menor tiempo.
- Toma de decisiones de control, de operación, de expansión o de mantenimiento de cada elemento de la REI.
- Pronóstico de recurso de viento y sol para generación renovable o pronóstico de consumo.

En México, derivado de la reforma energética y de la tendencia mundial, se realizan esfuerzos importantes en la integración de REI en el sistema eléctrico mexicano. En 2016, la Secretaría de Energía (SENER) publicó el programa de Redes Eléctricas Inteligentes (de Energía, 2016) que define los motivadores, catálogo de funciones y beneficios de tecnologías disponibles (Meneses-Ruáz y Carlos F. García-Hernández, 2016). En 2014 se instaló el grupo de trabajo para REI (GTREI), conformado por especialistas de SENER, de la Comisión Federal de Electricidad (CFE) y la Comisión Reguladora de Energía (CRE). Este grupo de trabajo plantea el mapa de ruta para la implantación de REI con el fin de modernizar sistemas, incorporar fuentes renovables y brindar un mejor servicio a los usuarios. Por la parte académica, en Conacyt, a través de los fondos sectoriales de sustentabilidad, crearon el Centro Mexicano de Innovación en Energía de redes eléctricas inteligentes, el CEMIE-Redes liderado por el Instituto Nacional de Electricidad y Energías Limpias (INEEL), (antes IIE, Instituto de Investigaciones Eléctricas) (<https://www.ineel.mx/cemie-redes.html>). Se espera que los proyectos del CEMIE-Red inicien este mismo año 2018 como

una oportunidad de integrar técnicas de AS/RP en la solución de problemas de las REI. Por ejemplo, en la gestión inteligente de activos como en el diseño de agentes inteligentes que ayuden en la toma de decisiones del mercado mexicano de energía.

Desde una visión desde la computación, se requiere de técnicas novedosas como Internet de las cosas (IoT) y “Big-data”. El reto que se tiene enfrente es la integración de técnicas en esas dos nacientes áreas con la tradicional de AS/RP. En realidad, se tendrá que utilizar las técnicas de AS/RP en IoT y Big-data, lo que requerirá de nuevos estudios y tesis sobre la fundición de estas técnicas. Por ejemplo, en el sector de la energía, el desarrollo de las redes eléctricas inteligentes está proveyendo información en grandes cantidades que antes no se imaginaba. Información de cada usuario, de cada colonia, ciudad o país. A nivel de poste, de subestaciones u de centrales generadoras. Esto empieza a producir enormes cantidades de información que deberá ser procesada para sacarle provecho. Se deberá obtener los patrones de usuarios en regla, los irregulares y los ilegalmente colgados. Se deberá encontrar el patrón de las redes que requieren mantenimiento de las que están a punto. En suma, se deberá integrar IoT + Big-data + AS/RP.

Referencias

- Alejos-Moo, E., Tziu-Dzib, J., Canto-Esquivel, J., y Bassam, A. (2018). Deep Learning for Photovoltaic Power Plant Forecasting. En *Intelligent computing systems. isics 2018* (Vol. 820, pp. 56–66). Springer International Publishing. doi: 10.1007/978-3-319-76261-6_5
- Álvarez del Castillo, A., Santoyo, E., y García-Valladares, O. (2012). A new void

- fraction correlation inferred from artificial neural networks for modeling two-phase flows in geothermal wells. *Computers and Geosciences*, 41, 25–39. doi: 10.1016/j.cageo.2011.08.001
- Arrieta-Paternina, M., Zamora-Mendez, A., Ortiz-Bejar, J., Chow, J., y Ramirez, J. (2018). Identification of coherent trajectories by modal characteristics and hierarchical agglomerative clustering. *Electric Power Systems Research*, 158, 170–183. doi: 10.1016/j.epsr.2017.12.029
- Bassam, A., Álvarez Del Castillo, A., García-Valladares, O., y Santoyo, E. (2015). Determination of pressure drops in flowing geothermal wells by using artificial neural networks and wellbore simulation tools. *Applied Thermal Engineering*, 75, 1217–1228. doi: 10.1016/j.applthermaleng.2014.05.048
- Bassam, A., Santoyo, E., Andaverde, J., Hernández, J., y Espinoza-Ojeda, O. (2010). Estimation of static formation temperatures in geothermal wells by using an artificial neural network approach. *Computers and Geosciences*, 36(9), 1191–1199. doi: 10.1016/j.cageo.2010.01.006
- Bush, S. F., Goel, S., y Simard, G. (2013, Sept). Ieee vision for smart grid communications: 2030 and beyond roadmap. *IEEE Vision for Smart Grid Communications: 2030 and Beyond Roadmap*, 1–19. doi: 10.1109/IEEESTD.2013.6690098
- Cadenas, E., Jaramillo, O., y Rivera, W. (2010). Analysis and forecasting of wind velocity in Chetumal, Quintana Roo, using the single exponential smoothing method. *Renew. Energy*, 35(5), 925–930. doi: 10.1016/j.renene.2009.10.037
- Cadenas, E., y Rivera, W. (2010, diciembre). Wind speed forecasting in three different regions of Mexico, using a hybrid ARIMA-ANN model. *Renew. Energy*, 35(12), 2732–2738. doi: 10.1016/j.renene.2010.04.022

- Cruz-May, E., Ricalde, L., Atoche, E., Bassam, A., y Sanchez, E. (2018). Forecast and Energy Management of a Microgrid with Renewable Energy Sources Using Artificial Intelligence. En *Intell. comput. syst. isics 2018* (Vol. 820, pp. 81–96). Springer. doi: 10.1007/978-3-319-76261-6_7
- Dabbaghchi, I., Christie, R., Rosenwald, G., y Liu, C. (1997). Ai application areas in power systems. *IEEE Expert*, 12(2), 58–66.
- de Energ  , S. (2016). *Programa de Redes El  ctricas inteligentes* (Technical Report). Ciudad de M  xico, M  xico: Secretar   de Energ  , M  xico.
- D  az-Gonz  lez, L., Hidalgo-D  vila, C., Santoyo, E., y Hermosillo-Valadez, J. (2013). Evaluaci  n de t  cnicas de entrenamiento de redes neuronales para estudios geotermom  tricos de sistemas geot  rmicos. *Revista Mexicana de Ingenier   Qu  mica*, 12(1), 105–120.
- Espinosa-Paredes, G., Nu  ez Carrera, A., Laureano-Cruces, A. L., V  zquez-Rodr  guez, A., y Espinosa-Mart  nez, E. (2008). Emergency management for a nuclear power plant using fuzzy cognitive maps. *Annals of Nuclear Energy*, 35(12), 2387–2396. doi: 10.1016/j.anucene.2008.07.007
- Espinosa-Paredes, G., Olea-Gonzalez, U., Espinosa-Martinez, E., V  zquez-Rodr  guez, A., Romero-Paredes, H., y V  zquez-Rodr  guez, R. (2010). Inverse simulation methods to obtain formation temperatures in Wellbores. *Petroleum Science and Technology*, 28(12), 1250–1259. doi: 10.1080/10916460903030458
- Flores, J., Graff, M., y Cadenas, E. (2005). Wind prediction using genetic algorithms and gene expression programming. *Proc. Int. Conf. Model. Simul. Enterp. AMSE 2005*, 1–9.
- Flores, J., Graff, M., y Rodriguez, H. (2012). Evolutive design of ARMA and ANN models for time series forecasting. *Renew. Energy*, 44, 225–230.

doi: 10.1016/j.renene.2012.01.084

Flores, J., Loaeza, R., Rodríguez, H., y Cadenas, E. (2009). Wind speed forecasting using a hybrid neural-evolutive approach. En *Micaí 2009 adv. artif. intell.* (pp. 600–609). doi: 10.1007/978-3-642-05258-3_53

Gamez, M., Sanchez, E., y Ricalde, L. (2011). Optimal operation via a recurrent neural network of a wind-solar energy system. En *Proc. Int. Jt. Conf. Neural networks* (pp. 2222–2228). doi: 10.1109/IJCNN.2011.6033505

Gamez-Urias, M., Sanchez, E., y Ricalde, L. (2015). Electrical Microgrid Optimization via a New Recurrent Neural Network. *IEEE Syst. J.*, 9(3), 945–953. doi: 10.1109/JSYST.2014.2305494

Graff, M., Peña, R., y Medina, A. (2013, jun). Wind speed forecasting using genetic programming. En *2013 IEEE Congr. Evol. Comput.* (pp. 408–415). IEEE. doi: 10.1109/CEC.2013.6557598

Graff, M., Peña, R., Medina, A., y Escalante, H. J. (2014). Wind speed forecasting using a portfolio of forecasters. *Renew. Energy*, 68, 550–559. doi: 10.1016/j.renene.2014.02.041

Guillen, D., Arrieta-Paternina, M., Ortiz-Bejar, J., Tripathy, R., Zamora-Mendez, A., Tapia-Olvera, R., y Tellez, E. (2018). Fault detection and classification in transmission lines based on a PSD index. *IET Generation, Transmission & Distribution*, 12(18), 4070–4078. doi: 10.1049/iet-gtd.2018.5062

Herbert-Acero J. and Franco-Acevedo J. and Valenzuela-Rendón M. and Probst-Oleszewski, O. (2009). Linear Wind Farm Layout Optimization through Computational Intelligence. *MICAÍ 2009 Adv. Artif. Intell. Lect. Notes Comput. Sci. Vol. 5845*, 692–703. doi: 10.1007/978-3-642-05258-3_61

- Ibargüengoytia, P., Reyes, A., Borunda, M., y García, U. (2018). Predicción de potencia eólica utilizando técnicas modernas de inteligencia artificial. *Revista Ingeniería Investigación y Tecnología*, 19(4), 1–11. doi: <http://dx.doi.org/10.22201/fi.25940732e.2018.19n4.033>
- Ibargüengoytia, P. H., Liñan, R. L., y Betancourt, E. (2010). Transformer diagnosis using probabilistic vibration models. En *2010 IEEE PES transmission and distribution*. New Orleans, U.S.A..
- Idarraga, G., Guillen, D., Ramirez, J., Zamora, A., y Arrieta-Paternina, M. (2015). Detection and classification of faults in transmission lines using the maximum wavelet singular value and Euclidean norm. *IET Generation, Transmission & Distribution*, 9(15), 2294–2302. Descargado de <http://digital-library.theiet.org/content/journals/10.1049/iet-gtd.2014.1064> doi: 10.1049/iet-gtd.2014.1064
- Laureano-Cruces, A., Ramírez-Rodriguez, J., Mora-Torres, M., y Espinosa-Paredes, G. (2006). Modeling a Decision Making Process in a Risk Scenario: LOCA in a Nucleoelectric Plant Using Fuzzy Cognitive Maps. *Research in Computer Science*, 26(1), 3–14.
- Laureano-Cruces, A. L., y Espinosa-Paredes, G. (2005). Behavioral design to model a reactive decision of an expert in geothermal wells. *International Journal of Approximate Reasoning*, 39(1), 1–28. doi: 10.1016/j.ijar.2004.08.002
- Lauritzen, S. L. (1995). The em algorithm for graphical association models with missing data. *Computational Statistics & Data Analysis*, 19, 191–201.
- May-Tzuc, O., Bassam, A., Escalante-Soberanis, M., Venegas-Reyes, E., Jaramillo, O., Ricalde, L., ... El-Hamzaoui, Y. (2017). Modeling and optimi-

- zation of a solar parabolic trough concentrator system using inverse artificial neural network. *Journal of Renewable and Sustainable Energy*, 9(1). Descargado de <http://dx.doi.org/10.1063/1.4974778> doi: 10.1063/1.4974778
- Meneses-Ruiz, J., y Carlos F. García-Hernández, J. C. M.-C. (2016). Factores impulsores para el desarrollo del marco regulatorio para redes eléctricas inteligentes en México. En *CIINDET, 2016*. Cuernavaca, Mor., México.
- Núñez Carrera, A., Prieto-Guerrero, A., Espinosa-Martínez, E. G., y Espinosa-Paredes, G. (2009). Analysis of a signal during bistable flow events in Laguna Verde Nuclear Power Station with wavelets techniques. *Nuclear Engineering and Design*, 239(12), 2942–2951. doi: 10.1016/j.nucengdes.2009.09.008
- Olea-Gonzalez, U., Vázquez-Rodríguez, A., García-Gutiérrez, A., y Anguiano-Rojas, P. (2007). Estimation of Formation Temperatures in Reservoirs: Inverse Method. *Revista Mexicana de Ingeniería Química*, 6(1), 65–74.
- Palacios, R., Valdes, E., y Batres, R. (2018). A multi-step multi-objective generation expansion planning model-A case study in Mexico. *Int. J. Smart Grid Clean Energy*, 7(1), 90–97. doi: 10.12720/sgce.7.2.90-97
- Pérez-Zárate, D., García-López, C., Bassam, A., Acevedo, A., Díaz-González, L., y Santoyo, E. (2015). Performance evaluation of neural network architectures for the multivariate analysis of gas compositions and the prediction of geothermal reservoir temperatures. En *Tercer Simposio Internacional sobre Energías renovables y Sustentabilidad* (p. 1). Descargado de <https://colecciondigital.cemiegeo.org/xmlui/handle/123456789/556?show=full>
- Reyes, A., Ibargüengoytia, P., Jijón, J., Guerrero, T., García, U., y Borunda, M.

- (2016). Wind power forecasting for the villonaco wind farm using a.i. techniques. En *15th Mexican International Conference on Artificial Intelligence* (pp. 1–10). Springer International Publishing.
- Ricalde, L. J., y Cruz, B. J. (2010). High Order Recurrent Neural Control for Wind Turbine with a Control Neuronal Recurrente de Alto Orden para Turbinas de Viento con. *Computación y Sistemas*, 14(2), 133–143.
- Rodriguez, H., Flores, J., Puig, V., Morales, L., Guerra, A., y Calderon, F. (2017). Wind speed time series reconstruction using a hybrid neural genetic approach. *IOP Conf. Ser. Earth Environ. Sci.*, 93(1), 0–9. doi: 10.1088/1755-1315/93/1/012020
- Román-Flores, A., Santamaría-Bonfil, G., Arroyo-Figueroa, G., y Díaz-González, L. (2019). Single imputation methods applied to a global geothermal database. En *17th Mexican International Conference on Artificial Intelligence* (pp. 1–13). Springer.
- Sánchez-Pérez, P. A., Robles, M., y Jaramillo, O. (2016). Real time Markov chains: Wind states in anemometric data. *J. Renew. Sustain. Energy*, 8(2). doi: 10.1063/1.4943120
- Santamaría-Bonfil, G., Reyes-Ballesteros, A., y Gershenson, C. (2016). Wind speed forecasting for wind farms: A method based on support vector regression. *Renew. Energy*, 85, 790–809. doi: 10.1016/j.renene.2015.07.004
- Santoyo, E., y Díaz-González, L. (2010). A New Improved Proposal of the Na / K Geothermometer to Estimate Deep Equilibrium. En *World geothermal congress 2010* (pp. 25–29).
- Valtierra-Rodriguez, M., Romero-Troncoso, R., Osornio-Rios, R., y Garcia-Perez, A. (2014). Detection and classification of single and combined power quality disturbances using neural networks. *IEEE Trans. Ind. Elec-*

tron., 61(5), 2473–2482. doi: 10.1109/TIE.2013.2272276

Zuñiga, M., Santamaría-Bonfil, G., Arroyo-Figueroa, G., y Batres, R. (s.f.). An association-rule method for short-term electricity demand forecasting and consumption pattern recognition. En *17th Mex. Int. Conf. Artif. Intell.* (pp. 1–10). in press, Springer.

Capítulo 7

Patrones en Medicina y Biología

FELIPE ORIHUELA-ESPINA, INAOE

JAVIER HERRERA-VEGA, INAOE

7.1. Introducción

Al sensado o extracción de datos de sistemas vivos, nos referimos colectivamente bajo el término imagenología biomédica, y es un área multidisciplinaria con aportaciones de la física, la matemática, la química y por supuesto las disciplinas de interés en este capítulo; la computación, la medicina y la biología. Estos datos biomédicos pueden tomar diferentes formas; señal (1D temporal), imagen (2D espaciales), volumen (3D espaciales) o video (2D espaciales y 1D temporal), 4D (3D espaciales y 1D temporal), o superiores (en casos multimodales y/o funcionales). La adquisición de los mismos puede tener una finalidad científica ya sea médica o biológica, clínica en el caso médico, o para explotación industrial en el caso de biología. El AS/RP es una componente fundamental de esta extracción de información mejorando la calidad de los datos, reduciendo

tiempos de análisis, o estimando descriptores que los caractericen para facilitar su interpretación. A esta subárea nos referimos comúnmente como procesamiento y/o análisis de señales biomédicas. En el mundo digitalizado actual, son pocas las aplicaciones que involucran datos biomédicos y que no están intervenidas en algún momento por algún proceso de AS/RP. Este capítulo pretende esbozar la magnitud de esta aportación. Para ello, hemos organizado el capítulo con secciones dedicadas acorde al fin para el que fueron obtenidos los datos biomédicos mediante la imagenología. El capítulo además incluye una sección adicional centrada expresamente en los trabajos que actualmente se desarrollan en México.

El desarrollo y uso de las técnicas de análisis de señales y reconocimiento de patrones (AS/RP) para el análisis de la información de las áreas de medicina y biología se agrupan a menudo bajo el nombre de procesamiento y/o análisis de señales biomédicas (Rangayyan, 2015; Theis y Meyer-Bäse, 2010), a su vez bajo el paraguas del área de imagenología biomédica. La **imagenología o imagenología biomédica** es el área de la ciencia que se encarga de extraer información de sistemas vivos. Esta incluye no sólo la parte de AS/RP que nos concierne, sino también conocimiento en física, química, matemáticas en general y por supuesto las propias biología y medicina. En principio, bajo el término de imagenología biomédica se consideran tanto señales, término normalmente asociado a series temporales unidimensionales, como imágenes, término normalmente asociado a fotografías bidimensionales, así como videos o 4D (volúmenes espaciales más tiempo), que son datos de más alta dimensionalidad. En este capítulo, cuando la dimensionalidad de los datos no sea crítica, la distinción entre estos términos será difusa, y se utilizarán los términos de señal e imagen de forma intercambiable.

La imagenología biomédica cubre la adquisición de imágenes: tanto *in-vivo* como *ex-vivo*, *estructurales* o anatómicas como *funcionales* o fisiológicas y metabólicas, así como, generadas mediante diferentes principios físicos (e.g. electrofisiología, resonancia magnética, ultrasonidos, etc). Estas imágenes pueden adquirirse con finalidad científica para generar conocimiento, o clínica para aplicar el conocimiento y en este segundo uso como veremos cubriendo las etapas de prevención, diagnóstico, tratamiento y monitoreo. La imagenología biomédica cubre todo el proceso del flujo de la información desde la formación y adquisición, a la interpretación (Herrera-Vega, Treviño Palacios, y Orihuela-Espina, 2017). El AS/RP está presente en todo este flujo en casi todas sus manifestaciones estadísticas y por supuesto computacionales; descomposición de fuentes, filtrados para la mejora de la calidad de las imágenes, extracción de información histofisiológica, inferencia, estimación de relaciones asociativas o causales, medición de calidad de los modelos, soporte a la toma de decisiones, etc. Su implementación puede darse *in-silico* en soporte hardware o software, y de manera interesante incluso *in-vivo* aprovechando las fuerzas que rigen los enlaces fisicoquímicos moleculares, estas últimas aún muy rudimentarias, y fuera del ámbito de este capítulo. En este capítulo no obstante, nos centraremos únicamente en el procesamiento *in-silico* y no prestaremos mucha atención al soporte en que se implemente o el lenguaje de programación con el que se operacionalice.

Aunque, las imágenes biomédicas existen desde el siglo XIX, la llegada de las computadoras, la digitalización de las imágenes y su tratamiento automatizado mediante las técnicas de AS/RP supusieron una revolución en nuestro conocimiento médico y biológico. Para la adquisición de la información biológica muchas técnicas de imageología están disponibles. Sin ser exhaustivos; Najarian and Splinter (Najarian y Splinter, 2006) proveen una introducción suave a

los principios subyacentes a la electroencefalografía (EEG), tomografía computarizada (CT) y rayos X, resonancia magnética (MRI) y tomografía por emisión de positrones (PET). Cada una de las modalidades de imagenología cuentan a su vez con muchas submodalidades con diferente aprovechamiento. Otras técnicas pueden consultarse en obras más especializadas; por ejemplo, microscopía y sus diferentes variantes, nanoscopía, tomografía computarizada de emisión de un positron (SPECT), ultrasonido, óptica difusa, o la innovadora CLARITY las diferentes variantes de electrofisiología; electroencefalografía, electrocorticografía, electrocardiografía (ECG o EKG), electromiografía, conductancia de la piel, etc, entre otras fuentes. Esta es la información de la que parte el AS/RP, y que el AS/RP debe respetar para que la interpretación de los resultados esté confinada a las particularidades de cada fuente de datos.

7.1.1. Aplicaciones

Las aplicaciones de las soluciones computacionales de AS/RP en imagenología biomédica son numerosas; estudios farmacológicos y anestesiología, bioquímicos, registro de información biológica, incluyendo por supuesto aplicaciones no directamente biomédicas pero si con información biomédica, como por ejemplo la seguridad basada en marcadores biométricos de lectura de retina o huellas dactilares, monitoreo continuo en quirófano o unidades de cuidados intensivos, aplicaciones de interacción humano computadora e.g. cómputo afectivo, incluyendo la siempre popular interfaz cerebro computadora (BCI), cuidado neonatal y compresión del neurodesarrollo, monitoreo del feto, estudio de enfermedades e.g. epilepsia, eventos cerebrovasculares, condiciones neurológicas y neurorehabilitación, desórdenes psiquiátricos y psicológicos, oncología, traumatología, las diferentes subespecialidades de cirugía, tanto en las variedades tra-

dicionales, como en las mínimamente invasivas y cirugía robótica, genética y epigenética así como la proteómica u otras -ómicas e incluso más recientemente en la conocida como medicina de precisión. En general, ya no hay rama de la medicina y la veterinaria, o la biología celular donde el AS/RP no se haya vuelto simplemente imprescindible.

7.1.2. Estructura y orientación del capítulo

No hay una única forma de organizar este capítulo de manera adecuada para cada lector, y no es posible la cobertura exhaustiva, debido a que el alcance y la riqueza del área hacen imposible que un sólo autor (¡o varios!) tengan un conocimiento tan amplio como sería requerido. Así pues, teniendo en cuenta estas dos consideraciones, hemos optado por organizar el capítulo según la finalidad; científica en la Secc. 7.2, médica en la Secc. 7.3 y biológica en la Secc. 7.4, primero de forma general, y luego en México en las Secc. 7.5. Así mismo, pretendemos no poner énfasis en alguna técnica de análisis, o modalidad de imagen en particular, la intención es simplemente narrar un panorama general dando algunos ejemplos en el camino, aunque, nuestro propio sesgo formativo sin duda resultará en una mayor abundancia de ejemplos de nuestra área, las neuroimágenes funcionales.

7.2. Uso para el desarrollo del conocimiento científico biomédico

La ciencia establece relaciones entre causas y efectos naturales para explicar nuestro universo. Estas relaciones se establecen a partir del registro sistemático de observaciones de los fenómenos. En las ciencias biomédicas estas observaciones

son las señales e imágenes biomédicas. La manipulación (principalmente digitalizada) automatizada de estas señales e imágenes biomédicas mediante técnicas de AS/RP es a veces simplemente conveniente, y en la mayoría de los casos absolutamente imprescindible por ejemplo por cuestiones de volumen, precisión o reproducibilidad.

El método científico formaliza dos grandes variantes para la construcción del nuevo conocimiento científico; el guiado por hipótesis y el guiado por datos, cuya diferencia fundamental es el momento en el que se formula la hipótesis o modelo del fenómeno observado; previo o posterior a la experimentación, con las sabidas implicaciones, virtudes y limitaciones de cada una de estas aproximaciones. Bajo cualquiera de estas dos variantes del método científico, tras la adquisición de las observaciones experimentales, la información adquirida debe ser:

- Transformada del registro crudo a la información histofisiológica de interés, lo que conocemos como *reconstrucción*
- Aislada de información no deseada pero adquirida de forma concomitante por el mecanismo de imagenología -reducción del ruido-, y reexpresada en un formato más conveniente para su estudio, lo que conocemos como *procesamiento* o pre-procesamiento (según las diferentes subáreas del conocimiento).
- Caracterizada y resumida en términos de estadísticas o descriptores, a menudo en la jerga del campo referidos como características, y sintetizadas en modelos matemáticos, lo que conocemos como *análisis*.

- Y explotadas para la extracción y generación de nuevo conocimiento confinada por los sesgos que afecten al modelo, lo que conocemos como *interpretación*.

7.2.1. Reconstrucción

Dependiendo del sistema de sensado, la información puede o no estar directamente disponible en las variables de interés. El registro crudo recogido por los transductores (a menudo voltajes o intensidades) debe ser traducido a información histofisiológica. Por ejemplo, en resonancia magnética la información cruda viene en el llamado espacio k , que es una descomposición en fase y frecuencia de la señal, y la transformada de Fourier permite la reconstrucción de la información anatómica. Otras técnicas de imagenología requieren de otras operaciones de reconstrucción más elaborados. Por ejemplo, en espectroscopía infraroja las intensidades son primero transformadas a propiedades ópticas, y a partir de estas, las concentraciones de algún cromóforo son reconstruidas mediante sistemas de ecuaciones o alternativas más elaboradas basadas en matemática inversa y cómputo científico. La reconstrucción del número y tamaño de las partículas dispersivas a partir de las señales de moteado es central en las variantes de espectroscopía de onda. En electroencefalografía, técnicas como LORETA reconstruyen la localización de los dipolos a partir del montaje en superficie. En tomografía, la transformada de Radon reconstruye la información estructural a partir de la impresión balística de los rayos X en varios planos. La capacidad de las redes convolucionales para capturar información espacial ha beneficiado el incremento de la resolución en técnicas de ultrasonido. Por supuesto, no podía faltar el ahora pervasivo aprendizaje profundo, por ejemplo reduciendo el espacio k .

7.2.2. Procesamiento

En el procesamiento, las técnicas se centran en la descomposición de los registros crudos en nuevos registros que separarán la parte de interés, llamada simplemente señal, de la que no era de interés, referida como ruido. A finales del XIX y e inicios del XX, ya se procesaba la información biomédica, aunque este procesamiento no era digitalizado. Por ejemplo, los nobeles Camilo Golgi y Santiago Ramón y Cajal, padres de las neurociencias, usaban tintes para resaltar estructuras histológicas en su estudio de los tejidos neurológicos. Poco a poco, y a medida que la capacidad de cómputo se incrementaba, el procesamiento y análisis de la información biomédica se fue sofisticando. Técnicas de filtrado clásico (Proakis y Manolakis, 2007) o adaptativo han permitido realzar la información de interés y suprimir el ruido indeseado. Este filtrado puede llevarse a cabo en el dominio original de la información, o en algún espacio transformado. Entre los primeros, podemos mencionar operaciones como;

- el filtrado de media móvil popular en las series temporales, y en general las técnicas de análisis de series temporales (Chatfield, 2004) e.g. modelos ARIMA, con aplicaciones biomédicas como epidemiología, política de salud o electroencefalografía.
- las operaciones de curado morfológico incluyendo erosiones, dilataciones, operaciones de sal y pimienta, clausuras, etc filtros de la mediana o la moda, con aplicaciones biomédicas como reconocimiento de la estructura esponjosa del hueso, o de capas retinales en oftalmología.

Entre los segundos, aquellos que operan el filtrado en un espacio transformado, podemos resaltar;

- el diseño de filtros en el espacio de frecuencias alcanzado mediante la transformada de Fourier, esencial para la reconstrucción de la resonancia magnética, o el estudio de la variabilidad del ritmo cardíaco en electrocardiografía, pero también popular en ultrasonido en aplicaciones de cardiología como la detección de infarto de miocardio, o su uso en angiografía tomográfica óptica coherente.
- más recientemente, otras transformadas como las onduletas (wavelets), con aplicaciones como la eliminación de ruido sistémico o la detección de movimiento en el óptodo en neuroimagenes ópticas funcionales, la reducción del ruido de Ricci en las resonancias magnéticas, en EEG la clasificación de eventos evocados o la corrección de artefactos oculares, entre muchas otras.
- las operaciones de regularización, y sus manifestaciones particulares en forma de supermuestreos o submuestreos para corregir desbalances estadísticos y/o inestabilidades matemáticas, con aplicaciones en la reconstrucción de tomografía óptica difusa, microscopía de fluorescencia, la clasificación y visualización de patrones codificados en redes neuronales, mejora de la resolución en tomografía de emisión de positrones y tomografía coherente.
- las factorizaciones o descomposiciones de espacios, con manifestaciones tan populares como análisis de componentes principales (PCA), análisis de componentes independientes (ICA), y en general otras variantes de la separación de fuentes ciega. Las aplicaciones biomédicas de estas técnicas en términos de filtrado han tenido repercusión en aplicaciones como neuroimagenes, genética, eliminación de autofluorescencia en imágenes

de microscopía, la detección de actividad muscular con electromiografía y muchísimas otras.

7.2.3. Análisis

Cuando se trata de AS/RP, en la etapa de análisis la tarea principal se convierte en una de selección y extracción de características, y la posterior generación de modelos regresivos (de salida continua) y clasificatorios (de salida discreta). Quizás es aquí donde el impacto de la AS/RP es más evidente. Los diferentes algoritmos de aprendizaje de máquina y minería de datos han permitido establecer relaciones, que por su complejidad hubiesen sido imposibles desenmascararlas con cálculos manuales. La gestión de los crecientes volúmenes de datos en aplicaciones de salud ha despolvado las viejas estrategias de divide y vencerás y las ha mezclado con las modernas formas de bases de datos relacionales y orientadas a objetos para dar soporte al repositorio de la información así como con técnicas agresivas de paralelización y uso de la capacidad de cómputo de los procesadores gráficos para poder atender las nuevas necesidades. Muchos análisis de la nueva generación de estudios oncológicos, neurológicos, genéticos, farmacológicos y epidemiológicos están basados invariablemente en el AS/RP mediante la bioestadística y la bioinformática. La mezcla de las técnicas geoestadísticas y de sistemas de información geográfica con epidemiología a dado lugar a nuevas ramas completas de conocimiento como la geomedicina, ha permitido una mejor planeación de las etapas translacionales de la investigación médica. En biología tanto vegetal como animal, el conocimiento de las sociedades de organismos hace un aprovechamiento intensivo de modelos de RP; por ejemplo en la detección de especies de plantas invasoras o conocer el estado de las poblaciones de coral en la gran barrera Australiana.

7.2.4. Interpretación

La interpretación se encarga de confirmar (o refutar) el emparejamiento entre un modelo o hipótesis dadas unas observaciones, normalmente determinando si existe alguna entrada que haga plausibles dichas observaciones dado el modelo. Composición de funciones, propagación de errores o análisis de confiabilidad y de residuos, son elementos importantes en la interpretación, especialmente cuando el objetivo es explicar y no simplemente predecir (Shmueli y cols., 2010). El nuevo conocimiento científico biomédico generado a través de la interpretación también se beneficia de las técnicas de AS/RP. Las técnicas de minería de datos y descubrimiento de conocimiento han repercutido en avances en la informática biomédica incluyendo la integración de datos biomédicos.

La exigencia de replicabilidad y reproducibilidad de los experimentos biomédicos ha pasado de ser una discusión teórica y ética, para dar paso a la operacionalización y exigencia práctica (Moher y cols., 2010). La computación en general, y la AS/RP juega un papel angular en este proceso a través de:

- los algoritmos de anonimización y cegado de datos que ayuda a mantener la privacidad de los sujetos experimentales
- la validación y curación de las observaciones con técnicas de imputado
- el depositado en plataformas informáticas y su posterior minería,
- la demanda de la exposición pública de los algoritmos (por ejemplo a través de las plataformas de versionado)
- las nuevas arquitecturas de tuberías (*pipelines* si se nos permite el anglicismo), que facilitan y garantizan la replicabilidad de procesamiento y analítica sobre los datos

- la estandarización de los formatos de ficheros, pensados para que sean a la vez entendibles por un humano y procesables por una máquina con particular énfasis en la estructuración de los metadatos esenciales en el RP no etiquetado.

7.3. Uso aplicado en medicina

Medicina (del latín *mederi*) significa curar, léase hacer desaparecer una enfermedad o daño físico. Las alteraciones estructurales normalmente nos referimos a ellas como traumas, y las fisiológicas como enfermedades o condiciones. El origen de las primeras son las lesiones generadas por un agente mecánico, como puede ser un golpe, e incluye fracturas, esguinces, luxaciones, etc. El origen de las segundas es una falla biológica ya sea causada por agentes internos (genéticos o adquiridos) o externos (e.g. virus y bacterias). En general, esta etiología tanto de los traumas como de las enfermedades y condiciones es muy dispar. La clasificación internacional de enfermedades (ICD por sus siglās en inglés - <https://icd.who.int/> -, actualmente en su versión 11) de la Organización Mundial de la Salud (WHO por sus siglas en inglés) contiene unas 155 mil entradas!. Además esta clasificación se revisa periódicamente incrementando este número; en la revisión anterior por ejemplo “solo” había unas 14400 entradas en comparación.

Para lograr su objetivo de curar, la medicina es la rama del conocimiento que se encarga de estudiar las alteraciones estructurales y fisiológicas del cuerpo humano con respecto a su estado basal y con este conocimiento desarrollar y aplicar soluciones que:

- *Preven*gan la ocurrencia de traumas y enfermedades,

- *Diagnostiquen* la presencia de alguna alteración y permitan su identificación,
- eliminen o alivien los signos, y si procede, los síntomas, de las enfermedades y traumas, lo que nos referimos como el *tratamiento*,
- permitan seguir la evolución del estado de salud de las personas a nivel individual y colectivo, lo que nos referimos como *seguimiento* o *monitoreo*.

7.3.1. Prevención

Hasta hace poco la respuesta médica era principalmente reactiva. Recientemente el énfasis es en una medicina activa que mejore el estado de salud de la población mediante acciones y campañas preventivas. La prevención de enfermedades agrupa a todas aquellas acciones orientadas a evitar la aparición de enfermedades y traumas; estudios de factores de riesgo, campañas de concienciación, políticas de vacunación, y por supuesto el monitoreo continuo de la actividad social. El AS/RP ha sido clave en cada una de estas. Así, ha permitido el establecimiento de relaciones entre comorbilidades previamente desconocidas, ayudando en la identificación de factores de riesgo de diferentes enfermedades tanto genéticos y epigenéticos. Permite una más amplia difusión, alcance y efectividad de las campañas gracias a los medios digitales, por ejemplo a través del cómputo ubicuo, el internet de las cosas o el cómputo afectivo. Favorece la planeación de los tratamientos antiretrovirales y la epidemiológica para mejorar el impacto de las campañas de vacunación. Finalmente, también ha sido un detonador en la vertiente del monitoreo preventivo con manifestaciones como el cómputo pervasivo persuasivo o la detección temprana de enfermedades con ejemplos en el desarrollo de soluciones de tamizaje. Todas estas ramas aparentemente dispares

de la computación parten de datos, ya sea de estudios aleatorizados y controlados, como de estudios observacionales, y llegan a conclusiones y decisiones a través de la manipulación controlada y dirigida de los conjuntos de datos con el fin de extraer información latente; el AS/RP.

7.3.2. Apoyo al diagnóstico

El diagnóstico de enfermedades ha sido un área particularmente fértil para la inteligencia artificial (IA) que ha logrado el diagnóstico de muchas enfermedades, incluso enfermedades raras (aquellas que afectan a menos de 5 a 10 de cada 10000 personas¹) a través de los sistemas expertos y los sistemas de apoyo al diagnóstico. Aunque la IA se ha llevado mercedamente los focos, no podemos obviar que el AS/RP es casi de forma imperativa el paso precedente soportando la preparación de los datos.

El diagnóstico moderno está principalmente basado en signos -aquellas magnitudes cuantificables mediante algún estudio o prueba clínica o la observación por un profesional de la salud- más que en síntomas -aquellas magnitudes cualitativas expresadas por el paciente-. Cuando el estudio de los signos llega a través de la imagenología biomédica, el AS/RP suele estar presente; en el análisis de imágenes histopatológicas de microscopía para el diagnóstico de cáncer, incluyendo linfomas, leucemia o seno entre otros, la eliminación de artefactos de movimiento en resonancias, la detección automatizada de arritmias a partir del electrocardiograma, en trabajos en detección de ataques epilépticos, en el análisis de estudios de polisomnografía para diagnóstico de alteraciones del sueño, segmentación de imágenes de OCT (Tomografía de óptica coherente) de retina

¹El número puede variar ligeramente dependiendo de la autoridad sanitaria.

patológicas, así como otras condiciones; e.g. diabetes, autismo, son sólo algunas de las tareas llevadas a cabo por el AS/RP en su apoyo al diagnóstico.

7.3.3. Tratamiento

Los tratamientos de las enfermedades y traumas pueden ser quirúrgicos, farmacológicos o rehabilitatorios (o una combinación de estos). Los quirófanos modernos son una exhibición continua del AS/RP con la demanda añadida de una respuesta en lo que informalmente conocemos como tiempo real; el procesamiento de las realidades aumentadas, la mejora de las imágenes de las cámaras que acompañan a las diferentes variedades de cirugías mínimamente invasivas, el control de los dispositivos laparoscópicos y robóticos, la restricción de los movimientos del cirujano tanto reduciendo movimientos no deseados e.g. la eliminación de temblores de la mano del cirujano en cirugía robótica es en esencia un filtro adaptativo, como mediante la compensación para estabilizar los instrumentos ante las respuestas musculares, e.g. el ritmo cardiaco, evitando punzadas a otros organos en la navegación imponiendo regiones de restricción activa o la detección de actividades quirúrgicas, entre muchas otras. En el caso de intervenciones quirúrgicas, el soporte de la AS/RP al tratamiento, viene desde la etapa preoperativa, por ejemplo, delineando lesiones de cara a la cirugía.

El apoyo del AS/RP a la farmacología podría verse como parte del uso en el ámbito científico, discutido en la Secc. 7.2, o aplicado. La AS/RP ha tenido y continua teniendo una importancia crítica en el desarrollo y pruebas metabólicas. Por ejemplo, imágenes de PET tratadas con AS/RP permiten un análisis cuantitativo de proteínas. Finalmente, los tratamientos rehabilitatorios son el campo de la medicina física y de rehabilitación. En estos la intervención es principalmente mecánica, pero a menudo está apoyada por medios robóticos y/o

virtuales. En estos casos, de nuevo el AS/RP es la encargada de generar la información de control, y procesar las mediciones digitales obtenidas a través del robot o del entorno virtual.

Una vez establecido el diagnóstico y cualquiera que sea el tipo de intervención, el AS/RP también puede contribuir a facilitar el pronóstico de recuperación de un paciente, e.g. en pacientes de dolor de espalda baja, y en general otras enfermedades crónicas.

7.3.4. Seguimiento

El seguimiento o monitoreo, ocurre una vez el paciente recibe el alta médica para abandonar el hospital, pero continúa bajo supervisión médica. De nuevo aparecen en escena variantes del AS/RP como el cómputo ubicuo (ahora aplicado al seguimiento). Aplicaciones como el envejecimiento saludable, el monitoreo de enfermedades crónicas degenerativas mediante las redes de sensores corporales, el procesamiento masivo de flujos de datos necesarios para el desarrollo de sistemas de alertas inteligentes ante emergencias de salud, etcétera.

Otra área fértil dentro del uso médico de AS/RP es la medicina del deporte. Esta rama de la medicina es un ejemplo claro de como el monitoreo de los atletas de élite pueden mejorar su rendimiento a partir del sensado de sus datos biológicos; mediciones de oxigenación muscular y cerebral, mediciones de su respuesta neuronal con EEG, perfilados iónicos a partir del sudor u otros fluidos corporales, análisis del movimiento biomecánico, medición de sinergias musculares con electromiogramas, establecimiento de tiempos de reacción y coordinación visuomotora son sólo algunas de las aportaciones del AS/RP a este área.

Cerramos esta sección dejando entrever como el AS/RP contribuye a la medicina por vías alternativas a las sugeridas relacionadas con la cura del paciente.

El AS/RP también está presente en aplicaciones de medición del desempeño del cirujano, la neuroergonomía quirúrgica o el análisis de las políticas de salud públicas.

7.4. Uso aplicado en biología

La biología es el estudio de la vida. Esto se extiende desde el conocimiento de la célula, el organismo vivo más elemental, sus organelos y su química incluyendo la genética, hasta el estudio de los ecosistemas a nivel macroscópico. Las aplicaciones industriales de la biología son contundentes; agrónoma, pecuaria incluyendo tanto la explotación ganadera, como la rama veterinaria, forestal, pesca y otras indirectas como la protección de la biomasa para fines atmosféricos o meramente para fines turísticos. Así pues la sección la hemos dividido en dos subsecciones. La primera describe el papel del AS/RP en los diferentes niveles de organización biológica.

- *A nivel celular* mediante el estudio de la biología molecular, genética y epigenética, metabolismo, y la reproducción celular.
- *A nivel organismo* mediante el estudio de los tejidos, órganos, sistemas y los propios organismos individuales, unicelulares y multicelulares.
- *A nivel ecosistema* estudiando las poblaciones, comunidades, propiamente los ecosistemas, hasta llegar al estudio de la ecosfera.

La segunda sección describe ejemplos de como el AS/RP impacta en la *explotación industrial* de la biología a través de las diferentes industrias.

7.4.1. A nivel celular

El estudio de las células, sus estructuras, su química, su genética y su metabolismo ha ido evolucionando desde estudios manuales a estudios guiados por técnicas computacionales de análisis de imágenes, primero microscópicas en sus diferentes modalidades, y más recientemente, nanoscópicas, donde las técnicas de AS, en particular de superresolución, han permitido esquivar el límite de difracción de Abbe. La proteómica ha visto como las técnicas de análisis de imágenes han incrementado su precisión. La detección de patrones de co-ocurrencia entre coctéles antiretrovirales y secuencias de mutaciones genéticas en el virus de la inmunodeficiencia humana (VIH) puede ser utilizado para mantener bajo control el avance de la enfermedad en un paciente con síndrome de inmunodeficiencia adquirida (SIDA). Incluso técnicas tan sencillas como las de falso color pueden suponer pasos críticos en la investigación a nivel celular, por ejemplo, permiten explorar el transporte interno involucrado en la digestión de macromoléculas.

7.4.2. A nivel organismo

En seres multicelulares las células que lo conforman se agrupan en tejidos, órganos y sistemas. Mucho del estudio de estos pasa por la identificación de estos a partir de las imágenes. En histología la segmentación de células individuales y/o tejidos permite establecer estimadores morfométricos, densidades, orientaciones o parámetros de crecimiento entre otros. Técnicas probabilistas de reconstrucción ayudan al oncólogo a identificar de forma visual la penetración de las células cancerígenas o la metástasis. Técnicas de seguimiento permiten observar el movimiento de células individuales en cultivos *in-vitro*.

7.4.3. A nivel ecosistema

El estudio de la biología a gran escala estudia poblaciones, comunidades, ecosistemas y a escala general la ecosfera. El AS/RP permite el monitoreo de poblaciones, el cual tiene implicaciones para el control de ecosistemas y la ayuda en la recuperación de especies. Así por ejemplo, las simulaciones de series temporales permiten estimar la evolución de poblaciones, a menudo en pares presa-predador. El estudio de las formas permite distinguir especies de peces. Es más, el estudio de las manchas de la piel de algunas especies permite la identificación de cada individuo, y en general, el RP es un elemento fundamental de la biometría animal. El estudio de las trazas de sensores de geolocalización en animales permite conocer las migraciones. El procesamiento de la señal de sonar permite el estudio del comportamiento de bancos de peces. Cambios en la tonalidad observada mediante espectroscopía en imágenes satelitales, permite determinar la especie forestal dominante en una región.

7.4.4. Explotación industrial

Ya hemos dejado entrever la importancia económica de la biología. Esta sección describe ejemplos más específicos de la penetración del AS/RP en estas industrias. En el uso veterinario, básicamente medicina animal, los usos y ventajas son relativamente similares a los obtenidos en la medicina humana. Algunos ejemplos son la detección de parásitos en animales domésticos, o su contribución en la oftalmología veterinaria.

La agricultura y la ganadería están viviendo una nueva revolución con la llegada de los métodos computacionales de AS/RP; control sistematizado de los invernaderos y granjas, control fitosanitario y de plagas, control inteligente de los sistemas de riego, agricultura inteligente incluyendo el sensado de imágenes

por drones o algo más tradicional por medios satelitales para la llamada agricultura de precisión, optimización de los procesos de empaque, trazabilidad y distribución logística de productos agrícolas y ganaderos, son solo algunos ejemplos en los que el uso intensivo de técnicas de AS/RP está marcando la diferencia.

La ingeniería de montes y forestal también se está viendo beneficiada de los métodos matemáticos de AS/RP orientados al control de plagas e incendios, imágenes satelitales para salud y medición de la masa forestal, para la detección de usos indebidos e.g. detección ilegal de tala de árboles, etc.

En pesca, las actuales explotaciones marinas o de piscifactoría a gran escala están totalmente supeditadas a la disposición de soluciones de AS/RP. Por ejemplo, en la pesca de alta mar, el procesamiento de la señal sonar permite la detección remota de los bancos de pesca. Las piscifactorías obtienen beneficios similares a los de la ganadería anteriormente mencionados; monitoreos de oxigenación y calidad del agua de las piscinas, conteos de población y estimación de tamaño, etc.

7.5. Vanguardia del campo en México

Algunos de los ejemplos que hemos citado arriba pertenecen a grupos Mexicanos. Aquí mencionamos algunos de los grupos más activos del área y destacamos su línea de trabajo. Sin un orden en particular;

- En la *Universidad Nacional Autónoma de México* varios grupos destacan en sus investigaciones relacionadas con las soluciones de AS/RP en aplicaciones en biología y medicina. Destacaremos, por ninguna razón en particular los laboratorios de Reconocimiento de Patrones y de Pro-

cesamiento de Imágenes del Instituto de Investigaciones en Matemáticas Aplicadas y en Sistemas (IIMAS) quienes han presentado recientemente trabajos en revistas relacionados con la detección del potencial P₃₀₀ (útil para las interfaces cerebro computadora) o el análisis de imágenes de retina. Históricamente han trabajado también imágenes de ultrasonido.

- Quizás uno de los grupos más veteranos en México es de la *Universidad Autónoma Metropolitana* en Iztapalapa, quienes realizan investigaciones en mapeo cerebral e interfaces cerebro computadoras así como otras bioseñales.
- El grupo del *Centro de Investigaciones en Matemáticas (CIMAT)* tiene una actividad destacada en neuroimágenes; tanto en electroencefalografía, como en resonancia magnética donde han explorado problemas de segmentación y tractografía entre otros.
- El departamento de computación del *Centro de Investigación y de Estudios Avanzados (CINVESTAV)* del *Instituto Politécnico Nacional (IPN)* no tiene como tal una línea específica de investigación en AS/RP en biomedicina. No obstante, sus investigaciones han tenido ramificaciones interesantes en este área.
- Por último, nuestro propio laboratorio de bioseñales y computación médica en el *Instituto Nacional de Astrofísica, Óptica y Electrónica (INAOE)* tiene actividad en varias áreas de la imagenología biomédica; por ejemplo la detección de características a partir del registro electroencefalográfico para consideraciones en epilepsia, clasificación de leucemia, pié diabético, interfaces cerebro computadora, comprensión del llanto de bebe, neuroimágenes o neurorehabilitación.

Por supuesto, otros grupos operan en México en temas relacionados e.g. el *Instituto Tecnológico de Sonora* trabaja en áreas de señales temporales con aplicación en ecología, o el *Centro de Investigación Científica y de Educación Superior de Ensenada (CICESE)* quienes, aunque alojado en el departamento de óptica en lugar del de computación, el grupo de procesamiento de imágenes, han presentado contribuciones claramente de corte computacional en biología, pero esperamos haber proporcionado un breve resumen de los grupos más destacados en el área. Inevitablemente, habrá grupos que no estemos mencionando por espacio o por nuestra propia ignorancia; para ellos una disculpa por adelantado.

7.6. Conclusiones

La contribución del AS/RP a las ciencias biomédicas es formidable; tanto desde el punto de vista de aportaciones teóricas como en la fabricación de tecnologías finales directamente disponibles a la sociedad. Este capítulo sólo ha mostrado la punta del iceberg. En el, hemos cubierto tanto técnicas como aplicaciones. No tenemos constancia de que se haya medido el impacto económico del AS/RP en su aplicación a las áreas de biología y medicina. No obstante, habiendo esbozado muchas de las aplicaciones actuales y las venideras, no parece osado pensar que dicho impacto podría medirse en los cientos o miles de millones de dólares a nivel mundial.

Referencias

Chatfield, C. (2004). *The analysis of time series: An introduction*. Boca Ratón: Chapman & Hall/CRC.

- Herrera-Vega, J., Treviño Palacios, C. G., y Orihuela-Espina, F. (2017). Neuroimaging with functional near infrared spectroscopy: from formation to interpretation. *Infrared Physics & Technology*, 85, 225-237. doi: 10.1016/j.infrared.2017.06.011
- Moher, D., Hopewell, S., Schulz, K. F., Montori, V., Gotzsche, P. C., Devereaux, P. J., ... Altman, D. G. (2010). Consort 2010 explanation and elaboration: uupdate guidelines for reporting parallel group randomised trials. *British Medical Journal*, 340, c869 (28 pp). doi: 10.1136/bmj.c869
- Najarian, K., y Splinter, R. (2006). *Biomedical signal and image processing*. Boca Raton. Florida. USA: CRC press. Taylor and Francis.
- Proakis, J. G., y Manolakis, D. G. (2007). *Digital signal processing: principles, algorithms and applications* (4th Ed. ed.). New Jersey: Pearson-Prentice Hall.
- Rangayyan, R. M. (2015). *Biomedical signal analysis* (Vol. 33). John Wiley & Sons.
- Shmueli, G., y cols. (2010). To explain or to predict? *Statistical science*, 25(3), 289-310.
- Theis, F. J., y Meyer-Bäse, A. (2010). *Biomedical signal analysis: Contemporary methods and applications*. MIT Press.

Capítulo 8

Interfaces Cerebro-Computadora

EDUARDO RIVAS-POSADA, ITCH

MARIO IGNACIO CHACÓN MURGUÍA, ITCH

JUAN ALBERTO RAMÍREZ QUINTANA, ITCH

El presente capítulo tiene como intención presentar al lector la situación actual de las interfaces cerebro-computadora en México y en el mundo, comenzando por describir el concepto de interfaz cerebro-computadora, así como las estrategias para la adquisición de las señales cerebrales y los modelos matemáticos utilizados. De igual manera, se exponen las aplicaciones más destacadas hasta el momento, tanto internacionales como nacionales. Finalmente, la sección de conclusiones muestra las áreas de oportunidad de AS/RP en México.

8.1. Interfaz cerebro-computadora

En esta sección se aborda el concepto de interfaz cerebro-computadora, así como las principales estrategias para la adquisición de las señales cerebrales y algunos modelos matemáticos para su implementación.

8.1.1. Concepto y adquisición de señales cerebrales

Una interfaz cerebro-computadora o interfaz BCI (*Brain Computer Interface*, por sus siglas en inglés) tiene como fin controlar acciones virtuales o dispositivos electrónicos externos –como una prótesis o un robot– utilizando, como entrada al sistema, señales cerebrales.

Al igual que cualquier sistema de clasificación, un sistema BCI incluye actividades de adquisición de las señales cerebrales, pre-procesamiento –a fin de mejorar la señal adquirida–, extracción de características, clasificación de los datos obtenidos y finalmente, el sistema de control que permite la ejecución de la acción deseada.

Como primer punto, para la adquisición de las señales cerebrales se pueden aplicar diversos métodos que se clasifican en *invasivos* y *no invasivos*. En un sistema *invasivo* se colocan electrodos directamente en la corteza cerebral por medio de un proceso quirúrgico para registrar las señales cerebrales. En cambio, un sistema *no invasivo* permite capturar las señales cerebrales mediante técnicas que no requieren de cirugía. La más utilizada en BCI es el electroencefalograma o EEG, que permite obtener la actividad eléctrica cerebral mediante diferencias de voltaje que se producen en los electrodos ubicados en el cuero cabelludo y una referencia. Dichos electrodos se posicionan en zonas específicas definidas

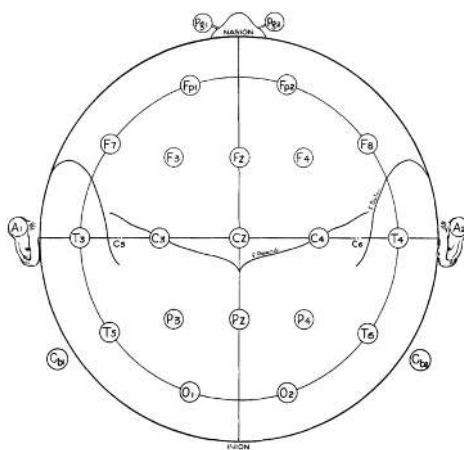


Figura 8.1: Posicionamiento de electrodos de acuerdo con el sistema Internacional 10-20.

por el Sistema Internacional o sistema 10-20 (Jasper, 1958), como se observa en la figura 8.1.

El Sistema Internacional 10-20 es denominado así debido a que los electrodos están espaciados entre el 10 % y el 20 % de la distancia total entre puntos reconocibles por el cráneo.

Asimismo, existe una nomenclatura que permite identificar cada electrodo y se rige de acuerdo con la región cerebral donde éste se ubica y una numeración que va de menor a mayor, desde zonas anteriores hacia zonas posteriores, con números impares del lado izquierdo y números pares del lado derecho. La tabla 8.1 reporta el área cerebral y la nomenclatura a la que pertenece dicha zona (Talamillo, 2011).

Esta nomenclatura y el Sistema 10-20 permiten que los diversos equipos electrónicos utilizados para la adquisición de las señales cerebrales sean posicionados

Tabla 8.1: Nomenclatura de los electrodos de acuerdo con el Sistema 10-20.

Zona cerebral	Hemisferio izquierdo	Línea Media	Hemisferio derecho
Frontopolar	FP ₁	-	FP ₂
Frontal	F ₃	F _z	F ₂
Frontotemporal	F ₇ C ₃	C _z	F ₈ C ₄
Temporal Medio y Parietal	T ₃ P ₃	P _z	T ₄ P ₄
Temporal posterior y Occipital	T ₅ O ₁	-	T ₆ O ₂

correctamente, permitiendo que se registre un EEG adecuado para su posterior interpretación.

Cabe señalar que existen equipos médicos (como *Medical Computer Systems* o *g.tec medical engineering*) especializados en la toma de muestras para revisar la actividad cerebral de un paciente y pueden ser muy costosos. Sin embargo, hoy en día la mayoría de las aplicaciones BCI recurren a las bondades de los avances tecnológicos de menor costo, dando paso a que nuevos sistemas para la adquisición de señales cerebrales –como diademas o cascos con electrodos (OpenBCI®, EMOTIV®, NeuroSky®, muse®)– sean implementados en interfaces cerebro-computadora.

Por otra parte, otro punto a considerar cuando se va a registrar la actividad eléctrica cerebral mediante EEG es el tipo de electrodos. Existen *electrodos húmedos* y *electrodos secos*. Los primeros requieren de un gel electrolítico u otras soluciones para facilitar la transmisión de la señal eléctrica cerebral, a través del cuero cabelludo al electrodo. Los *electrodos secos*, al contrario, no necesitan de

ninguna sustancia conductiva, lo que los hace más prácticos y reduce el tiempo de preparación del usuario de una interfaz BCI.

Una vez que se han colocado los electrodos sobre el cuero cabelludo, en el caso de un sistema *no invasivo* y mediante un EEG, se está en condiciones de analizar la actividad cerebral. Las principales ondas cerebrales se muestran en la tabla 8.2, donde se proporciona información del rango de frecuencia y el estado correspondiente a analizar en cada una de ellas.

Con dichas ondas cerebrales se pueden diseñar aplicaciones BCI que permitan, por ejemplo, la detección del nivel de concentración de un sujeto o si se encuentra en estado somnoliento.

Algunas de las señales comúnmente utilizadas para implementar una interfaz BCI son los Potenciales Evocados Visuales o VEP (por sus siglas en inglés *Visually Evoked Potentials*), los Potenciales Relacionados a Eventos o ERP (por sus siglas en inglés *Event Related Potentials*) y los ritmos sensoriomotores (que comprenden las ondas Mu y Beta) para detectar visualizaciones motoras.

Como estas señales cerebrales son el principio de funcionamiento del sistema BCI, deben ser sometidas a un pre-procesamiento para acondicionarlas y eliminar elementos indeseados. Después se continuará con las actividades de clasificación para concluir con el proceso de análisis y toma de decisiones y/o actuadores. Para una mejor comprensión del proceso anterior, en la siguiente sección se explica matemáticamente todo el proceso de clasificación, incluyendo los principales modelos utilizados para implementar una aplicación basada en BCI.

Tabla 8.2: Principales ondas cerebrales.

Onda	Rango de frecuencia	Estado
Beta	12 – 30 Hz	Estado de vigilia normal de la conciencia, cognición, enfoque, alerta, concentración.
Alfa	7.5 – 12 Hz	Relajación profunda, creatividad, visualización, ligera meditación con los ojos cerrados.
Mu	7.5 – 12 Hz	Aparecen en el área sensoriomotora del cerebro cuando el cuerpo está físicamente relajado.
Theta	4 -7.5 Hz	Sueño ligero, imágenes vívidas en la imaginación, meditación profunda.
Delta	0.1 – 4 Hz	Sueño profundo.

8.1.2. Modelos matemáticos

En esta sección se muestra el proceso que va desde la adquisición de las señales hasta la etapa de control, utilizando ecuaciones y modelos matemáticos que ayuden a su comprensión.

Como se abordó anteriormente, la etapa de adquisición comienza con la captura de actividad eléctrica cerebral. Para facilitar el análisis matemático se considerará una señal de un solo canal (un electrodo) y se asociará con $x(t)$, en el dominio del tiempo continuo al tratarse de una señal analógica.

Dentro de esta etapa es necesario realizar una conversión analógica a digital que incluye el muestreo, la cuantificación y la codificación (Proakis, Manolakis, y Castro, 1997), ya que eso permitirá realizar un acondicionamiento de la señal

y continuar las actividades de clasificación. En la figura 8.2 se observa el proceso de conversión análogo-digital.

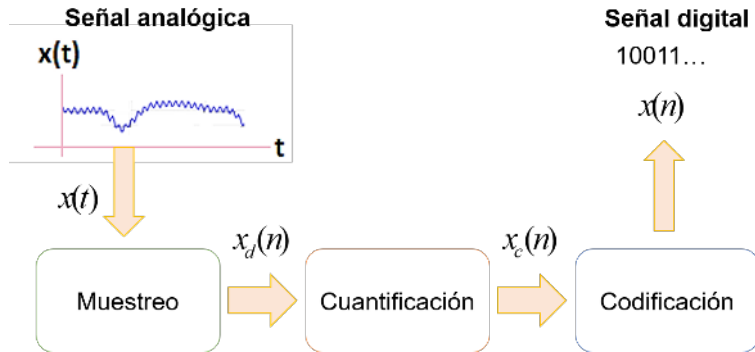


Figura 8.2: Proceso básico de conversión analógica a digital.

El *muestreo*, como su nombre lo dice, consiste en la toma de “muestras” para que la señal de entrada analógica, continua en el tiempo, sea convertida en una señal discreta en el tiempo. De forma que si $x(t)$, la señal cerebral, es la entrada del sistema de muestreo, la salida será $x_d(nT) \equiv x_d(n)$, donde T es el intervalo de muestreo.

La *cuantificación* consiste en la conversión de una señal de valores continuos tomados en instantes discretos de tiempo $x_d(n)$, en una señal con valores discretos en instantes de tiempo discreto, es decir, $x_c(n)$.

Finalmente, la *codificación* es asignar una secuencia binaria de n -bits a cada valor discreto $x_c(n)$, tomando en cuenta la resolución del convertidor análogo-digital o ADC (por sus siglas en inglés *Analog-to-Digital Converter*).

Así, a partir de $x(t)$, se obtiene una señal digital $x(n)$ correspondiente a los valores discretos en tiempo y amplitud de la actividad eléctrica cerebral que se va registrando. Con esto termina la etapa de adquisición y se da inicio al pre-procesamiento.

El pre-procesamiento consiste, básicamente, en el acondicionamiento de la señal digital para poder aplicar la extracción de características. Esto se logra mediante filtros digitales que eliminen ruido, artefactos indeseados y permitan delimitar al rango de frecuencias que se está buscando.

Hay una gran variedad de filtros digitales que se clasifican de acuerdo con su respuesta al impulso, ya sea finita (*FIR* del inglés *Finite Impulse Response*) o infinita (*IIR* del inglés *Infinite Impulse Response*).

Para obtener la salida $y(t)$ de un filtro digital tipo *FIR* se necesita conocer la entrada actual y entradas pasadas, como se muestra en la ecuación 8.1.

$$y(n) = \sum_{k=0}^M b_k x(n-k) \quad (8.1)$$

Donde k representa un índice que toma valores enteros positivos desde 0 hasta M ; $x(n-k)$ representa la entrada actual y entradas pasadas; y b_k son los coeficientes constantes correspondientes a la entrada del filtro.

En cambio, si se desea conocer la salida $y(n)$ de un filtro digital tipo *IIR* se necesitan conocer los valores de la entrada actual y entradas pasadas, así como de salidas pasadas. Esto hace que sea recursivo, como se observa en la ecuación 8.2.

$$y(n) = \sum_{k=0}^M b_k x(n-k) - \sum_{k=1}^N a_k y(n-k) \quad (8.2)$$

En donde la sumatoria de los elementos de la entrada van desde 0 hasta M y los de la salida desde 1 hasta N , tomando k valores enteros positivos. Los coeficientes de la entrada del filtro son expresados mediante b_k y los de la salida del filtro con a_k .

Al implementar estas expresiones matemáticas en el sistema BCI se realiza una actividad importante dentro de la etapa de pre-procesamiento que permitirá limpiar la señal y dejarla en las mejores condiciones para pasar a la siguiente fase, que es la extracción de características.

La extracción de características representa uno de los retos más grandes dentro de las actividades de clasificación. Ya que, si se tienen características sobresalientes, la clasificación resultará más evidente.

Existen muchas formas de extraer características, desde el manejo de datos en el dominio del tiempo mediante información estadística; en el dominio de la frecuencia para detectar *VEP* o *ERP*; o incluso en el dominio espacial para visualizaciones motoras.

Si se busca extraer características en el dominio temporal, herramientas estadísticas como la media, mediana, varianza, desviación estándar de la muestra serán de gran utilidad.

La media de una muestra es el promedio numérico que es representado mediante la ecuación 8.3 y es muy útil para observar el comportamiento general de la señal.

$$\bar{x} = \sum_{i=0}^n \frac{x_i}{n} = \frac{x_1 + x_2 + x_3 + \dots + x_n}{n} \quad (8.3)$$

Donde $\bar{x}$ es el valor de la media, x_i son las muestras que van desde la muestra 1 hasta la muestra n (número total de muestras) en enteros positivos.

Otra medida estadística es la mediana que se representa con $\tilde{x}$ y se obtiene mediante la expresión 8.4. La mediana refleja la tendencia central de la muestra

sin verse afectada por los valores extremos, ya que las muestras deben ser acomodadas en orden de magnitud creciente.

$$\tilde{x} = \begin{cases} x_{(n+1)/2}, & \text{si } n \text{ es impar,} \\ \frac{1}{2}(x_{n/2} + x_{(n/2)+1}), & \text{si } n \text{ es par.} \end{cases} \quad (8.4)$$

Además de las medidas estadísticas que permiten observar la tendencia central, existen otras que ayudan a calcular el grado de dispersión o variabilidad. Tal es el caso de la varianza de la muestra, que se representa σ y se expresa mediante la ecuación 8.5.

$$\sigma = \sum_{i=0}^n \frac{(x_i - \bar{x})^2}{n - 1} \quad (8.5)$$

De manera similar, si se calcula la raíz cuadrada de la expresión anterior, es decir, de la varianza, se obtendrá la desviación estándar de la muestra denotada por s como se observa en la ecuación 8.6.

$$s = \sqrt{\sigma} \quad (8.6)$$

Por otra parte, además de los elementos estadísticos, hay otras características que pueden apreciarse mejor si se maneja el dominio de la frecuencia. Para ello se utiliza, por lo general, la Transformada Discreta de Fourier que se muestra en la ecuación 8.7.

$$X(k) = \sum_{n=0}^N x(n) e^{-j2\pi \frac{k}{N} n} \quad (8.7)$$

Donde la entrada $x(n)$ en el dominio del tiempo discreto se multiplica por el factor $e^{-j2\pi \frac{k}{N} n}$, siendo k el índice de frecuencia y n la muestra actual. Así, la sumatoria da como resultado una señal en el dominio de la frecuencia $X(k)$.

Esta transformación permitirá extraer ciertas características en el dominio de la frecuencia que en el dominio del tiempo discreto sería más difícil su interpretación. Es por esta razón que, al cambiar de un dominio a otro, determinada información resulta más evidente y la etapa de clasificación será menos complicada.

Lo siguiente, una vez que se cuenta con las características a clasificar, es definir las clases del sistema BCI. Para ello se utilizan modelos matemáticos más complejos que permiten agrupar la información con base en los datos obtenidos de la etapa anterior. Algunos de estos modelos utilizan Máquinas de Vectores de Soporte o más conocidas como SVM (por sus siglas en inglés *Support Vector Machine*), redes neuronales artificiales, filtros adaptivos, *Common spatial pattern* (CSP, muy utilizado en visualizaciones motoras) o incluso herramientas como el Teorema de Bayes que recurre a la probabilidad.

Con dichos modelos matemáticos de clasificación se generarán comandos que son reinterpretados dentro de la etapa de control. Así, la señal cerebral original, una vez que ha sufrido todo el proceso de clasificación, es utilizada para decirle al sistema qué acción virtual generar o si un brazo robótico se debe mover hacia la izquierda o la derecha –en el caso de un sistema BCI con un dispositivo electrónico externo–.

Sin embargo, para no hacer suposiciones y ver realmente los sistemas BCI que se han implementado, la siguiente sección muestra aplicaciones a nivel internacional y nacional, categorizadas de acuerdo con el objetivo principal del sistema.

8.2. Aplicaciones

Una vez que se ha comprendido la parte teórica, se está en condiciones de explorar el estado del arte actual, en el mundo y en nuestro país, mostrando las áreas de mayor interés.

Las aplicaciones que se muestran a continuación, pero que no son las únicas, están categorizadas en tres áreas: médica, entretenimiento y cognitiva.

8.2.1. Médica

Dentro de esta categoría se encuentran aplicaciones que han sido enfocadas principalmente a personas que han sufrido alguna discapacidad, como la Esclerosis Lateral Amiotrófica, parálisis o que se encuentran en etapas de recuperación después de un ataque cerebrovascular. De forma que el objetivo principal será implementar una interfaz BCI que ayude en las labores diarias, de comunicación, de terapia para recuperación de movilidad.

El trabajo realizado por B. E. Olivas Padilla y M. Chacón Murguía en (Olivas-Padilla, 2018) propone un sistema BCI que utiliza visualizaciones motoras para manipulación de un ambiente de terapia virtual. En este caso, este tipo de sistemas pueden ser utilizados con personas que tienen alguna discapacidad motora, ya que la visualización motora genera actividad cerebral que puede ser registrada y almacenada sin mover propiamente una extremidad del cuerpo humano, simplemente con visualizarlo. Así, una persona con discapacidad motora puede ser entrenada y recibir retroalimentación visual a través de terapias virtuales llamadas TerapiaTec “Traffic Racer” y “Bate-Rapia Tec” (Olivas-Padilla, 2018), como se observa en la figura 8.3.

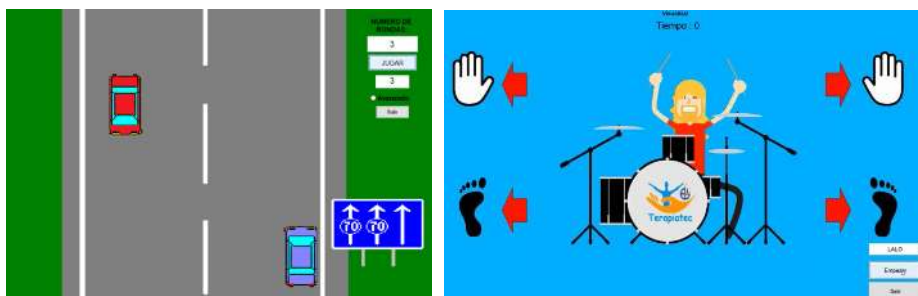


Figura 8.3: Terapias virtuales TerapiaTec “Traffic Racer” y “Bate-Rapia Tec” con visualización motora implementadas en un sistema BCI.

Otro trabajo, realizado por L. Madrid y J. Ramírez en (Madrid-Herrera, 2018), presenta un algoritmo para extracción de características de señal P300 para ser implementado en un sistema BCI. Una señal P300 es un tipo de ERP, los cuales son cambios en la actividad cerebral que se miden en el dominio del tiempo cuando se es expuesto a un evento aleatorio, y al tratarse de una onda P300 significa que se presenta una deflexión positiva alrededor de los 250-400 ms después del estímulo.

Este tipo de onda es muy común en los sistemas BCI y en este caso se utilizó un deletreador de Donchin modificado (Madrid-Herrera, 2018), como se muestra en la figura 8.4, para generar las señales P300 de entrenamiento y de prueba. Este deletreador permite que personas con parálisis o muy poca capacidad de movimiento, puedan establecer comunicación.

Además, se utilizaron Redes Neuronales Convolucionales para la clasificación y se presenta un escenario virtual para la selección de objetos, como se muestra en la figura 8.5.

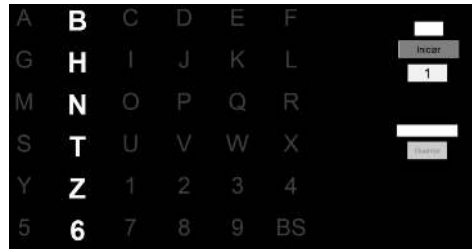


Figura 8.4: Deletreador de Donchin modificado.

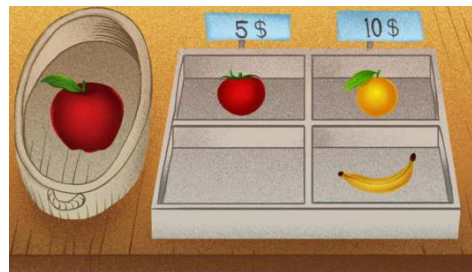


Figura 8.5: Escenario de frutas utilizado para la selección de objetos mediante onda P300 en un sistema BCI.

En este escenario parpadean las frutas y el usuario, al observar determinada fruta, genera la onda P300 en un tiempo específico, lo cual es utilizado para determinar qué fruta estaba viendo y depositarla en la canasta.

8.2.2. Entretenimiento

Esta segunda área está enfocada a generar interfaces BCI que permitan aumentar el nivel de emoción, control y jugabilidad, principalmente en videojuegos.

Una de los sistemas BCI más famosos es el implementado por la compañía “Neurable” (Neurable, 2018), que busca llevar los videojuegos a otro nivel



Figura 8.6: Prototipo de Neurable para implementar un sistema BCI en conjunto con lentes de realidad virtual.

de jugabilidad, utilizando realidad virtual en conjunto con un sistema BCI que detecta y clasifica señales cerebrales, como la onda P300, mediante aprendizaje de máquina. En la figura 8.6 se muestra el hardware completo del sistema (Neurable, 2018), que consiste en unos lentes de realidad virtual HTC *Vive* y la diadema de Neurable que contiene los electrodos para la adquisición de señales EEG.

Asimismo, en (Kerous, Skola, y Liarokapis, 2018) B. Kerous, F. Skola, y F. Liarokapis presentan un informe general de los avances en sistemas BCI con aplicación en videojuegos. Señalan que el 28 % de los artículos reportan el uso de visualizaciones motoras para implementar el sistema BCI; 20 % utilizan VEP de estado estable; 9.3 % manejan ERP mediante la onda P300 y la gran mayoría, 38.7 %, recurren a la neuroretroalimentación.



Figura 8.7: Lentes “Lowdown Focus” de la compañía SMITH.

8.2.3. Cognitiva

Finalmente, en esta categoría se encuentran trabajos que han ayudado a generar conocimiento acerca del comportamiento del cerebro humano al verse involucrado en diferentes situaciones como cambios emocionales, toma de decisiones, resolución de problemas, concentración.

Un producto internacional famosos son los lentes “Lowdown Focus” de la compañía SMITH (*Lowdown Focus*, SMITH, 2018). Estos lentes tienen incorporados electrodos en la zona temporal-parietal que permiten registrar la actividad cerebral para ser monitoreada mediante una aplicación para celulares, como se observa en la figura 8.7.

La página oficial señala que el uso de estos lentes y la aplicación de retroalimentación “Smith Focus” permiten observar y controlar el nivel de atención, así como combatir con factores estresantes diarios o estar concentrado en el cumplimiento de metas. También señalan que ayudan a controlar el nivel de ansiedad y aumentar la concentración en el aprendizaje.

De manera que este sistema BCI es utilizado por personas que buscan alto rendimiento, como deportistas o gente que necesita trabajar su nivel de atención y concentración en las labores diarias.

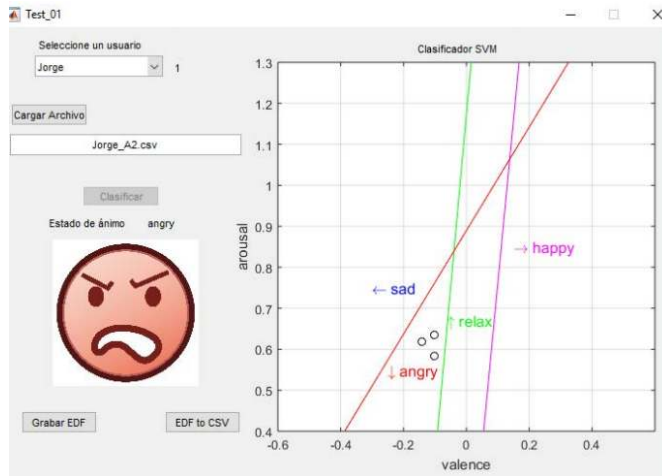


Figura 8.8: Interfaz de usuario para reconocimiento de emociones.

Otro sistema BCI es el implementado por O. Rivera y M. Chacón Murguía en (Rivera-Calderón y Chacón-Murguía, 2017), donde presentan un método para el análisis de señales EEG para el reconocimiento de cuatro emociones: enojo, alegría, relajación y tristeza. Las emociones son incentivadas mediante la proyección de estímulos audiovisuales y se generó una interfaz de usuario para mostrar la emoción detectada, tal como aparecen en la figura 8.8.

Esto se realiza mediante un análisis en las ondas alfa y beta y se utiliza SVM como clasificador.

8.3. Conclusiones

Con la situación actual de las BCIs, que se presentó en la sección anterior, se puede observar un panorama con muchos avances importantes dentro del AS/RP en México. Sin embargo, siempre hay áreas de oportunidad que permiten crecer

y fortalecer más el avance científico y tecnológico para el desarrollo de prototipos y aplicaciones en beneficio de la sociedad.

Se observan, principalmente, dos áreas de oportunidad: una que abarca el enfoque técnico y otra que consiste en el desarrollo de aplicaciones.

En cuanto al enfoque técnico, se cree que una de las mayores limitantes ha sido que falta diseñar e implementar nuevos algoritmos, que sean más robustos y permitan un mayor porcentaje certero de clasificación. Se sabe que la actividad cerebral es única en cada persona, por lo que cada usuario debe ser entrenado; pero si se tiene un buen algoritmo que permita extraer características y clasificar adecuadamente dicha actividad cerebral, la aplicación BCI final resultará más práctica, confiable y estable.

Por otra parte, hay muchos tipos de aplicaciones de sistemas BCI, pero la mayoría se disgrega entre las enfocadas a fines médicos o de salud –para personas con discapacidades o en recuperación– y las que tienen por objetivo personas sanas –con fines de entretenimiento o cognitivos–. Con esto se observa que un área de oportunidad es innovar en el tipo de aplicaciones donde se contemple a cualquier tipo de persona. Es decir, buscar un punto de equilibrio que contribuya a la situación actual de cada individuo, ayudándolos en sus labores diarias y promoviendo que las interfaces cerebro-computadora den un paso más en este largo camino que aún queda por recorrer.

Referencias

Jasper, H. (1958). Report of the committee on methods of clinical examination in electroencephalography: 1957. *Electroencephalography and Clinical Neurophysiology*, 10(2), 370 - 375.

- Kerous, B., Skola, F., y Liarokapis, F. (2018, 01 de Jun). Eeg-based bci and video games: a progress report. *Virtual Reality*, 22(2), 119-135.
- Lowdown focus, Smith* [En línea]. (2018). <https://www.smithoptics.com/us/lowdownfocus>. (Consultado: 2018-08-16)
- Madrid-Herrera, L. (2018). *Diseño de algoritmo para la extracción de características en un sistema de interfaz cerebro-computadora* (Tesis de Maestría). Instituto Tecnológico de Chihuahua.
- Neurable* [En línea]. (2018). <http://www.neurable.com>. (Consultado: 2018-08-16)
- Olivas-Padilla, B. E. (2018). *Aplicación de visualización motora para manipulación de un ambiente de terapia virtual* (Tesis de Maestría). Instituto Tecnológico de Chihuahua.
- Proakis, J., Manolakis, D., y Castro, J. (1997). *Tratamiento digital de señales*. Pearson Educación.
- Rivera-Calderón, O. G., y Chacón-Murguía, M. I. (2017). Reconocimiento de emociones mediante el análisis de señales EEG. *Congreso Internacional de Investigación Científica Multidisciplinaria*, 9(1), 58 - 72. Descargado de <http://www.chi.itesm.mx/icm/memorias2017/Ingenieria.pdf>
- Talamillo, T. (2011). Manual básico para enfermeros en electroencefalografía. *Enfermería docente*, 94, 29-33.

El reconocimiento de patrones y su aplicación a las señales digitales,

se terminó el 30 septiembre de 2019.

A partir del 1.º de diciembre de 2019 está disponible en formato PDF en la página de la

Academia Mexicana de Computación:

<http://www.amexcomp.mx>