
Resumen de la tesis:

Generación de Conocimiento basado en Aprendizaje Automático y Aplicación en Diferentes Sectores

AIPAKA

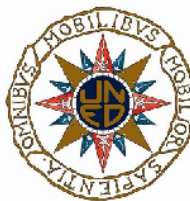
(Artificial Intelligence in a Process for Automated
Knowledge Acquisition and Applications)

Autor:

Fernando Pavón Pérez

Directores:

Dr. Jesús A. Vega Sánchez
Dr. Sebastián Dormido Canto



Madrid, Febrero de 2016

Resumen de la Tesis:

**Generación de Conocimiento basado en
Aprendizaje Automático y Aplicación en
Diferentes Sectores**

AIPAKA

**(Artificial Intelligence in a Process for Automated
Knowledge Acquisition and Applications)**

Fernando Pavón Pérez

Directores de Tesis:

Dr. Jesús A. Vega y Dr. Sebastián Dormido Canto.

Programa de Doctorado en Ingeniería de Sistemas y Control
Departamento de Informática y Automática
Escuela Técnica Superior de Ingeniería Informática (ETSI)
Universidad Nacional de Educación a Distancia (UNED)

Laboratorio Nacional de Fusión
Centro de Investigaciones Energéticas,
Medioambientales y Tecnológicas (CIEMAT)

Madrid, Febrero 2016

Resumen

El presente trabajo está basado en la utilización de técnicas del campo de la Inteligencia Artificial y más concretamente del Aprendizaje Automático, para la resolución de problemas complejos a los cuales se están enfrentando en la actualidad compañías de cualquier sector y muchos centros de investigación. Los problemas referidos se pueden definir por las siguientes características: disponibilidad de históricos, modelado complejo por la existencia de relaciones no lineales entre las variables, dinamismo, objetivos frecuentemente multivariantes y desconocimiento de las leyes fundamentales que gobiernan el sistema modelado o al que se refiere la información guardada.

En esta tesis, se definirá un procedimiento, denominado AIPAKA (*Artificial Intelligence in a Process for Automated Knowledge Acquisition and Applications*), que permitirá utilizar las mismas técnicas de Aprendizaje Automático a diferentes problemas. Usando para ello los siguientes elementos: clasificación de los problemas tratados, un conjunto de técnicas base para la construcción de la solución y un procedimiento para la creación y validación de los sistemas creados.

El objetivo fundamental es proveer de una herramienta adecuada para la solución de problemas complejos, basada en modelos predictivos, los cuales son obtenidos por inferencia automática de conocimiento a partir de históricos y de los datos producidos en tiempo real, mediante el uso de técnicas de Inteligencia Artificial.

El área de conocimiento donde se encuadra el presente trabajo puede ser lo que tradicionalmente se ha llamado la minería de datos o lo que es lo mismo: responder a las cuestiones de cómo descubrir y utilizar el conocimiento implícito en grandes colecciones de datos.

En los últimos tiempos se ha acuñado un nuevo término: macrodatos o inteligencia de los datos (del término en inglés Big Data). AIPAKA y las técnicas usadas pretenden ser una potente herramienta para la extracción y

uso del conocimiento a partir de los macrodatos. En estas ingentes cantidades de información es donde tiene todo el sentido la implementación de técnicas avanzadas de modelado predictivo, segmentación y análisis de patrones de comportamiento.

Índice general

1. Introducción	1
2. Estado del Arte	5
2.1. Introducción a la Inteligencia Artificial (IA) y a la Minería de Datos	5
2.1.1. Concepto y breve historia de la Inteligencia Artificial .	6
2.2. Necesidades de la industria y de la ciencia	18
3. AIPAKA	27
3.1. Introducción y Objetivos	27
4. Casos de Aplicación	31
4.1. Predicción de la demanda eléctrica a corto plazo	33
4.2. Segmentación en dispositivos experimentales complejos	35
4.2.1. Introducción	35
4.2.2. Tipo de problema	37
4.2.3. Datos	38
4.2.4. Modelización	39
4.2.5. Validación	39
4.2.5.1. Primeros estudios de las señales usando THE-FUMO	42
4.2.5.2. Sustitución de la primera capa de MSOM por una red OJA	46

4.2.5.3. Primeras conclusiones obtenidas con THE- FUMO	50
5. Conclusiones y desarrollos futuros	56

Índice de tablas

2.1. Macrodatos vs “pequeños datos”.	26
4.1. Señales muestreadas del dispositivo experimental de fusión. .	40
4.2. Algunas configuraciones y resultados de la SOM de la segunda capa. Se han señalado en negrita aquellas elegidas para estudiar las transiciones L-H.	45
4.3. Relación de señales y SOMs usadas en la primera capa de MSOM, correspondientes a la segunda capa de la MSOM com- puesta por la SOM 312. Para cada SOM y señal, la capacidad para discriminar los dos segmentos encontrados en la SOM de salida se cualificado con <i>Mala, Regular, Buena, Muy Buena</i> <i>y Excelente</i>	48
4.4. Vector de pesos de las neuronas 22, 38 y 39 de la SOM de salida 22105. Cada variables del vector de pesos corresponde a una de las 60 componentes principales que componen la salida de la red OJA (102). Se han señalado en negrita aquellos que pesos que difieren más entre las tres neuronas.	54
4.5. Señales con valores de θ más elevados en las componentes principales que influyen en activar la neuronas de la SOM de salida donde se concentran las transiciones L-H. Se indica el número de veces que esa seña ha resultado relevante para alguna de las estructuras de THEFUMO con red OJA que se han probado.	55

4.6. Comparación de algunas señales halladas como relevantes en THEFUMO con arquitectura MSOM y si coinciden con las halladas con arquitectura OJA.	55
---	----

Índice de figuras

4.1. Señales en una ventana de 300ms centradas en la transición L-H (en $x=150$). Se han representado la misma señal para varias descargas. Las señales que se representan son: <i>Density</i> , <i>TOG</i> , <i>TI</i> , <i>LID4</i> , <i>FDWDT</i> , <i>Bndiam</i> (tabla 4.1).	41
4.2. Formas de onda de dos neuronas de las SOMs de la primera capa de la MSOM, para el análisis de las señales Bndiam, LID4 y Rad.	43
4.3. Neuronas activadas en las SOMs del primer nivel. El color rojo indica que activan el grupo 1 de la SOM de salida (segundo nivel). La coloreadas en rojo son para aquellas neuronas que activan el grupo 2 de la SOM de salida. El color verde indica que existe “mezcla” para los dos segmentos encontrados en la capa de salida de la MSOM. El número en cada una de las celdas indican el número de transiciones L-H que han activado esa neurona.	47
4.4. Neuronas activadas por las transiciones L-H en la SOM de salida (22105), usando como capa de entrada un red OJA (102) de 60 neuronas.	49
4.5. Valores de thetas para una componente principal de la red OJA 102.	53

Acrónimos y Glosario

AGI *Artificial General Intelligence.* 14

AI *Artificial Intelligence.* 14

AIPAKA *Artificial Intelligence in a Process for Automated Knowledge Acquisition and Applications.* 22, 28, 30–32, 37, 57, 59–63, 65, 66

API *Application Programming Interface.* 64

BnDiam *Beta normalised with respect to the diamagnetic energy.* 56

Bndiam *Beta normalised with respect to the diamagnetic energy.* 41, 44

Bt *Toroidal magnetic field.* 41

Bt80 *Axial toroidal Magnetic Field at a $\psi=0.8$ surface.* 41

CAC *Centro de Atención al Cliente.* 66

CBR *Case Base Reasoning.* 31

CE *Conjunto de Entrenamiento.* 34

CN *Conjunto Nuevo.* 34

CP *Conjunto de Prueba.* 34

CRM *Customer Relationship Management.* 27

DalphaIn *Dalpha inner view.* 41

-
- DEC** *Digital Equipment Corporation.* 11
- DENDRAL** *Dendritic Algorithm.* 11
- Dens1** *Line integrated density.* 41
- ELM** *Edge-Localized Mode.* 38, 39
- ELO** *Elongation boundary.* 41
- FDWDT** *Time derivative of diamagnetic energy.* 41
- HLAI** *Human Level Artificial Intelligence.* 14
- IA** *Inteligencia Artificial.* 5, 7, 9–17, 31
- IBM** *International Business Machines.* 9
- Ipla** *Plasma current.* 41
- ITER** *International Thermonuclear Experimental Reactor.* 36, 38
- JET** *Joint European Torus.* 36, 37
- LI** *Plasma inductance.* 41
- LID4** *Outer interferometry channel.* 41, 44
- LSPri** *R coordinate inner lower strike point.* 41
- LSPro** *R coordinate outer lower strike point.* 41
- LSPzi** *Z coordinate inner lower strike point.* 41
- LSPzo** *Z coordinate outer lower strike point.* 41
- M2M** *machine two machine.* 2
- MGI** *McKinsey Global Insitute.* 2

Mode *Confinement mode from observations.* 41

MSOM *Multiple Self-Organized Maps.* 40, 43–45, 47–49, 53, 56

OJA Red de Neuronas Artificiales de Oja. 43, 47, 50–56

PCA *Principal Component Analysis.* 47

PRFL *Temperature profile from KK1.* 41

Ptot *Total heating power.* 41

Q80 *Safety factor at psi=0.8 surface.* 41

Q95 *Safety factor.* 41, 56

Rad *Radiated power.* 41, 44

RFID *Radio Frequency IDentification.* 14

RIG *Radial inner gap.* 41, 52

RNA Redes de Neuronas Artificiales. 16

ROG *Radial outer gap.* 41, 52, 56

ROI *Return on Investment.* 25

RXPL *R coordinate lower XP.* 52

SOAR *State, Operator And Result.* 13

SOM *Self-Organized Maps.* 43–53, 55, 56

Te *Electron temperature at the intersection between LID4 vertical view and PSI=0.8 surface.* 41

Te0 *Temperature in the center.* 41

THEFUMO *Thermonuclear Fusion Modelling.* 43, 47, 50, 52, 53, 56

- TI** Tecnología de la Información. 16
- Ti** *Ion temperature at $\psi = 0.8$* . 41
- TJ-II** Tokamak de la Junta de Energía Nuclear. 36
- TOG** *Top Outer GAP*. 41
- TRIL** *Lower triangularity*. 41, 49, 52, 56
- TRIU** *Upper triangularity*. 41
- USPri** *R coordinate inner upper strike point*. 41
- USPro** *coordinate outer upper strike point*. 41
- USPzi** *Z coordinate inner upper strike point*. 41
- USPzo** *Z coordinate outer upper strike point*. 41
- Wmhd** *Magnetohydrodynamic energy* . 41
- XPrl** *R coordinate lower XP*. 41
- XPru** *R coordinate upper XP*. 41
- XPzl** *Z coordinate lower XP*. 41
- XPzu** *Z coordinate upper XP*. 41
- ZXPL** *Z coordinate lower XP*. 52, 53, 56

Capítulo 1

Introducción

El reto principal del presente trabajo de investigación se puede plasmar en la siguiente cuestión: ¿cómo resolver problemas complejos de sectores y compañías reales a partir del Aprendizaje Automático utilizando repositorios de históricos y datos que se están produciendo en tiempo real?

La herramienta propuesta, AIPAKA, se basa en la definición de un procedimiento o marco de trabajo para, a partir de un problema, sus objetivos y los datos reales involucrados; definir una serie de pasos con el fin de construir los sistemas basados en modelos obtenidos por Aprendizaje Automático a partir de los datos, que solucionan el problema y optimizan los objetivos marcados. Se usará un conjunto de técnicas basadas en la Inteligencia Artificial para la automatización de la extracción del conocimiento implícito en los datos, creación de modelos predictivos y el auto-ajuste de los mismos con los nuevos datos conocidos.

La necesidad del presente trabajo viene apoyada fundamentalmente por los siguientes puntos:

- Necesidad de las compañías y centros de investigación de un mejor aprovechamiento de la información contenida en sus bases de datos y en otros repositorios de información públicos o privados. En 2012 se alcanzó la cantidad de 2,5 zettabytes [Yiu12], para hacerse una idea gráfica de qué significa esta cantidad, se puede poner el siguiente

ejemplo: si se almacena esta información en DVD's y se apilaran, la "pila" de DVDs sería una vez y media, la distancia de la Tierra a la Luna.

Se estima que sólo el 0.5 % de esta información se está usando realmente [GR12], preveyéndose que el universo digital alcance un mínimo de 40 zettabytes en 2020, lo que equivaldría a multiplicar por 50 la información existente a principios del 2010. El potencial económico y de servicios es enorme para las instituciones públicas y privadas. Por ejemplo: para el gobierno británico se estima que una mejor explotación de los datos disponibles, ahorrarían entre 16.000 y 31.000 millones de libras [Yiu12]. O generar un valor de 300.000 millones de dólares en el servicio de salud de Estados Unidos, o aumentar un 60 % el margen de las operaciones de los "retailers".

En los últimos tiempos, se ha acuñado un nuevo término: conjunto de datos demasiado grandes o macrodatos (del término en inglés Big Data). La popularización del término viene indudablemente, ligada al documento publicado por el *McKinsey Global Institute* (MGI) en Junio de 2011 [MCB⁺11], donde se define como "conjuntos de datos cuyo tamaño va más allá de la capacidad de captura, almacenamiento, gestión y análisis de las herramientas de base de datos".

- Complejidad crecientes de los sistemas actuales, incrementándose las aplicaciones máquina a máquina (*machine two machine* - M2M), automatización y necesidad de una mayor velocidad de reacción y flexibilidad ante ciertos eventos. En el mundo actual, no cabe la menor duda, de que cada vez se está más interconectados a nivel empresarial y personal. Es mucho más difícil aislar y medir la influencia de cada una de las variables que son susceptibles de ser recogidas y guardadas. A la vez, aumenta la petición por parte de los clientes o usuarios de servicios públicos, para que se les ofrezcan soluciones personalizadas y en el momento que las necesiten. Por tanto, parece claro que no queda

más remedio que utilizar un análisis multivariable capaz de definir la importancia de cada una de las variables para un cliente concreto.

- Rapidez en la aparición de nuevas tecnologías, que generan a su vez, gran cantidad de información utilizable: tarjetas de crédito, nuevos servicios de telefonía, web, redes sociales, etc ... A los sistemas informáticos actuales, se les exige con mucha frecuencia que sean robustos, flexibles o fácilmente adaptables y escalables. La razón no es más que buscar una rápida adaptación a problemas, restricciones y objetivos en continuo cambio.
- Necesidad de anticipación a la evolución futura de índices o comportamientos. No sólo se le exigen a los sistemas informáticos que sean flexibles, sino que también sigan automáticamente las evoluciones y cambios producidos en el índice o comportamiento de interés. En esta tesis, tendrá una especial relevancia, el descubrir patrones de comportamiento que identifiquen individuos o situaciones importantes para el negocio, y que serán utilizados para predecir la evolución de los clientes o usuarios de un determinado servicio o producto. Permitiendo personalizar los bienes ofrecidos por la empresa, y adelantarse a los requerimientos o necesidades de los clientes.
- Ausencia de procedimientos sencillos e “industrializados” para la aplicación de las técnicas de Aprendizaje Automático a los grandes conjuntos de datos disponibles. Que, además, permitan desplegar los modelos hallados en los sistemas actuales de las empresas, y evolucionarlos a lo largo del tiempo.

El procedimiento y técnicas desarrolladas en esta tesis, tiene muy en cuenta los puntos anteriores, no hay que olvidar la “vocación” práctica de este trabajo. Se quiere proveer de una herramienta de trabajo capaz de crecer con la empresa, enfrentarse a problemas complejos, cambiantes, de los que existen abundantes datos pero que las relaciones entre las diferentes

variables no son triviales, así como no es trivial prever la evolución futura de los índices de interés para el negocio.

Las técnicas que van a servir de base a este desarrollo son: las Redes de Neuronas Artificiales, el Razonamiento Basado en Casos (implementado mediante las Historias Vecinas Análogas), la segmentación automática y la visualización de los patrones de comportamiento y los segmentos que forman. En el siguiente apartado, se hará un estudio en profundidad del estado del arte con respecto a estas técnicas.

Capítulo 2

Estado del Arte

En este capítulo, se hará una breve introducción, a lo que se entiende por Aprendizaje Automático dentro del área de la IA, se definirán las técnicas y algoritmos más relevantes que han sido usados para el desarrollo del procedimiento descrito en esta tesis y aplicados a los casos prácticos que se presentarán en el capítulo 4.

2.1. Introducción a la Inteligencia Artificial y a la Minería de Datos

Nuestro cerebro está continuamente infiriendo nuevo conocimiento de cantidades ingentes de información que recibimos sin interrupción, quizás no se puedan recordar exactamente los datos recibidos, pero sí, a partir de esos datos, tomar decisiones acertadas en nuevos contextos. Es más, si se estudia, la velocidad de procesamiento de la neurona, y su capacidad de “almacenamiento”, éstas son muy inferiores a las capacidades de cualquier ordenador personal de los que se suelen usar en la actualidad. Aún así, la unión de muchas neuronas, es decir, de muchas pequeñas unidades de procesamiento, permiten capacidades de aprendizaje e inferencia infinitamente mayores que las mejores bases de datos, programas y procedimientos de visualización de la información existentes (que no suelen tener ninguna, y sólo

hacen lo que se le ha programado que haga: almacenar información, gestionarla y representarla).

La Inteligencia Artificial, y más concretamente el Aprendizaje Automático, han desarrollado algoritmos de aprendizaje basados en los sistemas nerviosos naturales que permiten “emular” el comportamiento biológico de los procesos de reconocimiento, aprendizaje y generalización. Con lo cual, se puede explotar la información contenida en las bases de datos y obtener, de manera automática, conocimiento útil para la toma de decisiones. Se podría resumir, que a partir de históricos los sistemas de Aprendizaje Automático son capaces de descubrir relaciones entre las variables, generalizando este conocimiento y respondiendo de manera adecuada a nuevos casos, presentados en nuevos contextos y “no vistos” con anterioridad.

Existen algunos desarrollos, que de una manera muy gráfica pueden dar idea de la potencia de algunas implementaciones basadas en el Aprendizaje Automático: por ejemplo, el juego de la web: <http://www.20q.net/>. Juego que a partir del aprendizaje obtenido de millones de jugadores, y de las preguntas que se les hacen a un nuevo jugador sobre un animal, vegetal o cosa en el que está pensando, “adivina” qué es eso en lo que el jugador está pensando. Al final, al ser capaz de procesar toda la información que los jugadores aportan, no es tan difícil adivinar en qué se piensa, no se es tan original, como para pensar en algo en que nadie antes ha pensado o para responder a las preguntas de una forma totalmente nueva, alejada de todas las respuestas dadas por otras personas.

2.1.1. Concepto y breve historia de la Inteligencia Artificial

La Inteligencia Artificial es uno de los campos más nuevos en las ciencias e ingenierías, su estudio empezó después de la Segunda Guerra Mundial acuñándose el término en 1956. La definición de la Inteligencia Artificial no es del todo fácil, se puede tener una idea bastante aproximada de lo que se trata, pero su definición formal ya es más complicada. Siguiendo la división dada por Russell y Norvig [RN10] se pueden encuadrar las definiciones dadas,

en las siguientes categorías: pensamiento humano y racional, y actuación humana y racional. Esta clasificación es ciertamente artificial pero ayuda a estructurar las diferentes definiciones que se han dado sobre la Inteligencia Artificial, a continuación se presentan algunas de ellas:

1. Pensando humanamente:

- “El esfuerzo por hacer que los ordenadores piensen... máquinas con mente, en el sentido completo y literal” [Hau85].
- “La automatización de las actividades que se asocian con el pensamiento humano, actividades como la toma de decisiones, solución de problemas, aprendizaje, ... ” [Bel78].

2. Pensando racionalmente:

- “El estudio de las facultades mentales a través del uso de modelos mentales” [CM85].
- “El estudio de la computación que hace posible percibir, razonar y actuar” [Win92].

3. Actuando humanamente:

- “El arte de crear máquinas que lleven a cabo funciones las cuales requieren inteligencia cuando son llevadas a cabo por personas” [Kur92].
- “El estudio de cómo hacer que los ordenadores hagan las cosas en las que, actualmente, las personas son mejores” [RK91].

4. Actuando racionalmente:

- “Inteligencia computacional es el estudio del diseño de agentes inteligentes” [PMG98].
- “IA ... trata sobre el comportamiento inteligente de artefactos” [Nil98].

Para tener una perspectiva general de la Inteligencia Artificial, es necesario presentar una cronología histórica. Esta visión en el tiempo, también es necesaria para valorar en su justa medida ciertos avances o técnicas, que en un momento dado pueden estar más o menos de “moda”. En este campo, no sería la primera vez que, a grandes expectativas le han seguido periodos de gran frustración. Usando la división histórica dada por Rusell y Norving, la cual se actualizará para presentar la historia más reciente de los últimos años, se pueden distinguir las siguientes etapas históricas:

1. La gestación de la Inteligencia Artificial (1943 - 1955)

El primer trabajo que es reconocido generalmente como perteneciente a la Inteligencia Artificial fue realizado por Warren McCulloch y Walter Pitts [MP43]. Ellos propusieron un modelo de neurona artificial, en el cual, cada neurona se caracterizaba por un estado de “on”-“off”; el cambio a “on” ocurría en respuesta a la estimulación hecha por un número suficiente de neuronas vecinas. Ellos mostraron que cualquier función computable puede ser programada por una red de neuronas conectadas y que todas las conexiones lógicas (and, or, not, ...) pueden ser implementadas por estructuras de red simples. McCulloch y Pitts también sugirieron que las Redes de Neuronas Artificiales podrían aprender.

Donald Hebb [Heb49] desarrolló una regla simple para modificar el peso de las conexiones entre las neuronas. Su regla (Hebbian Learnig) sigue siendo un modelo útil a día de hoy. En 1950, Marvin Minsky y Dean Edmons construyeron el primer ordenador neuronal: el SNARC, que simulaba una red de 40 neuronas. Minsky siguió estudiando la computación universal usando redes de neuronas, siendo bastante escéptico en cuando a las posibilidades reales de las Redes de Neuronas Artificiales [MP69]. Fue el autor de influyentes teoremas que demostraban las limitaciones de las Redes de Neuronas Artificiales.

No se puede terminar este breve repaso a los principios de la Inteligencia Artificial sin nombrar el influyente trabajo de Alan Turing

“Computing Machinery and Intelligence” [Tur50] donde introdujo el famoso test de Turing además de los conceptos de Aprendizaje Automático (machine learning), algoritmos genéticos (genetic algorithms) y aprendizaje reforzado (reinforcement learning).

2. El nacimiento de la Inteligencia Artificial (1956)

El “nacimiento oficial” de la Inteligencia Artificial se puede situar en el verano de 1956 en el Dartmouth College de Stanford. El padre fue John McCarthy que convenció a Minsky, Claude Shannon, y Nathaniel Rochester para reunir a los investigadores más eminentes en los campos de la teoría de autómatas, redes neuronales y del estudio de la inteligencia, con el fin de organizar unas jornadas de trabajo durante los dos meses del verano de 1956. Las jornadas de Dartmouth [MMRS55] no introdujeron ninguna línea rompedora, pero el nuevo campo de la Inteligencia Artificial estuvo dominado por los participantes y sus alumnos durante las siguientes dos décadas. En Dartmouth, se definió por qué es necesaria una nueva disciplina en vez de agrupar los estudios en IA dentro de alguna de las ya existentes: teoría del control, de la toma de decisiones, operaciones, matemáticas ... La primera razón es porque la IA trata de duplicar facultades humanas como la creatividad, el auto-aprendizaje o el uso del lenguaje. Otra razón, es porque la metodología usada parte de la ciencia de la computación y la Inteligencia Artificial es la única especialidad que trata de hacer máquinas las cuales puedan funcionar autónomamente en entornos complejos y dinámicos.

3. Grandes expectativas (1952 - 1969)

Estos años fueron de un gran entusiasmo, apareciendo trabajos bastante prometedores: IBM realizó algunos de los programas basados en Inteligencia Artificial, se creó un sistema capaz de probar teoremas geométricos que hasta para los propios estudiantes de matemáticas suponían alguna dificultad. Arthur Samuel en 1952 creó programas para jugar a las damas, capaces de aprender a jugar, acabando por

jugar mejor que su creador (este programa fue mostrado en televisión en 1956). En 1958 McCarthy creó el lenguaje Lisp que se convirtió en el lenguaje dominante para la IA en los siguientes 30 años. Las redes de neuronas introducidas por McCulloch and Pitts también sufrieron importantes desarrollos, apareció la red Adeline, basada en la regla de aprendizaje de Hebb. Se desarrolló el teorema de convergencia del Perceptrón que aseguraba que el algoritmo de aprendizaje puede ajustar los pesos de las conexiones de un perceptron de manera que ajuste cualquier función definida por las variables de entrada.

4. Una dosis de realidad (1966 - 1973)

Muchos investigadores en el nuevo campo de la IA, hicieron predicciones realmente aventuradas, que en absoluto llegaron a cumplirse. Incluso llegaron a predecir que las máquinas podrían pensar, aprender y crear; teniendo un rápido desarrollo y llegando a superar en poco tiempo a la propia mente humana (Herbert Simon en 1957). Evidentemente, se ha demostrado que es falso, además hubo sonoros fracasos como intentos de crear traductores automáticos del ruso al inglés en los años 60, fracasos que provocaron la retirada de fondos en 1966 del gobierno americano a las investigaciones sobre la creación de traductores. Las explosiones combinatorias de muchos de los problemas abordados por la IA demostraron que eran computacionalmente irresolubles. Los algoritmos evolutivos o algoritmos genéticos, eran computacionalmente muy costosos y en muchos casos no llegaban a ninguna conclusión. Otra dificultad estribaba en limitaciones fundamentales de las estructuras básicas usadas para generar comportamiento inteligente. Por ejemplo, en 1969 Minsky y Papert [MP69] probaron que aunque el perceptrón puede aprender cualquier cosa que pueda representar, la realidad es que puede representar muy pocas cosas.

5. Sistemas basados en el conocimiento (1969 - 1979)

En 1969 nacen los sistemas expertos con los cuales se cambiaba el en-

foque de la Inteligencia Artificial hasta el momento, que estaba basado en encontrar una solución al problema completo a partir de un proceso de “razonamiento” desde principios simples. Los sistemas expertos se basan en reglas o principios más complejos, de un campo de conocimiento mucho más específico, y que en muchos casos prácticamente significa que casi se conoce la respuesta al problema planteado. Uno de los primeros sistemas expertos, fue el programa DENDRAL (*Dendritic Algorithm*) desarrollado en Stanford [LBFL93] que resolvía el problema de determinar la estructura molecular a partir de la información proveniente de un espectrómetro de masas.

6. La Inteligencia Artificial comienza a ser una industria (1980 - presente)

Al principio de los años 80 la IA comenzó a ser una industria, fundamentalmente en Estados Unidos aparecieron compañías con grupos de trabajo dedicados a desarrollos basados en sistemas expertos, robótica y visión artificial; además de la fabricación del hardware y software necesario. Por ejemplo, el primer sistema experto comercial llamado R1, empezó a funcionar en DEC (*Digital Equipment Corporation*) en 1982 [McD82], el programa ayudaba en la configuración de las órdenes de nuevos sistemas de computación. En 1986, la compañía estimaba que el sistema había ahorrado 40\$ millones en un año. Para 1988 DEC había desarrollado 40 sistemas expertos, DuPont tenía 100 en uso y 500 en desarrollo, ahorrando unos 10\$ millones por año.

7. El retorno de las Redes de Neuronas Artificiales (1986 - presente)

A mitad de la década de los 80, desde varios grupos de investigación, se avanzó en el algoritmo de aprendizaje de “back-propagation” para las redes neuronales, en concreto, para el Perceptrón Multicapa, desarrollado originalmente en 1969 [BH69]. Este algoritmo fue aplicado en muchos problemas de aprendizaje y la difusión de los resultados en los

artículos de precesamiento paralelo y distribuido (Parallel Distributed Processing) [RM86] causaron una gran expectación.

En la actualidad, se está avanzando en el uso de herramientas que implementan las redes neuronales, incluso utilizando desarrollos en la nube (del inglés “cloud computing”) [SM13]. Lo que permite usar herramientas de entrenamiento, validación y uso de Redes de Neuronas Artificiales, así como “compartirlas” entre investigadores o desarrolladores de todo el mundo.

8. La IA adopta el método científico (1987 - presente)

A partir de los últimos años de los 80 y hasta el presente, se ha producido una revolución tanto en el contenido como en la metodología de trabajo de la Inteligencia Artificial. Últimamente, es más común construir a partir de teorías ya existentes que desarrollar nuevas, dotando a estas teorías de rigor matemático y mostrando su eficiencia en problemas reales más que en simulaciones o ejemplos simples de laboratorio.

En términos metodológicos, la Inteligencia Artificial ha adoptado firmemente el método científico. Para que una hipótesis sea aceptada, debe estar sujeta a experimentos empíricos rigurosos, y los resultados deben ser analizados estadísticamente para medir su importancia. Ahora es posible replicar los experimentos usando repositorios de datos compartidos, así como datos y código de testeo.

Un ejemplo para ilustrar lo anterior sería el reconocimiento del habla. En la década de los 70, una amplia variedad de arquitecturas y aproximaciones fueron probadas; muchas de ellas fueron hechas “ad-hoc” y con un planteamiento teórico muy débil. Siendo probadas en sólo unos pocos experimentos muy limitados. En los últimos años, aproximaciones basadas en los modelos ocultos de Markov (Hidden Markov Models) ([BP66] y [BE67]) parecen dominar este área. Dos características son importantes en los modelos de Markov: están basados en una teoría matemática rigurosa y son generados a partir de un gran corpus

de datos reales del habla.

Las redes neuronales artificiales han seguido un proceso similar: en un principio el enfoque de muchos desarrollos era mostrar cómo las redes de neuronas diferían de las técnicas “tradicionales”. Con el desarrollo de la metodología y de unos marcos teóricos robustos, se consigue comparar las redes neuronales a las técnicas estadísticas, reconocimiento de patrones y, en general a las técnicas más relevantes de cada aplicación. A raíz de estos desarrollos, la minería de datos se ha convertido en una nueva y vigorosa industria.

Por último, es importante resaltar el papel del razonamiento probabilístico [Pea88] en muchos campos de la IA, basado fundamentalmente en las redes bayesianas [Che85]. Las redes bayesianas fueron inventadas para permitir una representación eficiente y un razonamiento riguroso del conocimiento incierto (uncertain knowledge).

9. La aparición de los agentes inteligentes (1995 - presente)

Quizás por el avance de la Inteligencia Artificial en la solución de problemas muy específicos, se ha vuelto a plantear la cuestión de la solución de problemas generales o desde un punto de vista holístico. Esto da lugar a los sistemas de agentes inteligentes, donde agentes autónomos y especializados en ciertas tareas colaboran entre sí, generando un conocimiento mucho más global. Uno de los ejemplos más conocidos de arquitectura basadas en agentes es el sistema SOAR (*State, Operator And Result*) [LNR87]. Y uno de los entornos más importantes para los agentes inteligentes es Internet: motores de búsquedas, sistemas de recomendación o sistemas de agregación de sitios web.

A pesar de todo, algunos autores de los más influyentes en el campo de la IA (John McCarthy, Marvin Minsky, Nils Nilsson y Patrick Winston) han expresado su descontento con los progresos de la Inteligencia Artificial, ellos piensan que más que seguir mejorando el rendimiento en ciertas áreas o ejemplos concretos; la IA debe retornar al principio expresado por Simon: “máquinas que piensan, que aprenden

y que crean”. De esta corriente han surgido nuevas líneas de trabajo: Inteligencia Artificial Humana (Human-Level AI: HLAI) [MSS04], Inteligencia Artificial General (*Artificial General Intelligence* - AGI) [GP07], e IA amigable (Friendly AI) [Yud08] y [Omo08].

10. La disponibilidad de conjuntos de datos muy grandes (2001 - presente)

En los 60 años de historia de la computación, el énfasis se ha puesto en el algoritmo. Pero trabajos recientes muestran como en muchos problemas tiene mucho más sentido preocuparse por la cantidad de datos disponibles que por el algoritmo a aplicar [Yar95], [BB01], [HE07].

11. La aparición del concepto de macrodatos (Big Data) (2011 - presente)

En los últimos 4 años ha surgido con mucha fuerza un nuevo concepto: macrodatos (del inglés Big Data) [MCB⁺11]. Detrás de este concepto se encuentra la enorme velocidad con la que actualmente se producen y almacenan nuevos datos provenientes de múltiples fuentes, con datos estructurados y no estructurados [GR12].

Haciendo un poco de historia quizás los primeros que usaron el término “Big Data” fueron Michael Cox y David Ellsworth [CE97], refiriéndose al uso de grandes volúmenes de datos científicos usados para la visualización. Actualmente una de las definiciones más usadas de macrodatos es la dada por IBM [ZdP⁺12]: donde se afirma que los macrodatos vienen caracterizados por las tres “V”: Volumen, Variedad y Velocidad. Veamos cada uno de estos tres conceptos:

- Volumen: referido a los grandes volúmenes de datos generados desde una gran variedad de fuentes: la web, el internet de las cosas, etiquetas RFID (*Radio Frequency IDentification*) que identifican las mercancías a lo largo de la cadena de suministro, redes sociales, etc.
- Variedad: se usan múltiples tipos de datos para analizar una si-

tuación o evento. La información puede ser tanto estructurada como no estructurada.

- Velocidad: cada vez aparece una mayor cantidad de datos en menos tiempo y a la vez se incrementa la necesidad de tomar decisiones basados en esos datos. Por ejemplo, la velocidad de producción de datos de las redes sociales es vertiginosa, sólo Twitter produce más de 250 millones de mensajes al día. Mientras que los almacenamientos de datos clásicos (del inglés Data Warehouses), son bastante estáticos, están pensados para recopilar una determinada información a lo largo del tiempo. En los macrodatos todo es más dinámico, la información que actualmente se está recogiendo y analizando puede influir en los siguientes datos que se van a recoger y analizar.

A estas tres “Vs” clásicas, cada vez toma más fuerza el añadir la cuarta “V” de “Valor”, hay que justificar el valor de recoger tanta información, ser capaces de cuantificar los beneficios para la empresa o para la investigación. No se puede creer que el recopilar información tiene un valor “per se”, el esfuerzo de recopilar y almacenar datos debe perseguir un objetivo bien definido.

La Inteligencia Artificial bajo situaciones de grandes volúmenes de datos, como ocurre en los macrodatos, permite abordar problemas de reconocimiento de patrones, aprendizaje y creación de modelos predictivos. Además, los sistemas basados en IA, posibilitan las tomas de decisiones muy rápidas que pueden servir de apoyo a otras decisiones de más alto nivel. Tampoco hay que olvidar uno de los usos más extendidos de las técnicas de IA en el contexto de los macrodatos: “estructurar” información no estructurada.

Hay que señalar, que en el contexto de los macrodatos es muy importante la paralelización, una de las tecnologías más usadas y novedosas en los macrodatos es *MapReduce* [DG04], usada por Google y basada

en dividir un problema en pequeños trozos, cada uno de ellos ejecutado en un nodo de un grupo de ordenadores. MapReduce proporciona una interfaz que permite distribuir y ejecutar en paralelo en un conjunto de ordenadores. Hadoop (<http://hadoop.apache.org>) es una versión libre de MapReduce. Uno de los principales promotores de MapReduce es Yahoo (<http://developer.yahoo.com/hadoop>). Las Redes de Neuronas Artificiales, y muchos de los algoritmos usados en la IA son fácilmente paralelizables, aunque las implementaciones tradicionales hayan estado basadas en un solo ordenador. En el caso de las Redes de Neuronas Artificiales (RNA) su base son nodos que realizan operaciones muy simples y que se encuentran interconectados. Cada nodo o grupo de nodos pueden ser ejecutados en una máquina simple. De hecho desde hace mucho tiempo ha habido implementaciones de RNA paralizadas en componentes electrónicos ([EDT89], [Cau93] o [SBB⁺92]).

También para los organismos públicos, el aprovechamiento de la ingente cantidad de información disponible está tornándose en una prioridad para mejorar los servicios que prestan: sanidad, seguridad, etc. [Yiu12]. Según un informe de la “TechAmerica Foundation” [Tec13], los principales puntos en los que los macrodatos pueden ayudar a las administraciones públicas son:

- Remplazar o dar soporte a las decisiones con algoritmos automatizados.
- Reducir las ineficiencias de las agencias públicas.
- Dar transparencia en la gestión pública.
- Mejorar la puesta en marcha de iniciativas permitiendo la experimentación para descubrir necesidades y explorar las diferentes posibilidades.
- Mejorar el retorno de la inversión de las inversiones en Tecnología de la Información (TI).
- Mejorar la toma de decisiones y la “inteligencia operacional”.

- Proveer de capacidades predictivas para mejorar los resultados.
- Reducir las amenazas de seguridad y el crimen.
- Eliminar el despilfarro, el fraude y el abuso.
- Innovación en nuevos modelos de negocio y de servicios.

Actualmente, el sector de los macrodatos es uno de los que más negocio potencial puede generar en las Tecnologías de la Información [VNB⁺13], manejándose cifras de decenas de miles de millones de euros para los próximos años [DVN⁺13].

Para terminar esta sección, se hará una breve mención a los riesgos de la Inteligencia Artificial que han identificado algunos autores. Por ejemplo, para Omohundro [Omo08] el fin de todos los desarrollos en IA es alcanzar sus objetivos interactuando con el entorno; por lo que podría llevar a cabo acciones, que le permita un mejor ajuste a sus objetivos y para las cuales no fue programado en un principio, por ejemplo, evitar que sea apagado o desconectado, duplicarse en otras máquinas u obtener recursos sin importar la seguridad de bienes o personas.

En [Yud08] se avala la tesis de que las capacidades y riesgos de la Inteligencia Artificial deben ser mejor conocidas, casi todo el mundo piensa que conoce en profundidad lo que puede dar de sí la IA y que esta suele prometer grandes cosas que se quedan en casi nada a la hora de llevarlas a cabo. Pero, al menos, en potencia, la Inteligencia Artificial puede llegar a desarrollar un auto-aprendizaje y capacidad de acción que hagan que ciertas acciones ante nuevas situaciones no se corresponda con los fines para que el software fue desarrollado.

Dejando a un lado las elucubraciones más o menos futuristas anteriores, sí es cierto que la implementación masiva de sistemas dotados de cierta autonomía y capacidad de aprendizaje puede dar lugar a riesgos o situaciones no deseadas. Por ejemplo, en un mercado de valores, cuyas operaciones en una gran mayoría son automáticas, controladas por software que invierte en

los mercados. O de sistemas encargados del control del tráfico aéreo y de los sistemas de seguridad de vehículos. O sistemas que analizan gran cantidad de datos y presentan conclusiones erróneas o los analistas y a los decisores: por ejemplo en un hospital, defensa, etc.

Se impone una cierta prudencia y ser muy conscientes a la hora de desarrollar sistemas basados en Inteligencia Artificial, de controlar el ámbito de actuación de estos sistemas y la redundancia de los mismos (con personas u otros sistemas automáticos), en el caso de aplicaciones críticas. Además, esta prudencia también debe gobernar a la hora de implantar este tipo de sistemas automáticos, ya que, pueden afectar a los procedimientos del negocio y por tanto, a las personas que trabajan actualmente en la empresa, pudiendo en el extremo ser sustituidas por automatismos.

2.2. Requerimientos y necesidades de la industria y de la ciencia

El interés de la industria y de los ámbitos científicos en las tecnologías basadas en la Inteligencia Artificial ha aumentado de manera muy importante en los últimos tres o cuatro años. Basta ver las cifras de más de 2.000 millones de dólares que han movido las grandes tecnológicas (IBM, Google, Facebook, Apple y Amazon), entre 2011 y 2014 [DS14]. Adquiriendo más de 100 empresas emergentes en el campo de la Inteligencia Artificial. Estas grandes tecnológicas han creado departamentos muy potentes basados en la Inteligencia Artificial, donde están trabajando científicos de primer nivel. Estos grupos de trabajos llevan a cabo investigaciones en el ámbito de la excelencia investigadora, trabajando sobre los límites actuales de la ciencia en este campo. En [VL15] y [LMD⁺11] se muestra como el gigante Google desarrolló sistemas muy complejos basados en Redes de Neuronas Artificiales para el reconocimiento del habla y la detección de características relevantes en imágenes no etiquetadas. Usando para ello redes extremadamente complejas de hasta 1.000 millones de conexiones.

Las tecnologías del conocimiento están actualmente en expansión en prácticamente cualquier sector industrial. En la banca ayudan a la detección de fraude, automatización de la atención al cliente mediante las tecnologías de reconocimiento del habla o incluso para el reconocimiento y verificación de la persona que llama y quiere realizar alguna operación. En la salud, los reconocedores automáticos del habla se están usando en las clínicas de psicología de Estados Unidos para transcribir notas al mismo tiempo que se dictan; se están creando sistemas para el análisis automático de imágenes médicas. Incluso se empieza a probar sistemas que usan la literatura médica existente para realizar un diagnóstico y aprender de los resultados con el fin de ir mejorando dicho diagnóstico en futuros pacientes. En las ciencias de la vida el Aprendizaje Automático está siendo usado para predecir las relaciones entre la causa y el efecto, o estimar la actividad de los compuestos acelerando los procesos de la industria farmacéutica. Ejemplos como los anteriores se pueden encontrar cada vez en un mayor número en las compañías de entretenimiento, petroleras, sector público, “retailers” y tecnológicas.

Los beneficios potenciales para la empresa del uso de las tecnologías del conocimiento va mucho más allá de la mera automatización de los procesos:

- Rapidez en las acciones y toma de decisiones.
- Mejorar las salidas del sistema (por ejemplo, en la predicción de la demanda, diagnóstico médico, ...).
- Mejorar la eficiencia tanto de los recursos humanos más preparados como de equipamientos costosos.
- Reducir los costes.
- Hacer posible el escalar y realizar tareas que son imposibles manualmente.
- Innovación en los productos y servicios.

En cuanto al ámbito de la ciencia, se está planteando la aplicación de los futuros desarrollos en campos como el de la computación cuántica a la Inteligencia Artificial [Man13]. En muchos ámbitos científicos, la aplicación de las capacidades del Aprendizaje Automático posibilitan que los científicos de ese ámbito puedan trabajar con datos producidos en entornos altamente complejos, donde la explicación de los fenómenos observados y por tanto, la explicación usando las leyes fundamentales del sistema no es en absoluto trivial. Se podrá ver, una aplicación práctica de esto último, a los dispositivos experimentales de fusión termonuclear en la sección 4.2.

El impacto real de las tecnologías basadas en la Inteligencia Artificial, está creciendo de manera muy significativa. Uno de los factores es la mejora del rendimiento de las técnicas usadas y el que cada vez sea más factible su implementación debido al mayor rendimiento y disponibilidad del hardware necesario.

No hay que olvidar que, quizás el factor más relevante son los miles de millones de dólares que se ha invertido en comercializar estas tecnologías. La mayoría de esta inversión se ha llevado a cabo en la construcción de herramientas analíticas. Las cuáles están orientadas a dar capacidades avanzadas de análisis estadístico, minería de datos y visualización.

El desarrollo de las plataformas de *big-data* o macrodatos para el almacenamiento masivo y distribuido de los datos, también ha sido una de las tecnologías donde más se ha invertido; a pesar de que, gran parte de las soluciones son de código abierto (como todo el universo Hadoop¹). Es en este ecosistema de Hadoop, donde se han creado empresas privadas que ofrecen sus propias distribuciones, u ofrecen distribuciones gratuitas basadas en Hadoop, pero cobrando los servicios de consultoría, desarrollos de proyectos e integración.

En general, la industria está en un momento de saturación debido a la

¹<https://hadoop.apache.org/>

“moda” del *big-data*, parece que toda empresa de cierto tamaño, debe tener *big-data*, porque ahí, en la recogida masiva de los datos se encuentran las claves para hacer que el negocio crezca. Es cierto que en los datos, se encuentran, en una mayoría de casos, el conocimiento implícito necesario para optimizar los procesos del negocio, el servicio al cliente y la rentabilidad. Pero para obtener este conocimiento y para operarlo, no basta con guardar masivamente todos los datos que sean factibles de obtener.

En la ciencia, no existe esa “moda”, desde hace tiempo se disponen de los recursos para monitorizar sistemas industriales o experimentales muy complejos y se trabaja en la modelización de esos sistemas para optimizar los procesos. Las posibilidades que se abren ahora con la Inteligencia Artificial, es poder manejar cantidades ingentes de información para tener una primera aproximación a la descripción de los fenómenos experimentales observados.

Los departamentos de empresas y organismos actualmente dedicados a la investigación, necesitan las siguientes capacidades en el tratamiento masivo de los datos:

1. Capacidad de explotación de la información, no sólo almacenarla *per se*. El almacenar gran cantidad de datos, no sólo puede no resolver los problemas, sino hacer de estos un problema aún mayor, al aumentar significativamente el número de variables y registros para encontrar una solución.
2. Las herramientas analíticas disponibles están pensadas, en su mayoría, para ser usadas por técnicos especialistas en la minería de datos. Estos técnicos pueden usar las aplicaciones analíticas para encontrar modelos que expliquen la realidad modelada, pero esto no asegura que esos modelos sean útiles para la organización, incluso en el caso en que los modelos sean muy precisos describiendo la realidad. Esto es así porque los modelos predictivos deben poder integrarse dentro de los procesos del negocio para que sus soluciones sean operables o utilizables por los sistemas y personas de la empresa. En muchas ocasiones, las soluciones obtenidas dentro del ámbito del *big-data* o el Aprendi-

zaje Automático quedan restringidas a la plataforma analítica donde se creó, siendo muy difícil la integración con el resto de la compañía y, por tanto, del uso por aquellos actores que requieren de la solución. Actores que pocas veces conocerán el “lenguaje” técnico, por lo que las tecnologías de Aprendizaje Automático usadas en las soluciones deben ser transparentes al usuario o sistema final.

Lo anterior, es extrapolable al ámbito científico, simplemente que aquí suele ser necesario encontrar una explicación científica a lo *encontrado* mediante el uso de la Inteligencia Artificial con el fin de comprender el fenómeno observado y poder aprovechar esa mejor comprensión para realizar avances significativos y mucho más rápidos que los que se podrían haber conseguido sin la aplicación del Aprendizaje Automático.

3. Rapidez en la generación del modelo y de su despliegue o uso. Aparecen, realidades cada vez más cambiantes, no sólo de aquellos factores que influyen en el negocio, sino también de los procesos internos del negocio. Donde una ventaja competitiva clara, es ser capaces de adaptarse y prever las necesidades de los clientes, optimizar procesos internos y relaciones con los proveedores.

Esta rapidez, no sólo depende de las capacidades técnicas de los algoritmos basados en la Inteligencia Artificial, sino también en la identificación rápida de qué estructura de aprendizaje es la más óptima para la solución del problema comentado, su implementación e integración en los sistemas y servicios finales que exploten los nuevos datos conocidos haciendo uso de los modelos creados.

El procedimiento AIPAKA trata precisamente de definir un marco de trabajo para la elección de las técnicas de Inteligencia Artificial más adecuadas para resolver un problema, su combinación creando una *estructura de conocimiento*, facilitando su rápida implementación y despliegue en sistemas y servicios reales.

Quizás en el ámbito de la ciencia, este dinamismo en la creación y uso de modelos no es tan acuciante. Suele existir más tiempo para realizar

los experimentos y estudiarlos. Sí suele suceder, especialmente en los ámbitos de las ciencias físicas que estudian comportamientos a nivel atómico y sub-atómico, que se requiera una gran velocidad en la lectura y tratamiento o modelización de las señales. Esto es debido a que esas señales son producidas por partículas con una vida muy pequeña o con dinámicas extremadamente complejas, cambiantes y difíciles de medir.

4. Ajuste rápido de los sistemas basados en modelos. La obtención de buenos modelos por Aprendizaje Automático y su integración en los sistemas o servicios que lo hagan accionables por los diferentes actores o sistemas de la empresa, es en realidad el primer paso del ciclo de explotación del conocimiento. Los modelos implementados deben poder ajustarse rápidamente a realidades complejas y muy cambiantes. Las técnicas de Aprendizaje Automático son fundamentales para este cometido al permitir implementar mecanismos que reajusten los modelos ya obtenidos con anterioridad dependiendo de los nuevos datos conocidos.

Respecto a la industria y analizando algo fundamental para cualquier negocio: la rentabilidad. No se puede caer en la tiranía de la “moda tecnológica” del momento sin tener en cuenta si las inversiones van a solucionar problemas relevantes para el negocio, y por tanto, van a ser rentables. Por ejemplo, en una gran mayoría de empresas no son necesarias grandes inversiones en tecnologías relacionadas con el *big-data*, con el fin de almacenar más datos provenientes de múltiples fuentes internas y externas. Bastaría poder explotar el conocimiento implícito que ya contienen los datos con los cuales cuentan en la actualidad.

Los datos a los que se hace referencia en el párrafo anterior, que ya están disponibles por la compañía, es lo que algunos autores han venido a llamar “los pequeños datos” (del inglés “Small Data”). Datos que actualmente están disponibles (accesibles), se entienden y se enfocan a ciertas tareas (accionables). Este término de “pequeños datos” contrasta con el de macrodatos,

el cual, en ciertas ocasiones puede resultar una exageración para el tipo de sistemas, capacidades y necesidades reales de la empresa. Las diferencias entre conjuntos “pequeños de datos” y macro conjuntos de datos se pueden apreciar en la tabla 2.1, donde se comparan las “Vs” que definen las características de los macrodatos (ver sección 2.1.1), con esas mismas características en los “pequeños datos”.

En general, la idea de los “pequeños datos” es que las compañías pueden obtener resultados sin tener la necesidad de adquirir los sistemas que normalmente serían necesarios para las analíticas usando macrodatos. Muchas empresas se dan cuenta que en algunos casos pueden obtener grandes resultados usando conjuntos de datos robustos y mucho más pequeños que los que normalmente se manejan cuando se habla de implementar técnicas o procesos de basados en los macrodatos.

Los “pequeños datos” es una de las vías, a las cuales las empresas están volviendo para usar los recursos de los que ya disponen de manera eficiente y evitar gastar mucho más de lo necesario (no sólo en presupuesto sino también en tiempo) en ciertos tipos de tecnologías.

Resumiendo, se podría preguntar, ¿por qué los “pequeños datos”? y algunas respuestas podrían ser:

- *Los macrodatos son complejos.* Implementar tecnologías basadas en el almacenamiento masivo de datos, su análisis y esperar que de beneficios puede ser realmente lento. Sin mencionar que en muchas aplicaciones no son necesarios. Por ejemplo, muchas de las estrategias de marketing actuales, incluso focalizando campañas de manera personalizada, no necesitan del uso de macrodatos.
- *Los “pequeños datos” están realmente disponibles “entorno a nosotros”.* Los canales web ofrecen una gran riqueza para la recolección de “pequeños datos”, realmente útiles para informar de las decisiones de los compradores a los departamentos de marketing. Por ejemplo, cada vez que un potencial comprador entra en una tienda on-line y busca, compra, etc. está creando una *firma personalizada y digital*.

- *Los “pequeños datos” son el centro de los sistemas de gestión de los clientes.* Los sistemas de gestión del cliente son usados para conocer a los clientes de las empresas, segmentar el mercado, analizar los productos y servicios de los competidores. Combinando su información con los datos recogidos desde la web y de los sistemas transaccionales se puede tener un perfil completo y rico de los consumidores.
- *ROI - Return on Investment.* La rentabilidad de la inversión y la rapidez en el retorno de la misma, en una mayoría de casos sigue siendo superior en implementaciones en torno a los “pequeños datos” que en los sistemas construidos en torno a los macrodatos.
- *Los “pequeños datos” son sobre el usuario final.* Cuando no se tiene en cuenta tanto los macrodatos, intentando recogerlo absolutamente todo, de cualquier fuente disponible, se tiende a focalizar mucho más en el usuario final, qué es lo que necesita y cómo “entran en acción”. Focalizándonos primero en los usuarios, muchas de las decisiones en tecnología comienzan a estar más claras.
- *Simple.* Los “pequeños datos” pueden ser los datos correctos, y algunos de estos datos serán el comienzo de los macrodatos. Pero con la ventaja de que no es necesario ser un “genio” de los macrodatos y de su análisis para comprender o aplicar las conclusiones obtenidas de un volumen de datos menor y mejor conocidos por la empresa a las tareas de cada día, es todo mucho más simple.

Para muchas cuestiones y problemas, los “pequeños datos” son suficientes. Los datos de mi proveedor de energía, los horarios del transporte público, los gastos públicos, ... Todos ellos no son macrodatos pero pueden ayudar a solucionar muchos problemas del día a día a una mayoría de usuarios.

Quizás el problema no esté tanto en la distinción por el volumen y complejidad de los datos sino en la accesibilidad a los mismos, en la capacidad de extraer el conocimiento implícito y poder usar ese conocimiento de manera fácil y dinámica.

Categoría	Macrodatos	Pequeños Datos
Fuentes de Datos	Datos generados fuera de la empresa desde las fuentes de datos no tradicionales: ■ Redes sociales. ■ Sensores de datos. ■ Datos provenientes de los “logs”. ■ Datos provenientes de los dispositivos (por ejemplo, en el internet de las cosas). ■ Vídeos. ■ Imágenes. ■ Etc.	Datos que tradicionalmente forman parte de los datos empresariales: ■ Datos de los sistemas transaccionales. ■ Sistema de gestión de las relaciones con los clientes (CRM - <i>Customer Relationship Management</i>). ■ Transacciones realizadas a través de la web. ■ Datos financieros y de costes. ■ Etc.
Volumen	■ Terabytes (10^{12}) ■ Petabytes (10^{15}) ■ Exabytes (10^{18}) ■ Zettabytes (10^{21})	■ Gigabytes (10^6) ■ Terabytes (10^{12})
Velocidad	■ A menudo tiempo real ■ Requiere respuesta inmediata	■ Procesamiento por lotes ■ Casi tiempo real ■ No siempre requiere respuesta inmediata
Variedad	■ Estructurado ■ No estructurado	■ Estructurado ■ No estructurado

Tabla 2.1: Macrodatos vs “pequeños datos”.

;

Capítulo 3

AIPAKA

Artificial Intelligence in a Process for Automated Knowledge Acquisition and Applications

En este capítulo, se desarrollará el proceso AIPAKA, proceso para la construcción de modelos que constituyan la base de sistemas inteligentes, capaces de aprender automáticamente a partir de los datos obtenidos de las fuentes de información. Este procedimiento debe ser genérico y con capacidad de ser usado para la automatización de las decisiones y tareas de modelado. Permitiendo abordar cualquier problema de creación de sistemas basados en la inferencia del conocimiento implícito en los datos, de una manera “normalizada” y sin necesidad de grandes conocimientos en Inteligencia Artificial.

3.1. Introducción y Objetivos

Actualmente, el trabajar con técnicas de Inteligencia Artificial con el fin de construir sistemas *inteligentes* es una especie de *arte*, donde cada *artista*¹ usa su saber, para encontrar la solución a los problemas que se le van planteando. Pero no se suele seguir una metodología para resolver estos problemas indicando, por ejemplo, qué técnicas y arquitectura de conocimiento

¹En la actualidad a las personas que se dedican a la explotación analítica de la información se les suele llamar *científico de datos* (del inglés *data scientist*)

usar para llegar a una solución adecuada.

Alguna de las consecuencias que conlleva esta situación se pueden enumerar como siguen:

1. Dificultad en transmitir el conocimiento de *cómo se hace*, y de generalizar ese conocimiento, con el fin de abarcar el mayor número de problemas posibles.
2. Imposibilidad de trabajar de manera conjunta con otros investigadores o grupos en la resolución de los problemas, ya que, cada uno tiende a tener su parcela de conocimiento, técnicas, y formas de hacer que dificulta el avanzar en las soluciones de manera conjunta.
3. Dificultad creciente en el mantenimiento de los proyectos, ya que, cada desarrollo se ha afrontado y realizado de una manera diferente.
4. No poder formar a técnicos en la Inteligencia Artificial de manera rápida y “coherente”, enseñando unos principios básicos, procedimiento y herramientas que faciliten la productividad de los equipos de trabajo y la flexibilidad en las responsabilidades asignadas.
5. Desconfianza de la dirección del negocio, ante la falta de profesionales cualificados y un procedimiento de trabajo definido. Con grandes dudas, en el enfoque e inicio de los proyectos: herramientas a usar, técnicas de Inteligencia Artificial más adecuadas, tiempo de desarrollo, etc.
6. Falta de ambición en los proyectos. La propia dificultad en el establecimiento de una “base mínima” sobre la que sustentar los desarrollos, hace que no se quiera que estos sean complejos o ambiciosos en cuanto al alcance. Llegando en muchos casos, a no pasar de pruebas iniciales, con poco convencimiento e implicación de las diferentes partes de la compañía.
7. Despliegues en producción muy costosos. Si se consiguen desarrollar unos modelos basados en el Aprendizaje Automático, suele existir una

gran distancia entre la consecución y las pruebas de esos modelos en el entorno de desarrollo, y su despliegue en los entornos de producción. Debido a la propia falta de procedimiento que hace que la forma de trabajar y herramientas usadas por el analista no sean conocidas por las personas encargadas de la implementación de los modelos en producción.

Esta falta de implementación en los procesos productivos de las empresas hace que la dirección siga sin confiar en estas tecnologías, alimentando el círculo vicioso y evitando la aplicación real de la Inteligencia Artificial en un gran número de empresas.

En resumen, las empresas suelen ver la Inteligencia Artificial y la explotación de la información implícita en grandes cantidades de datos como algo bastante lejano en cuanto al beneficio real que pueden obtener. Haciendo alguna tímida aproximación, en cuanto a la adquisición de licencias de herramientas analíticas, que suelen usarse como herramientas estadísticas, y para el reporte de información. O en el mejor de los casos, para empezar a introducir tecnologías de almacenamiento de *big-data*, introduciendo las bases de datos no-Sql y distribuidas. Pero siguen con las mismas cuestiones de negocio sin resolver: qué le puedo vender a qué clientes, qué clientes me van a abandonar, cómo puedo detectar y prevenir el fraude, cómo puedo optimizar mis procesos operativos, etc. En el mejor de los casos, la empresa cuenta con más datos almacenados, o con más estadísticos, informes y cuadros de mandos generados. Lo cual, puede complicar aún más el hallar una solución, debido a la abundancia de datos e información disponible donde buscar esa solución.

AIPAKA como indica su nombre traducido del inglés: Inteligencia Artificial en un Proceso para la Adquisición Automática de Conocimiento y sus Aplicaciones, trata de normalizar las decisiones que se toman a la hora de definir una arquitectura de extracción de conocimiento y los pasos a seguir para la creación de modelos predictivos.

AIPAKA está basada, principalmente, en Redes de Neuronas Artificiales y

en el Razonamiento Basado en Casos (CBR - *Case Base Reasoning*). Aunque no es un proceso estático y rígido, sino que, por el contrario está abierto a la incorporación de nuevos algoritmos y técnicas de Inteligencia Artificial. De hecho, en los ejemplos de aplicaciones que se verán en el capítulo siguiente (4), se usan otras técnicas como, por ejemplo, la *segmentación automática*.

El **objetivo** principal de AIPAKA es: *dado un problema de minería de datos, disponer de un procedimiento para ser capaz de diseñar y llevar a cabo, una solución basada en la construcción de modelos predictivos, sin necesidad de poseer grandes conocimientos de técnicas de Inteligencia Artificial*. El procedimiento trata de simplificar la solución a los problemas planteados, aplicando una serie de principios o reglas que orienten en el uso de técnicas de IA para la construcción de modelos predictivos que den solución al problema que se quiere resolver.

Capítulo 4

Casos de Aplicación

En este capítulo, se presentan varios casos de aplicación a diferentes sectores, siguiendo el procedimiento de AIPAKA o aplicando las técnicas de Inteligencia Artificial contenidas en este procedimiento. Algunos de estos casos han ayudado a definir AIPAKA. Fundamentalmente, en cuanto a la definición de los tipos de problemas considerados, y a los algoritmos y arquitectura de los sistemas de conocimiento recomendados en la fase de modelado del procedimiento.

Los casos están basados en experiencias reales en los que se ha trabajado y pertenecen a diferentes sectores: telecomunicaciones, energía, banca, producción y distribución de bienes de consumo.

Con estos ejemplos de modelado reales, se quiere mostrar que el proceso de extracción de conocimiento presentado en este trabajo, no se queda en lo meramente teórico o conceptual, sino que ha sido desarrollado a partir de la experiencia. Esta experiencia ha mostrado que lo importante no es almacenar cada vez más información, antes estructurada y ahora cada vez más no estructurada, sino ser capaces de crear modelos predictivos usando las herramientas provistas por la Inteligencia Artificial. Desplegando estos modelos en los entornos de producción, manteniendo el rendimiento medido en las fases de desarrollo. Con este despliegue “se materializa” el conocimiento extraído, pudiéndose utilizar para operar el negocio o proceso.

4.1. Predicción de la demanda eléctrica a corto plazo

Un problema importante en los sistemas eléctricos es la imposibilidad de almacenar la energía con un nivel aceptable de rendimiento. Esto significa que, es necesario producir la energía que será consumida. Si hay diferencias entre la energía producida y consumida, el sistema no está aislado y el rotor de los generadores podría acelerarse o desacelerarse, por lo que la frecuencia de la electricidad (50Hz) cambiaría. Si el sistema no está aislado, puede ocurrir una asignación o salida de energía con otros sistemas eléctricos con los cuales está conectado.

Debido a lo anterior, existen sistemas de regulación para producir la misma energía que se va a consumir. Pero el sistema de regulación, sólo puede trabajar en un rango pequeño, cuando la demanda cambia rápidamente otras centrales deben empezar a producir electricidad. El problema es cuando hay diferentes tipos de centrales: las hidráulicas son muy rápidas en empezar a entregar energía, en el otro extremo, las centrales nucleares son muy lentas en comenzar a producir electricidad.

Es necesario predecir el consumo de electricidad a corto plazo para incrementar la eficiencia de la generación. La predicción de la demanda depende de variables que no se pueden conocer: el número y tipo de los equipos eléctricos conectados a la red eléctrica y el momento en el que fueron encendidos o apagados. Pero se puede hacer una estimación de las necesidades futuras de potencia, usando los datos históricos para construir modelos predictivos de la demanda de energía eléctrica a corto plazo.

El problema presentado en este caso, es para una de las más importantes empresas proveedoras de electricidad en España. Se han usado datos reales de consumo eléctrico de cuatro años y medio. Los primeros tres años han sido usados para el *Conjunto de Entrenamiento* (CE), el cuarto año para el *Conjunto de Prueba* (CP) y el *Conjunto Nuevo* (CN) contiene los últimos

datos correspondientes a medio año.

4.2. Fusión Termonuclear: segmentación en dispositivos experimentales complejos

4.2.1. Introducción

En todos los casos de aplicación descritos hasta el momento en este capítulo, se han tratado problemas en los cuales se aplican algoritmos de Aprendizaje Automático para crear soluciones que predigan comportamientos o evoluciones de índices, segmenten o clasifiquen.

En este último caso se va a cambiar el enfoque: no se quiere crear una solución, sino crear y testar algoritmos que sean capaces de enfrentarse a problemas realmente complejos. Mucho más allá de lo que comúnmente se llaman macrodatos (ver 2.1.1). Estos algoritmos deben permitir explotar los datos producidos en las descargas, y posibilitar que los científicos que trabajan en este campo puedan dar “sentido físico” a los fenómenos descubiertos por el modelado.

El problema con el que se está experimentando es el de la fusión termonuclear. Cualquier dispositivo de fusión actual dispone de bases de datos con millones de registros, por ejemplo: el *stellarator* español TJ-II¹ y el JET (*Joint European Torus*)². En el caso del *stellarator* TJ-II se gestionan más de tres millones de señales por las más de 30.000 descargas producidas. La mayoría de estas señales son de evolución temporal, con un número de muestras entre 10.000 y varios millones.

El caso del dispositivo de fusión más grande en la actualidad: JET, la base de datos es de 100 Tbytes y su tendencia de crecimiento refleja que la cantidad de datos adquiridos se duplica cada dos años. Los magnitudes anteriores van a ir incrementándose con el tiempo ya que las técnicas de medida van generando día a día un número creciente de señales, es fácil imaginar lo que supondrá la futura base de datos del ITER (*International Thermonuclear Experimental Reactor*)³. El número de señales por descarga que se espera

¹http://www-fusion.ciemat.es/New_fusion/es/

²<http://www.jet.efda.org/jet/>

³<http://www.iter.org/>

recoger estará entre 500.000 y 1.000.000, por lo que la tasa esperada de datos es de Pbytes/año. En la actualidad, sólo se han podido analizar el 10 % de la información almacenada en las bases de datos de JET.

Existen problemas muy importantes que resolver para la construcción de reactores comerciales de fusión, cuyas claves para ser resueltos pueden estar implícitas en la información guardada. Pero se sigue sin tener los mecanismos para explotar dicho conocimiento implícito. Cualquier colección de datos carece de valor sin unos mecanismos eficientes para extraer información y conocimiento de los mismos.

En conclusión, lo que se pretende es probar algunos algoritmos de Aprendizaje Automático usados por AIPAKA en *condiciones extremas*, con el fin de asegurar su correcto desarrollo en entornos mucho más *normales* en los diferentes sectores y problemáticas donde se pueden aplicar. Los problemas de modelado de procesos físicos relacionados con la fusión mediante confinamiento electromagnéticos, son sistemas muy complejos donde se monitorizan miles de señales con comportamientos altamente no lineales, donde los problemas de ruidos son muy relevantes, por lo tanto, constituyen un banco de pruebas extremo donde validar la robustez y capacidad de los algoritmos de Inteligencia Artificial.

En el caso presentado, habrá que centrarse en el estudio de las transiciones L-H⁴ que se dan en la fusión. Pero se está trabajando de manera muy intensa en la selección de características que ayuden a determinar las *disrupciones*⁵. Las cuales, tienen unos efectos muy negativos en las reacciones de fusión, interrumpiéndola y pudiendo dañar de manera importante el dispositivo donde se está llevando a cabo. En [PVM⁺15] se muestran algunos de estos trabajos, utilizando predictores de Venn [VSN03] y algoritmos genéticos para la selección de señales o características relevantes medidas de

⁴En la siguiente sección se describen los modos L y H. Baste decir aquí que son los regímenes de confinamiento del plasma (Low - High).

⁵La disrupción es un fenómeno complejo, relacionado con inestabilidades magnetohidrodinámicas del plasma, en el que provoca una pérdida rápida de energía y la consiguiente terminación brusca de la descarga.

la reacción de fusión y su uso para predecir si habrá una disrupción o no.

Este tipo de problemas, van a ser enfrentados en la explotación del futuro ITER el cual, va a ser un dispositivo de pulso largo (1000 s) que va a generar millones de señales con muy alta dimensionalidad. Las grandes masas de datos que se adquieran en ITER, solamente podrán ser analizadas de manera eficiente usando técnicas de minería de datos como las propuestas en el presente caso.

Antes de continuar con el desarrollo de este caso, indicar que será referenciado como THEFUMO (Thermonuclear Fusion Modelling). Nombre que se le ha dado al proyecto sobre el que se apoyan estos estudios.

4.2.2. Tipo de problema

El objetivo de la solución a implementar es el detectar los diferentes tipos de transición L-H que se dan en la reacción de fusión. No se conocen a día de hoy, si las transacciones L-H tienen diferentes tipos o sólo hay una forma de transitar desde el estado L al H. Por supuesto, tampoco se saben las variables que influyen en los diferentes tipos de la transición L-H. Según lo anterior, se está ante un problema de segmentación y de detección de las variables importantes.

Se verá un poco más en detalle, el problema que se presenta: En fusión es muy importante el estudio de los regímenes de confinamiento del plasma (modo H y L) y especialmente, las transiciones entre ambos estados. Este estudio está orientado a comprender cuántos tipos de transiciones L-H y H-L existen y cuáles son las variables más relevantes que influyen en cada uno de los tipos de transiciones. Uno de los objetivos que los científicos persiguen con este estudio es comprender por qué se producen los ELMs (*Edge-Localized Mode*) lo cual sería muy importante para lograr reacciones de fusión termonuclear estables.

A continuación, se describen brevemente los términos usados en la descripción de este problema⁶:

⁶<http://fusedweb.pppl.gov/Glossary/glossary.html>

Modo H (H-mode). Régimen de alto confinamiento, que ocurre en los plasmas por encima de un cierto umbral de potencia de calentamiento, y que se caracteriza por un gradiente alto de presión en el borde del plasma, denominado pedestal, un aumento del tiempo de confinamiento de la energía de más de un factor 2 y la aparición de modos localizados en el borde, denominados *Edge-Localized Mode*. El nombre procede de la inicial del término inglés que significa “alto”, high.

Modo L (L-mode). Régimen de confinamiento bajo, el normal de operación de un tokamak con calentamiento adicional. El nombre procede de la inicial del término inglés que significa “bajo”, low.

ELM (*Edge-Localized Mode*). Inestabilidad magnetohidrodinámica que se manifiesta durante el modo H. Generalmente, se designa mediante el acrónimo ELM. Produce pérdidas transitorias de energía y partículas que en el caso de los modos de gran tamaño pueden dar lugar a daños en los materiales. Los ELMs más pequeños pueden ser positivos pues contribuyen a la eliminación del helio producido por las reacciones de fusión.

Parece claro que, el tipo de problema planteado, es de *segmentación*. Se necesitan segmentar los patrones definidos por las señales procedentes de los sensores que monitorizan el proceso de fusión termocuclear, con el fin de identificar los diferentes tipos de transiciones L-H que se tienen.

4.2.3. Datos

Se está trabajando con más de 800 descargas del dispositivo experimental de fusión. Se disponen de entorno a 1,5 millones de patrones para la construcción de los conjuntos de entrenamiento. Cada patrón está compuesto por vectores de 100 muestras de la señal tratada. Para el caso del estudio de transiciones L-H, se tomarán muestras de las señales en una ventana de

entre 300 ms y 2 seg. centradas en el momento de la transición⁷.

Lógicamente, el número total de patrones dependerán del tamaño de la ventana usada y del número de patrones para entrenamiento y validación que se quieran usar en la construcción y validación de cada modelo. Las señales usadas se listan en la tabla 4.1.

A modo de ejemplo, en la figura 4.1 se dibuja la forma de onda de algunas de las señales. En estas gráficas, la transición L-H ocurre en el instante 150 (valor en el eje de abscisas). Son ventanas de 300 ms centradas en el momento de la transición L-H.

4.2.4. Modelización

La arquitectura de modelado que se está probando está basada en una MSOM (*Multiple Self-Organized Maps*).

Como ya se sabe, se pretende averiguar cuántos tipos de transiciones L-H existen y cuáles son las variables más relevantes que influyen en cada uno de los tipos de transiciones. El sistema diseñado permite estudiar, en general, cualquier instante de tiempo durante el proceso de fusión, bastaría con usar como entradas de la MSOM aquellos vectores de señales obtenidos alrededor del instante sobre el que se quiere trabajar (que en este caso es principalmente L-H, aunque la H-L también resulte especialmente interesante).

4.2.5. Validación

Al trabajar con el sistema construido e intentar elegir los elementos de cada una de las capas hay que ir resolviendo los problemas de cada uno de los *niveles* desde arriba (*preparación de las entradas*) a abajo (*segmentación automática*). En los siguientes puntos se verán los resultados obtenidos para cada una de las capas, se sustituirá la primera capa de la red MSOM,

⁷Los datos de muestras, ventanas y otros que vayan apareciendo a lo largo de la sección son configurables por software. Uno de los objetivos es estudiar también aquellos parámetros que optimizan los resultados obtenidos.

<i>Nombre</i>	<i>Descripción (en inglés)</i>
Bndiam	Beta normalised with respect to the diamagnetic energy
Bt	Toroidal magnetic field
Dens1	Line integrated density
ELO	Elongation boundary
FDWDT	Time derivative of diamagnetic energy
Ipla	Plasma current
LI	Plasma inductance
LID4	Outer interferometry channel
Ptot	Total heating power
Q95	Safety factor
Wmhd	Magnetohydrodynamic energy
Te	Electron temperature at the intersection between LID4 vertical view and PSI=0.8 surface
Ti	Ion temperature at $\psi = 0.8$
TRIL	Lower triangularity
TRIU	Upper triangularity
XPrl	R coordinate lower XP
XPzl	Z coordinate lower XP
XPru	R coordinate upper XP
XPzu	Z coordinate upper XP
LSPri	R coordinate inner lower strike point
LSPzi	Z coordinate inner lower strike point
LSPro	R coordinate outer lower strike point
LSPzo	Z coordinate outer lower strike point
USPri	R coordinate inner upper strike point
USPzi	Z coordinate inner upper strike point
USPro	R coordinate outer upper strike point
USPzo	Z coordinate outer upper strike point
RIG	Radial inner gap
ROG	Radial outer gap
Q80	Safety factor at $\psi=0.8$ surface
Bt80	Axial toroidal Magnetic Field at a $\psi=0.8$ surface
DalphaIn	Dalpha inner view
TOG	Top Outer GAP
PRFL	Temperature profile from KK1
Rad	Radiated power
Te0	Temperature in the center
Mode	Confinement mode from observations

Tabla 4.1: Señales muestreadas del dispositivo experimental de fusión.

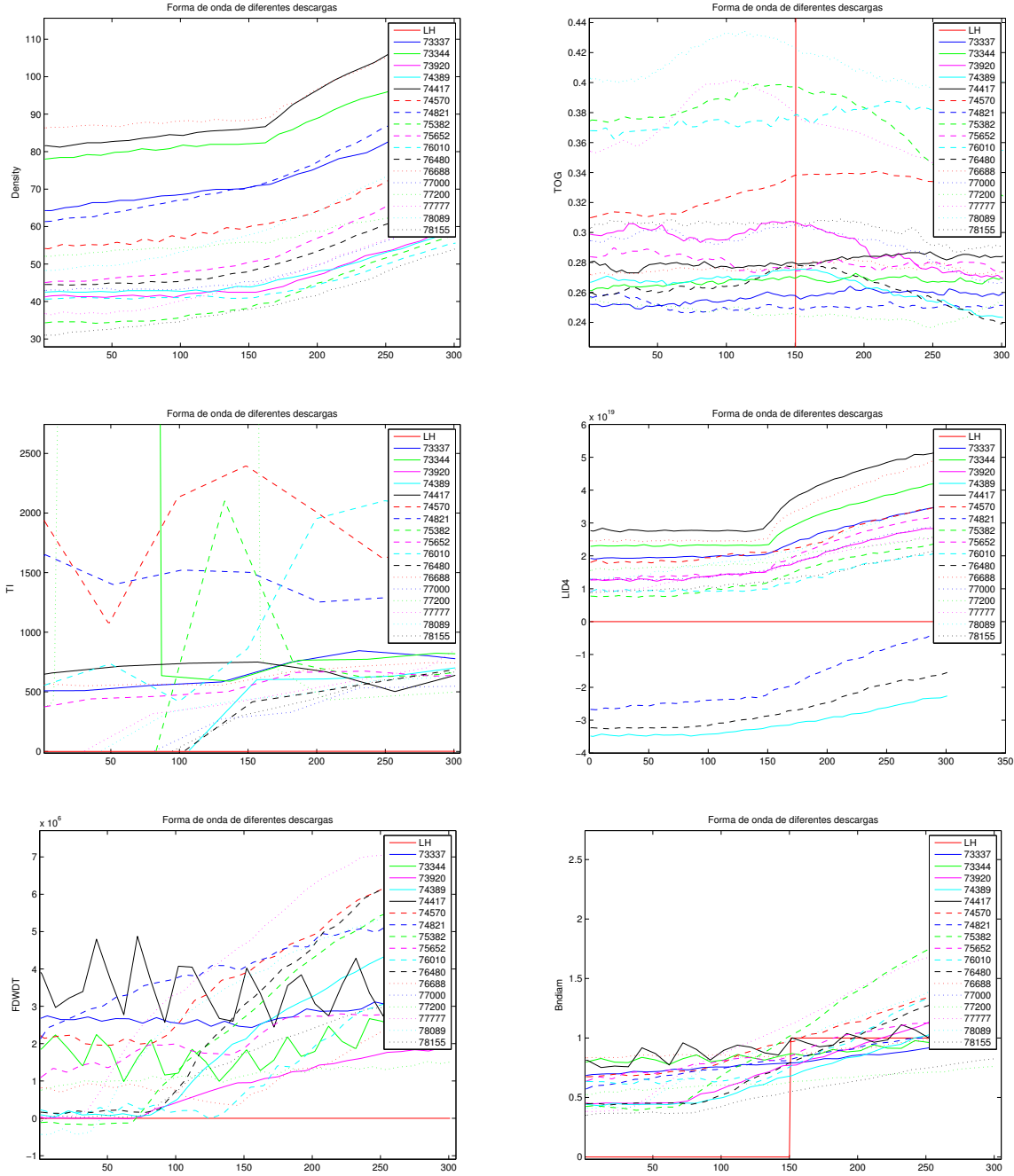


Figura 4.1: Señales en una ventana de 300ms centradas en la transición L-H (en $x=150$). Se han representado la misma señal para varias descargas. Las señales que se representan son: *Density*, *TOG*, *TI*, *LID4*, *FDWDT*, *Bndiam* (tabla 4.1).

compuesta por SOMs especializadas para cada señal, por una red OJA. Y se terminará con unas conclusiones generales, una vez construido el sistema y analizados mínimamente los resultados.

- *Análisis de la primera capa de la MSOM.*

En este apartado, se muestran los resultados de algunas de las SOM de la primera capa. Se comprueba que las redes creadas distinguen las diferentes formas de onda que conforman las señales tratadas. En la figura 4.2 se presentan ejemplos de clasificación de formas de onda para los patrones caídos en una neurona de una SOM de la primera capa la cual trata una de las señales leídas del dispositivo experimental de fusión. En estos ejemplos se puede apreciar como para una misma señal las redes SOMs han *agrupado* en una celda aquellos vectores con una forma de la señal similar. En cierta medida esta primera capa de la MSOM sirve para describir la forma de onda de las señales de entrada.

- *Análisis de la segunda capa de MSOM .*

Como ya se ha indicado, las neuronas ganadoras de la primera capa constituirán el vector de entrada para la SOM de la capa segunda.

Se están probando varias redes para todas las SOMs que componen la MSOM tanto en la primera capa como en la segunda. Eligiéndose aquellas, que parecen se ajustan mejor a las diferentes formas de onda que componen las señales, o diferencian bien los grupos o segmentos donde se encuentran un número apreciable de patrones de comportamiento. Esta última parte, es la fundamental para la SOM del segundo nivel que es usada para distinguir los diferentes segmentos que existen.

4.2.5.1. Primeros estudios de las señales usando THEFUMO

Para estudiar la segmentación conseguida mediante THEFUMO e intentar obtener resultados que puedan explicar cuántos tipos de transiciones L-H existen y cuáles son las señales más influyentes en cada uno de esos tipos, se analizará THEFUMO en sentido ascendente: desde la salida hasta

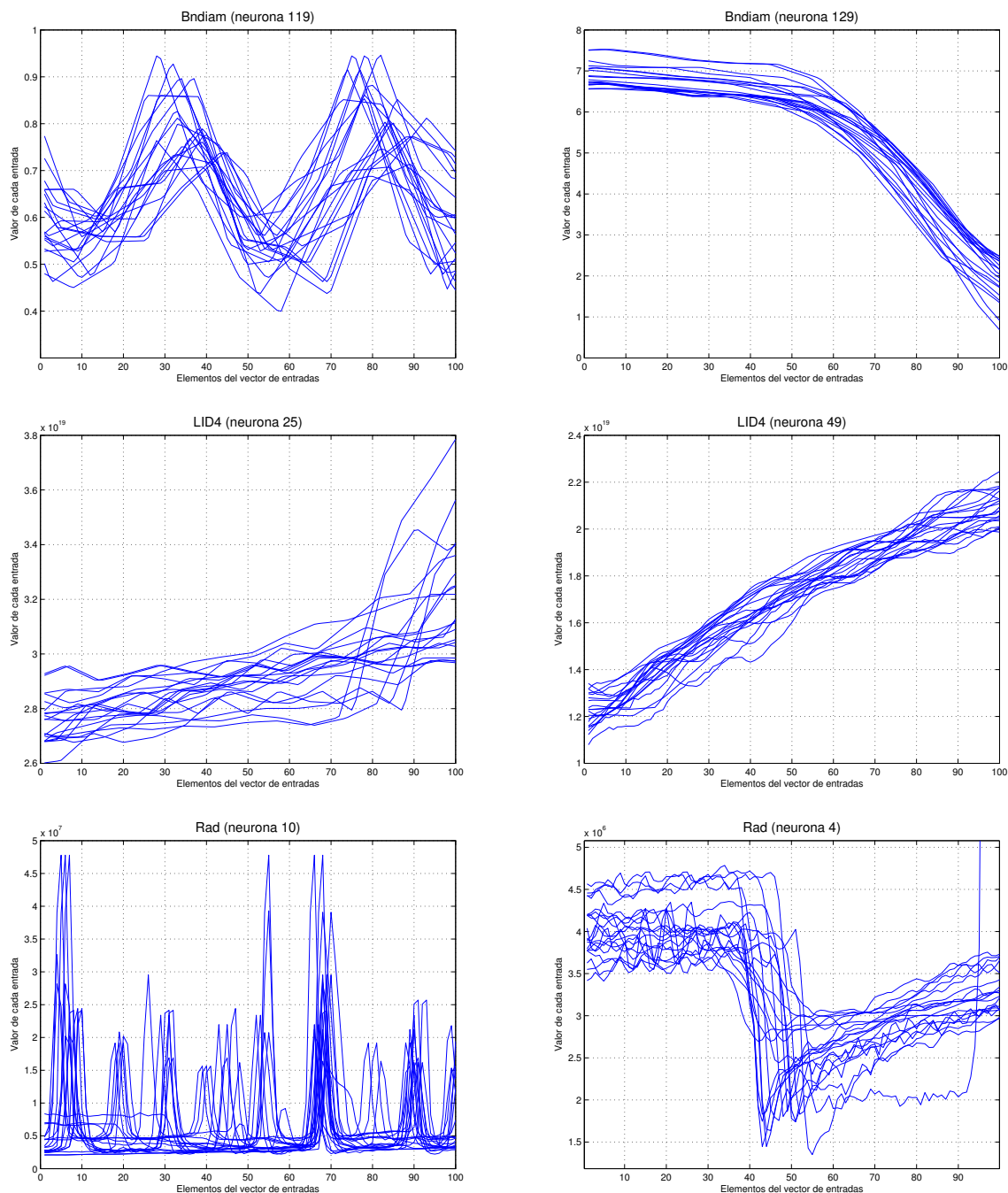


Figura 4.2: Formas de onda de dos neuronas de las SOMs de la primera capa de la MSOM, para el análisis de las señales Bndiam, LID4 y Rad.

las señales de entrada. Se estudia la salida obtenida por la segmentación automática, y en sentido ascendente, se hallarán qué neuronas han “caído” en cada uno de los segmentos de la SOM de la segunda capa de la MSOM, y se comprobarán qué neuronas de las SOMs de la primera capa de la MSOM han activado cada neurona de la segunda capa.

La descripción que puede parecer compleja, pero no es más que recorrer desde la salida hacia la entrada el sistema construido con el fin de, una vez identificados en la salida los diferentes segmentos que se han formado, descubrir qué señales originales son las que más han influido en la formación y diferenciación de esos segmentos.

Como ya se ha indicado, para cada una de las capas de la MSOM se han probado múltiples redes. Para la SOM de la segunda capa de la MSOM se han escogido 6 redes de todas las entrenadas; de dimensiones 5x5, 7x5 y 7x7, obteniéndose los resultados recogidos en la tabla 4.2. Los casos que se han seleccionado, han sido escogidos porque, las neuronas seleccionadas de forma automática por el algoritmo de inundación recursiva, se *amoldan* convenientemente a la distribución observada en los gráficos de casos caídos en cada neurona para las transiciones L-H.

En la tabla 4.2 se han buscado aquellos segmentos donde el número de patrones correspondientes al momento de la transición L-H, caídos en las neuronas que forman cada uno de los segmentos, es relevante.

Para cada una de las redes seleccionadas, se procede a comparar señal a señal, en la capa primera, las neuronas activadas para cada grupo de descargas. Se considerarán *buenas* aquellas señales que distingan claramente los grupos hallados por la SOM de la segunda capa, y que presenten pocos casos para los cuales se solapan los grupos de la segunda capa. En definitiva, se intentan encontrar aquellas señales que son discriminantes de los diferentes grupos hallados por la segmentación automática.

Mediante un análisis gráfico, se comparan las neuronas activadas en las redes de la capa 1 para cada grupo de neuronas detectado en la capa 2,

<i>Id SOM</i>	<i>Dimensión</i>	<i>Concordancia entre casos caídos en las neuronas y las su- perficies de Clusot</i>	<i>Factor de inundación</i>	<i>Segmentos encontra- dos</i>
307	5 x 5	Buena	0.8	2
		Regular	0.75, 0.7, 0.6	1
307	5 x 5	Regular	0.9	2
		Regular	0.8, 0.7	1
309	5 x 5	Regular	0.9, 0.8, 0.2	1
310	7 x 5	Regular	0.9	2
		Regular	0.8	2
		Regular	0.7	2
311	7 x 5	Regular	0.9	2
		Regular	0.8	2
		Regular	0.7	1
312	7 x 5	Regular	0.9	1
		Regular	0.8	2
		Regular	0.7	2
		Buena	0.6	2
313	7 x 7	Regular	0.9	2
		Regular	0.8	2
		Regular	0.7	2
		Buena	0.6	2
315	7 x 7	Regular	0.9	3
		Regular	0.8	3
		Regular	0.7	3
		Regular	0.6	2
		Regular	0.5	2

Tabla 4.2: Algunas configuraciones y resultados de la SOM de la segunda capa. Se han señalado en negrita aquellas elegidas para estudiar las transiciones L-H.

obteniendo ejemplos como los mostrados en la figura 4.3. El fin es, de una manera cualitativa, observar qué señales son discriminantes en cada segmento hallado por THEFUMO, ya que en la SOM asociada a esa señal en la primera capa de la MSOM, se activan las neuronas para los patrones de entrada. Y estos se pueden relacionar fácilmente, con los dos grupos detectados por la SOM de salida (estas neuronas se representan en la figura 4.3 con colores rojo y azul). Y, por contra, aquellas señales cuyas neuronas activadas no son capaces de distinguir los dos grupos de la SOM 313, sino que hay “mezcla” serán representadas con colores *intermedios*, siendo el verde cuando la mezcla es en igual número entre patrones que activan el primer segmento, como los que activan el segundo segmento. En la tabla 4.3, se especifican las señales, las redes SOMs usadas en la primera capa de la MSOM y la *apreciación* de si la señal observada puede ser discriminante en cuanto a los dos segmentos donde se han agrupado principalmente las transiciones L-H.

Antes de analizar los resultados obtenidos, se puede ver otro ejemplo de estructura de THEFUMO, con otras redes SOM formando la primera y segunda capa de la MSOM. En concreto, la SOM de salida será la que tiene el $id = 313$.

4.2.5.2. Sustitución de la primera capa de MSOM por una red OJA

Una de las pruebas que se están realizando, es la de sustituir la primera capa de la MSOM por una red neuronal de OJA [Oja92].

Las redes de OJA se utilizan para realizar Análisis de Componentes Principales (PCA - *Principal Component Analysis*). Se trata de una red directa (feedforward) de una sola capa capaz de extraer las componentes principales de los vectores de entrada.

El Análisis de Componentes Principales es una técnica esencial para la comprensión de los datos y para la extracción de características. Se trata de un método para reducir el número de variables de entrada, descartando

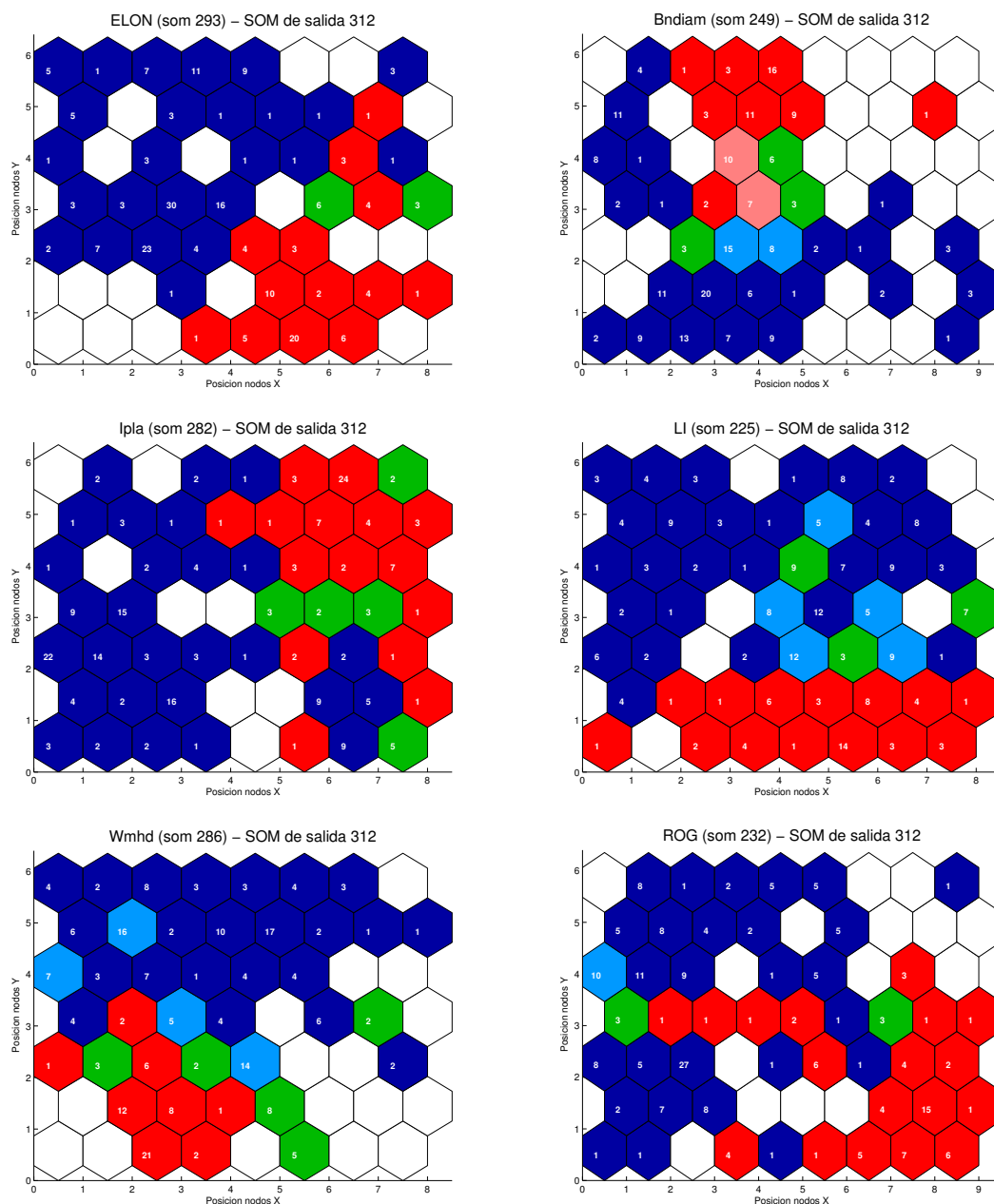


Figura 4.3: Neuronas activadas en las SOMs del primer nivel. El color rojo indica que activan el grupo 1 de la SOM de salida (segundo nivel). La coloreadas en rojo son para aquellas neuronas que activan el grupo 2 de la SOM de salida. El color verde indica que existe “mezcla” para los dos segmentos encontrados en la capa de salida de la MSOM. El número en cada una de las celdas indican el número de transiciones L-H que han activado esa neurona.

<i>Señal</i>	<i>Id SOM</i>	<i>Distinción entre los dos segmentos encontrados en la SOM de salida (312)</i>
Bndiam	249	Muy Buena
Bt	250	Muy Buena
Density	251	Buena
FDWDT	252	Muy Buena
Ipla	282	Muy Buena
LI	225	Muy Buena
Ptot	284	Regular
Q95	227	Regular
Wmhd	286	Buena
Rad	229	Muy Buena
RXPL	230	Regular
ZXPL	289	Excelente
RSIL	299	Excelente
ZSIL	271	Excelente
RSOL	243	Excelente
ZSOL	273	Muy Buena
ROG	232	Excelente
RIG	262	Buena
TOG	263	Muy Buena
ELON	293	Excelente
TRIL	294	Regular
TRIU	295	Buena
Te	247	Buena
TI	267	Mala
DalphaIn	239	Buena
LID4	298	Muy Buena
PRFL	303	Mala
GradientTI	246	Mala
GradientTe	277	Buena

Tabla 4.3: Relación de señales y SOMs usadas en la primera capa de MSOM, correspondientes a la segunda capa de la MSOM compuesta por la SOM 312. Para cada SOM y señal, la capacidad para discriminar los dos segmentos encontrados en la SOM de salida se cualificó con *Mala*, *Regular*, *Buena*, *Muy Buena* y *Excelente*.

aquellas combinaciones lineales que provocan pequeñas variaciones y dejando aquellas que son las responsables de las variaciones importantes.

Se ha creado una red de Oja de 60 componentes principales, es decir, se reducen las 2.500 variables de entrada originales (en el caso de 25 señales por 100 muestras por señal) a tan solo 60. En la figura ?? se representa la nueva arquitectura de THEFUMO con la red OJA.

El principal problema, radica en hallar cuáles son las señales originales que más influyen en los segmentos que agrupan la mayoría de transiciones L-H en la SOM de salida. Veámoslo con un ejemplo: en la figura 4.4 se representan las neuronas de la SOM de salida activadas por las transiciones L-H.

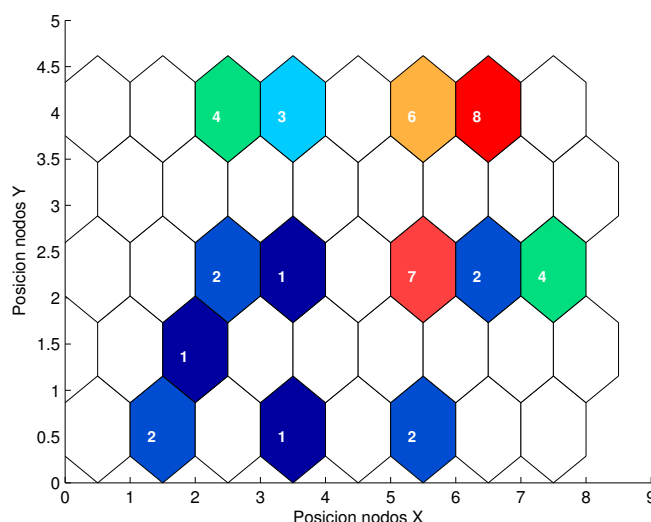


Figura 4.4: Neuronas activadas por las transiciones L-H en la SOM de salida (22105), usando como capa de entrada un red OJA (102) de 60 neuronas.

Una vez identificadas cuáles son las neuronas activadas por las señales en el momento de la transición L-H, en la figura 4.4 serían las clases 22, 38 y 39. Nótese que vuelven a aparecer dos grupos diferenciados de transiciones L-H: por un lado la neurona 22 y por otro las vecinas 38 y 39.

Uno de los procedimientos que se pueden seguir para “subir” en la estructura de THEFUMO hasta llegar a las señales de entrada que más están influyendo en que se active un grupo de transiciones L-H u otro puede ser:

1. Comparar los vectores de pesos de las 3 neuronas que más transiciones L-H tienen. Estos vectores estarán compuestos por 60 variables, que son las 60 componentes principales halladas por OJA. Se comprueba qué valores del vector de pesos son más diferentes entre sí, siendo estos los que pueden discriminar entre una neurona y otra.
2. A partir de las componentes principales o variables del vector de pesos hallados como “discriminantes”, se analizan los pesos (θ_s) de las componentes principales. Pesos que ponderan cada una de las variables de entrada, que son en total 2.500, repartidas en segmentos de 100 muestras de cada una de las señales monitorizadas del dispositivo experimental de fusión.

Para ayudar en este tedioso proceso, se pueden utilizar gráficas, como la 4.5 donde se pueda visualizar qué thetas tienen los valores más elevados. También se aprecia los segmentos de 100 valores de thetas que corresponden a las 100 muestras de una determinada señal.

En las tablas 4.4 y 4.5 se pueden comprobar los resultados de los dos pasos anteriores.

En la siguiente sección, se presentan algunas conclusiones que se obtienen de este trabajo de creación de una estructura para el sistema THEFUMO donde se usa una red OJA en vez de una primera capa formada por redes SOM.

4.2.5.3. Primeras conclusiones obtenidas con THEFUMO

Realizando los estudios mostrados en las anteriores secciones para diferentes configuraciones de las SOMs que componen las dos capas usadas en THEFUMO, se pueden obtener unas primeras conclusiones, a la *espera* de que esta información sea evaluada por los físicos e intenten dar el significado físico a los diferentes segmentos y señales con las que se ha trabajado.

Algunas de las conclusiones obtenidas son:

- En la tabla 4.2 se puede comprobar que el número de segmentos siem-

pre es entre 1 y 3. Siendo el valor que más se repite el 2. Con lo cual, se podría tener una primera conclusión hallada por THEFUMO: parece que existen 2 tipos de transiciones L-H o al menos es lo más probable.

- En una mayoría de modelos se observa una distribución de las transiciones L-H en dos grupos. Existen varias señales que suelen permanecer como influyentes en una gran mayoría de los modelos desarrollados: RXPL, ZXPL, ROG, RIG y TRIL son las más constantes en este sentido.
- El sistema MSOM desarrollado no solo es válido para la clasificación de los posibles grupos de transiciones L-H, sino para la clasificación de cualquier instante de tiempo dentro de los intervalos estudiados. Dicho de otra forma: si en vez de estudiar la transición L-H, se desea conocer con más detalle, por ejemplo, un instante 200 milisegundos anterior (o posterior) a dicha transición, este mismo sistema resultará válido sin necesidad de entrenamiento adicional.
- En las pruebas que se han hecho con la red OJA (ver 4.2.5.2), se comprueba que muchas de las señales influyentes en los, típicamente, dos segmentos o clases de transiciones L-H se repiten con las señales halladas con la primera capa de la MSOM, compuesta por una SOM por cada señal tratada. En la tabla 4.6 se muestran las cinco señales que en la estructura original de THEFUMO más veces se repiten como influyentes o discriminantes de los diferentes segmentos de transiciones L-H; comparándose con los resultados obtenidos con OJA. Se comprueba que todas excepto ZXPL coinciden en ambas estructuras.
El problema de usar una red OJA es la complejidad para, a partir de los segmentos hallados por la segmentación automática, ver qué variables de entrada son las más influyentes para distinguir ambos comportamientos.
- THEFUMO se está creando como una herramienta de minería de datos para los físicos que trabajan en el desarrollo de la energía por fusión

termonuclear, son ellos los que tienen que validar y dotar de significado físico los segmentos y señales relevantes hallados por el sistema.

Este es el paso actual en el cual se encuentra el desarrollo, siendo muy importante los resultados que se obtengan para poder ir ajustando las capacidades de THEFUMO a las necesidades de la ciencia de la fusión.

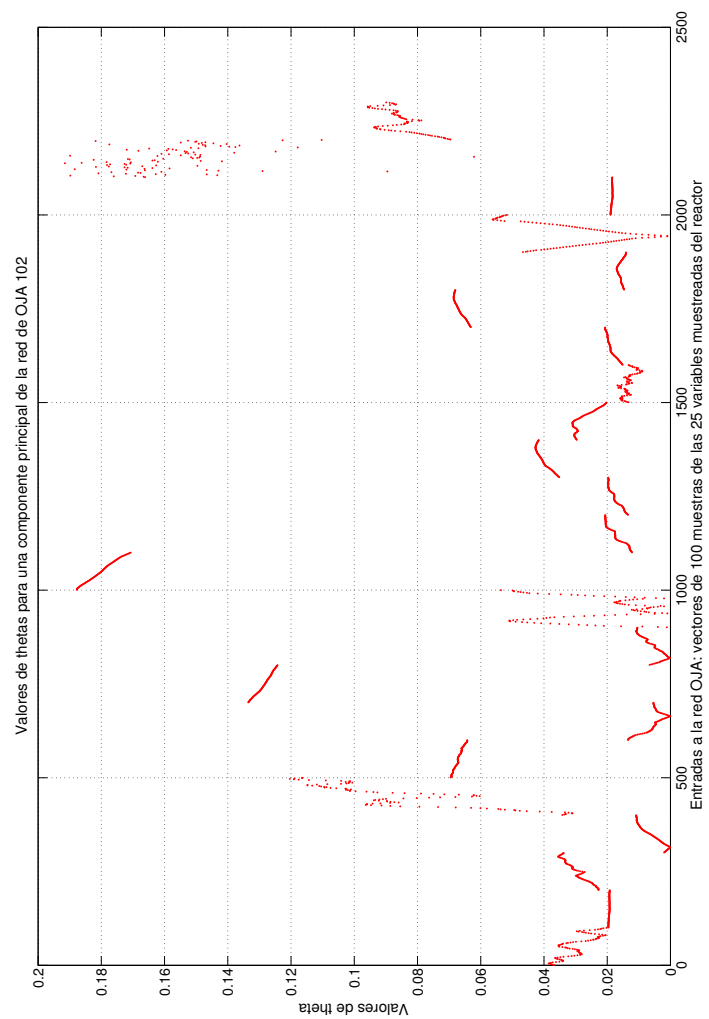


Figura 4.5: Valores de thetas para una componente principal de la red OJA 102.

Componente P.	W(22)	W(38)	W(39)
PC1	1.0981	-0.9114	-0.5306
PC2	0.1075	0.8618	0.5758
PC3	1.1862	-0.8471	-0.8471
PC4	-0.9713	1.0330	-0.0665
PC5	-0.9479	-0.6541	1.0584
PC6	1.0405	-0.9671	0.6868
PC7	0.6648	-0.8229	-0.5858
PC8	0.5966	0.8425	0.8555
PC9	-0.3887	-0.9416	-0.1892
PC10	-0.0686	-0.5231	-0.8056
PC11	0.4352	0.8446	0.8446
PC12	0.5564	0.3434	0.9045
PC13	-1.0946	0.4238	0.5813
PC14	-0.7809	0.9271	-0.6412
PC15	0.9280	0.8799	-0.7462
PC16	-0.1804	-0.6625	-0.9652
PC17	-1.1451	0.5859	0.0620
PC18	1.0392	-0.9809	0.3277
PC19	-0.5882	-0.9208	0.4424
PC20	0.2543	-0.8603	0.4146
PC21	-0.5059	0.9937	-0.6794
PC22	-1.1198	0.9119	0.5040
PC23	-1.1853	-0.3764	0.6872
PC24	0.9831	0.8598	0.9830
PC25	-0.9166	0.9654	0.9860
PC26	-0.9693	0.9217	0.3974
PC27	-0.9597	0.9324	0.4954
PC28	-0.8922	0.4049	0.6355
PC29	1.1372	-0.9235	-0.6698
PC30	0.0580	0.7313	0.0481
PC31	1.0813	-0.8937	-0.8937
PC32	0.9890	0.9267	-0.6450
PC33	-1.1424	0.9173	-0.4993
PC34	-0.2711	0.9358	0.4266
PC35	0.8764	-0.2312	-1,2
PC36	-0.8175	0.6842	0.3484
PC37	-0.9246	0.8270	0.5156
PC38	1.0550	0.6525	-0.4073
PC39	1.1355	-0.9178	-0.8974
PC40	-0.9920	0.1007	-0.8210
PC41	0.2335	-0.8916	-0.8915
PC42	0.7572	-0.8843	-0.3471
PC43	-0.6587	-0.8817	-0.8818
PC44	-0.1846	-1.0036	-1.0019
PC45	0,78	-0.1168	-0.7821
PC46	-1.0256	1.0231	0.7547
PC47	1.0734	0.8870	-0.9823
PC48	-0.9596	-0.9608	-0.8500
PC49	1.1686	-0.8346	-0.8345
PC50	0.3464	-0.7939	-0.9189
PC51	0.6764	0.6399	0.9455
PC52	-0.2276	0.9020	0.9020
PC53	-1.1949	0.8200	0.6311
PC54	1.1623	-0.8767	-0.8767
PC55	-0.6344	1.1186	-0.0582
PC56	0.9923	-1.0098	0.3754
PC57	0.2670	0.2895	-1.0791
PC58	-1.1969	0.8519	0.8311
PC59	-1.1689	-0.2257	0.8371
PC60	-1.2129	0.8519	0,42

Tabla 4.4: Vector de pesos de las neuronas 22, 38 y 39 de la SOM de salida 22105. Cada variables del vector de pesos corresponde a una de las 60 componentes principales que componen la salida de la red OJA (102). Se han señalado en negrita aquellos que pesos que difieren más entre las tres neuronas.

Señal	Nº de veces relevante
DalphaIn	8
TI	5
Rad	5
Ipla	5
ELON	3
Q95	3
TRIU	1
TRIL	1
TOG	1
RIG	1
ROG	1
TeO	1
LI	1

Tabla 4.5: Señales con valores de θ más elevados en las componentes principales que influyen en activar la neuronas de la SOM de salida donde se concentran las transiciones L-H. Se indica el número de veces que esa seña ha resultado relevante para alguna de las estructuras de THEFUMO con red OJA que se han probado.

Señales influyentes con MSOM	¿Coincide con OJA?
BnDiam	Sí
Q95	Sí
ZXPL	No
ROG	Sí
TRIL	Sí

Tabla 4.6: Comparación de algunas señales halladas como relevantes en THEFUMO con arquitectura MSOM y si coinciden con las halladas con arquitectura OJA.

Capítulo 5

Conclusiones y desarrollos futuros

A partir de lo desarrollado en esta tesis, se pueden obtener algunas conclusiones sobre la capacidad de AIPAKA de proveer de un entorno de trabajo que permite la aplicación de técnicas de Aprendizaje Automático en gran número de situaciones, tanto en el campo industrial como en los ámbitos más científicos o experimentales. Posibilitando trabajar con estas técnicas incluso a personas que no tengan un conocimiento o experiencia relevantes sobre Inteligencia Artificial o el modelado predictivo.

De una manera estructurada, las conclusiones más relevantes y trabajos futuros a los que puede dar lugar esta tesis se enumeran a continuación.

Como **conclusiones** destacadas se presentan las siguientes:

- Disponer de un procedimiento de trabajo, para la aplicación de las capacidades del Aprendizaje Automático en gran número de problemas. Y accesible a un amplio número de profesionales e investigadores no necesariamente relacionados directamente con los campos de la Inteligencia Artificial. Pero que se pueden beneficiar de las capacidades de estas técnicas de inferencia de conocimiento en sus respectivas profesiones o campos de investigación.

No se trata tanto de que con AIPAKA se construyan sistemas *inte-*

ligentes a partir de la experiencia acumulada en las diferentes fuentes de datos, haciendo innecesaria la labor de los analistas. Sino, al contrario, hacer más eficiente y productiva la labor de los profesionales, al facilitarles el estudio de la realidad reflejada en los datos. Esto se ve especialmente, cuanto más complejos son los sistemas abordados, intentar explicarlos recurriendo a las leyes fundamentales que los gobiernan puede ser en la práctica imposible. Debido al número de variables involucradas y a la complejidad de sus relaciones. Sin embargo, si se dispone de un procedimiento para estructurar y modelar la información disponible, modelando los diferentes comportamientos observados; se facilitará en gran medida la labor de los analistas e investigadores expertos en el campo de actuación del sistema. Este trabajo con sistemas complejos con el fin de ayudar a los expertos en el campo del saber al que pertenece el sistema, se puede ver en esta tesis en los casos de aplicación a los dispositivos de fusión (4.2), determinación de los riesgos en las solicitudes y concesiones de créditos, etc.

- Estructuración de los procesos de modelado, permitiendo la trazabilidad, repetibilidad, monitorización, implantación, mantenimiento y desarrollo.

El proceso de creación de un modelo predictivo, mediante la aplicación de algoritmos de Aprendizaje Automático puede ser más o menos complejo. Pero no hay que olvidar, que ese proceso debe ser trazable y repetible. No vale de nada, en un extremo, hallar un modelo y no ser capaces de conocer exactamente cómo se calcularon las entradas, de qué variables se obtuvieron originalmente, cómo fue parametrizado el algoritmo de aprendizaje, etc. Se estaría ante *un caso único de modelo* que dejaría de ser válido a lo largo del tiempo, siendo muy difícil repetirlo y adaptarlo a los cambios ocurridos en el sistema real.

En muchas ocasiones, la persona que realiza el modelado no es quien implanta ese modelo software en el proceso o sistema de producción

final. Por tanto, los modelos deben poder ser implantados fuera de la plataforma o sistema que se usó para su creación, es decir, en los sistemas de producción industriales o de experimentación. Permitiendo la monitorización de los modelos desplegados para asegurar que el rendimiento es el mismo en el entorno de producción que en el de desarrollo. Por último, ningún modelo predictivo de sistemas complejos va a permanecer *estable* por mucho tiempo, siempre habrá cambios en el sistema modelado o ajustes que se quieran hacer, como por ejemplo, añadir nuevas variables que han aparecido durante la operación del modelo y que no estaban disponibles en el entrenamiento de éste. AIPAKA facilita enormemente las tareas anteriores, ya que, provee de un marco estructurado desde la fase inicial del modelado hasta que los modelos son liberados para su despliegue o integración en entornos de producción o experimentales. Conociendo en todo momento cuál ha sido el camino seguido para la obtención del modelo: datos usados, parámetros de entrenamiento, medición del rendimiento y validación de los modelos creados.

- Durante la elaboración de AIPAKA se han creado una serie de herramientas muy útiles para la estandarización de los procesos de modelado, y para la propia creación de los modelos predictivos. Son de destacar aquellas capacidades desarrolladas para el trabajo y manejo de las variables disponibles:
 - *Entradas Potenciales*. Es algo fundamental en AIPAKA poder tener la capacidad de definir todas las variables, que pueden ser usadas como entradas a un modelo para cualquier clase de problema. Cuestión que no estaba resuelta y para la cual se ha creado el mecanismo de las *Entradas Potenciales*. Convirtiéndose en una potente herramienta, fácilmente implementable en cualquier lenguaje de programación. Las *Entradas Potenciales* permiten la organización de las variables, su cálculo, el uso en los vectores

de entrada y la trazabilidad de todo el proceso: desde las primeras descargas de las bases de datos originales, hasta las variables finalmente usadas en los entornos de producción.

- Detección de la importancia de variables. Se ha trabajado en varias vías para la selección características y variables ([ABCP09] y [PVM⁺15]), detectar qué variables son independientes y cuáles combinaciones de otras en espacios de muy alta dimensionalidad [FP15]. Estos trabajos son importantes en la fase de *Modelización* de AIPAKA, pero también son muy relevantes para cualquier problema que se aborde mediante el uso de las técnicas de Aprendizaje Automático. Donde siempre se tendrá que decidir qué variables son las que se usan de todas las disponibles, evitando perder información relevante y tampoco caer en problemas de sobre-ajuste o tener que manejar variables que no son relevantes para el problema que se quiere solucionar.
- AIPAKA permite focalizar los esfuerzos, de las empresas e instituciones dedicadas a la investigación, en la obtención de datos relevantes y en su area de conocimiento específico. Sin gastar recursos en el campo de la Inteligencia Artificial, que puede ser un ámbito que no resulte de su interés principal. Permitiendo avanzar más rápidamente en aquellas áreas en las cuales son diferenciales y centrándose en el núcleo de su trabajo de investigación o negocio.
En definitiva, AIPAKA permite concentrar los esfuerzos en aquellos puntos que más valor aportan o son más complejos de resolver. Al tener “resueltos” gran número de los problemas de modelado predictivo, segmentación y clasificación a partir de grandes repositorios de datos.
- Modularidad. Permite crear sistemas complejos a partir de módulos más simples. AIPAKA ayuda a afrontar grandes problemas, dividiéndolos en partes y tratando cada parte de manera estructurada dependiendo de su tipo. Lo anterior, es una ventaja no sólo a la hora de crear

el sistema sino también en el momento del despliegue en producción y en el mantenimiento posterior. Al tener un sistema modular se puede optimizar el despliegue, asignando a cada módulo aquellos recursos que lo optimicen y evolucionar cada módulo por separado haciendo el mantenimiento más sencillo y óptimo.

Esta modularidad también facilita la aplicación de las capacidades actuales de procesamiento distribuido, permitiendo optimizar los tiempos de ejecución de las técnicas de Inteligencia Artificial. Las cuales suelen ser computacionalmente muy costosas durante la ejecución del entrenamiento de los modelos.

- Sistema abierto y escalable. A lo largo de los años de trabajo, que han dado lugar a esta tesis se han ido incorporando y afinando algoritmos basados en técnicas de Inteligencia Artificial. No se imponen límites al tipo de técnicas a usar para resolver cada clase de problema: segmentación, clasificación y predicción. Por otro lado, tampoco está restringido el número de partes que pueden formar un sistema dado. Permitiendo elegir la arquitectura de módulos del sistema que sea más adecuada para dar una solución óptima.

Por último, AIPAKA es un sistema abierto en cuanto a facilitar la integración de los modelos desarrollados en los sistemas finales, pudiendo en todo momento tener claras cuáles son las variables de entrada que se han usado y cómo se han obtenido o calculado (recordemos las entradas potenciales).

- Prototipado rápido a partir de grandes cantidades de datos. Permite comprobar, entre otras cosas, qué variables de las muchas que se guardan son realmente relevantes. AIPAKA habilita mecanismos para la explotación rápida de la información guardada con el fin de resolver los problemas que se plantean. No se quedará en un marco teórico de *ciclo de trabajo* sino que se dan herramientas prácticas para construir modelos predictivos por Aprendizaje Automático. También se huye

del mero *guardar mucha información porque se puede hacer*, cayendo en una especie de síndrome de Diógenes moderno. Sino que se guarda información porque se puede usar para avanzar en los objetivos planteados por el negocio o por la investigación.

Por otra parte, están las **líneas futuras de trabajo** surgidas a partir de las líneas de investigación llevadas a cabo durante la elaboración de esta tesis, quizás las más interesantes sean las siguientes:

- Automatización completa del procedimiento. AIPAKA puede permitir la creación automática de modelos predictivos a partir de las fuentes de datos y de los objetivos a alcanzar. Habría que crear un meta-lenguaje con el que se permitiera definir las diferentes fuentes de datos, tipo de variables que contienen, y cuál es el objetivo que se quiere conseguir. A partir de ahí el sistema crearía los modelos que cumplen con el objetivo marcado usando para ello el Aprendizaje Automático.

Una parte importante de esta automatización es la integración de las *Entradas Potenciales* en ella, junto con la capacidad de elegir automáticamente las mejores variables para componer los vectores de entrada de todas las disponibles.

Algunos avances en este sentido para ir automatizando el proceso ya se han hecho. Por ejemplo, en los siguientes dos casos:

- Implementación del Razonamiento Basado en Casos. A partir de los datos usados, el algoritmo *distingue* si se encuentra ante un caso de clasificación, o de estimación de valores continuos. Estimando si la salida es una variable discreta, en cuyo caso sería una clasificación, o es una variable continua, por lo que se trataría de una predicción; creando en consecuencia el modelo más adecuado.
- Predicción de indicadores a partir de series temporales. Donde a partir de la serie temporal histórica del indicador que se quiere predecir y el horizonte temporal con el adelanto requerido, el algoritmo de aprendizaje construye la estructura del modelo y

ajusta los parámetros del mismo. Lo anterior ya se ha implementado con éxito en alguno caso, por ejemplo, en los de predicción de la demanda (uno de ellos ha sido la predicción de llamadas entrantes).

- Fase de despliegue. En la figura donde se muestra el esquema de AIPAKA, el último nivel o fase, es el de *despliegue*. Esta fase no ha sido desarrollada en esta tesis. Se trataría de desplegar los modelos hallados en producción, interactuando con los sistemas y procesos finales para los cuales fueron construidos. O llevarlos a entornos de explotación, aunque en un principio sólo fueron desarrollados para realizar estudios experimentales.

En el despliegue es importante tener en cuenta las siguientes cuestiones:

- Independiente de herramientas analíticas que no estén disponibles en los entornos de producción. Los entornos de análisis y desarrollo de modelos, en una mayoría de casos, son diferentes a los entornos de producción u operación. Esto genera una serie de problemáticas a la hora de “extrapolar” modelos obtenidos por Aprendizaje Automático desde la herramienta analítica usada a producción. En la mayoría de casos, se requiere un análisis y desarrollos específicos para *exportar* los modelos al entorno de producción definitivo.
- Capacidad de comprobar que realmente lo que se ha desplegado es el mismo modelo creado. Parece una obviedad, pero en muchas ocasiones es realmente complejo llevarlo a la práctica. Hay que prever las herramientas de chequeo suficientes para asegurar que el despliegue ha sido correcto y que los sistemas creados en el *laboratorio* son los que realmente se han desplegado.
- Interacción con otros sistemas. Actualmente muy pocas aplicaciones trabajan de manera completamente aislada. Normalmente

necesitan interactuar con otras. En definitiva, se debiera proveer de interfaces de comunicaciones para poder interconectarse de manera fácil con otros sistemas. Se ha presentado un caso de cómo sistemas basados en modelos predictivos pueden comunicarse con el resto de aplicaciones clientes y fuentes de datos a través de unas interfaces web (o usando el término inglés: Web APIs).

- Ajuste en tiempo real. Las realidades modeladas son muy dinámicas. Ese dinamismo se verá reflejado en los datos, obligando a los modelos desplegados a ajustarse para seguir manteniendo el rendimiento original e incluso mejorarlo. No hay que olvidar que la base es el Aprendizaje Automático, y por tanto hay que llegar a aprender de las entradas que se están produciendo durante la operación del sistema.
- Protección de la propiedad intelectual y de los derechos de explotación. Todo software desplegado en entornos operativos debe instalarse conociendo cuáles son los requerimientos de uso, incluida la parte legal. Hay que definir si es una licencia, un servicio o proyecto llave en mano. Según lo anterior, puede tener más o menos connotaciones en el despliegue de los modelos. Por ejemplo, ante una licencia habría que definir el software licenciado, para cuántos usuarios o instancias en las máquinas, etc.

Desde el punto de vista de la ciencia tiene poca relevancia más allá de proteger la propiedad intelectual de los investigadores.

- Lenguaje natural. Incorporar análisis provenientes de documentación o información escritas en lenguaje natural. Yendo más allá de proveer simplemente una serie de librerías para la Minería de Textos (del término inglés *Text Mining*); llegando a construir los mecanismos para proveer de una *validación* adecuada, que mida la bondad de estos modelos. En otras palabras, que se pueda validar el rendimiento de los reconocedores semánticos construidos. Esto se podría hacer extrapolando los conjuntos de entrenamiento, prueba y nuevo a este contexto,

creando un *corpus* de textos previamente etiquetados que posibilite un entrenamiento supervisado y la medida de la bondad de las conclusiones obtenidas por los modelos de reconocimiento semántico.

Desde otro punto de vista, también podría proveer a AIPAKA de una interfaz en lenguaje natural, mediante el que interactuar para facilitar la creación de los modelos.

- Desarrollo de mecanismos de optimización compleja. La idea es desarrollar una capa de optimización transversal a las fases de AIPAKA, sobre todo a aquellas fases de datos y modelización. Esta capa recibiría las entradas de los modelos predictivos (modelado) y de los datos; añadiendo otras entradas para definir los objetivos que se deben alcanzar y las restricciones del sistema. Dando como resultado una *función de coste* que haya que minimizar (o maximizar), encargándose la optimización de encontrar una solución óptima que ajuste lo máximo posible los objetivos, sin violar las restricciones, teniendo en cuenta las predicciones y estimaciones de los modelos, y manteniendo siempre un coste mínimo.

Las aplicaciones son innumerables, pensemos en algunos de los casos que hemos visto en el capítulo 4:

- Detección de oportunidades comerciales: cómo se optimizan las rutas de los comerciales, rentabilidad del cliente y del producto, optimización de sus procesos productivos, etc.
- Predicción del riesgo, cómo hacemos que el cliente sea lo más rentable posible: cuáles son los parámetros óptimos del crédito para que sea aceptado por el cliente y, a la vez, optimice la rentabilidad para el banco: tiempo de amortización, intereses, etc.
- Optimización del número de agentes que deben atender las llamadas en el CAC o en qué tipos de llamadas o colas serán más necesarios.

Lo anterior son sólo tres ejemplos de los muchos en los que se podría

aplicar la optimización. Para que pudiera ser factible esta aplicación es necesario haber desarrollado mecanismos de paralelización y cálculo masivo sobre las diferentes técnicas de modelización y optimización en las que se apoya AIPAKA. Ya que los problemas de optimización son computacionalmente muy costosos.

Bibliografía

- [A⁺03] Hussein Abbass et al. Pareto neuro-evolution: Constructing ensemble of neural networks using multi-objective optimization. In *Evolutionary Computation, 2003. CEC'03. The 2003 Congress on*, volume 3, pages 2074–2080. IEEE, 2003.
- [AAH15] Muhammad Arif, Khubaib Amjad Alam, and Mehdi Hussain. Application of data mining using artificial neural network: Survey. *International Journal of Database Theory and Application*, 8(1):245–270, 2015.
- [ABCP09] Manuel R Arahal, Manuel Berenguel, Eduardo F Camacho, and Fernando Pavón. Selección de variables en la predicción de llamadas en un centro de atención telefónica. *Revista Iberoamericana de Automática e Informática Industrial RIAI*, 6(1):94–104, 2009.
- [ABPC03] MR Arahal, M Berenguel, F Pavón, and EF Camacho. Comparing the performance of some neural fraud detectors in telecommunications. 2003.
- [AC95] Bruce H Andrews and Shawn M Cunningham. Ll bean improves call-center forecasting. *Interfaces*, 25(6):1–13, 1995.
- [ACC02] Manuel R Arahal, Alfonso Cepeda, and FE Camacho. Input variable selection for forecasting models. In *15éme IFAC World*

- Congress on Automatic Control, Barcelone, Espagne juillet, 2002.*
- [AdTR08] Eliana Angelini, Giacomo di Tollo, and Andrea Roli. A neural network approach for credit risk evaluation. *The quarterly review of economics and finance*, 48(4):733–755, 2008.
- [Aka74] H. Akaike. A new look at the statistical model identification. *IEEE Transactions on Automatic Control*, 19:716–723, 1974.
- [Ali12] Aleem Ali. A concise artificial neural network in data mining. *International Journal of Research in Engineering & Applied Sciences*, 2(2):418–428, 2012.
- [AM⁺10] Javad Abdi, Behzad Moshiri, et al. Comparison of rbf and mlp neural networks in short-term traffic flow forecasting. In *Power, Control and Embedded Systems (ICPCES), 2010 International Conference on*, pages 1–4. IEEE, 2010.
- [APC04] Manuel R Arahall, Fernando Pavon, and Eduardo F Camacho. Models for incoming calls forecasting in a customer attention center. In *System Identification (SYSID’03): A Proceedings Volume from the 13th IFAC Symposium on System Identification, Rotterdam, the Netherlands, 27-29 August 2003*, page 205. Elsevier, 2004.
- [AS99] Javier Aracil Santonja. ¿ es menester que los ingenieros filosofen? *Argumentos de razón técnica: Revista española de ciencia, tecnología y sociedad, y filosofía de la tecnología*, (2):29–50, 1999.
- [Ati01] Amir F Atiya. Bankruptcy prediction for credit risk using neural networks: A survey and new results. *Neural Networks, IEEE Transactions on*, 12(4):929–935, 2001.

- [BA91] Sherif M Botros and Christopher G Atkeson. Generalization properties of radial basis functions. In *Advances in Neural Information Processing Systems*, pages 707–713, 1991.
- [Bar93] A. R. Barron. Universal approximation bounds for Superpositions of a Sigmoidal Function. *IEEE Transactions on Information Theory*, 39:930–944, 1993.
- [BB01] Michele Banko and Eric Brill. Scaling to Very Very Large Corpora for Natural Language Disambiguation. In *ACL*, pages 26–33. Morgan Kaufmann Publishers, 2001.
- [BBR08] D. Brugger, M. Bogdan, and W. Rosenstiel. Automatic Cluster Detection in Kohonen’s SOM. *Neural Networks, IEEE Transactions on*, 19(3):442–459, March 2008.
- [bbv13] Big Data. Now’s the time to create business value with data. Technical report, June 2013.
- [BdTP07] Claudia Biancotti, Leandro d’Aurizio, and Raffaele Tartaglia-Polcini. A neural network architecture for data editing in the bank of italy’s business surveys. *Bank of Italy Economic Research Paper*, (612), 2007.
- [BE67] Leonard E. Baum and J. A. Eagon. An inequality with applications to statistical estimation for probabilistic functions of Markov processes and to a model for ecology. *Bulletin of the American Mathematical Society*, 73(3):360–363, May 1967.
- [Bel78] Richard Bellman. *An introduction to artificial intelligence. Can computers think?* Boyd & Fraser Pub. Co., 1978.
- [Ben09] Yoshua Bengio. Learning Deep Architectures for AI. *Foundations and Trends in Machine Learning*, 2(1):1–127, 2009.

-
- [BFG95] Monica Bianchini, Paolo Frasconi, and Marco Gori. Learning without local minima in radial basis function networks. *Neural Networks, IEEE Transactions on*, 6(3):749–756, 1995.
- [BH69] A. E. Bryson and Y. C. Ho. *Applied Optimal Control*. Blaisdell, New York, 1969.
- [BP66] Leonard E. Baum and Ted Petrie. Statistical Inference For Probabilistic Functions Of Finite State Markov Chains. *Ann. Math. Statist.*, 37(6):1554–1563, 1966.
- [BR01] Martin Bogdan and Wolfgang Rosenstiel. Detection of cluster in Self-Organizing Maps for controlling a prostheses using nerve signals. In *ESANN*, pages 131–136, 2001.
- [BRSY03] AM Bagirov, AM Rubinov, NV Soukhoroukova, and J Yearwood. Unsupervised and supervised data classification via nonsmooth and global optimization. *Top*, 11(1):1–75, 2003.
- [BSA83] A. G. Barto, R. S. Sutton, and C. W. Anderson. Neuronlike adaptive elements that can solve difficult learning control problems. *IEEE Transactions on Systems, Man, and Cybernetics*, (13), 1983.
- [BSV98] Barbro Back, Kaka Sere, and Hannu Vanharanta. Analyzing financial performance with self-organizing maps. In *Neural Networks Proceedings, 1998. IEEE World Congress on Computational Intelligence. The 1998 IEEE International Joint Conference on*, volume 1, pages 266–270. IEEE, 1998.
- [Cau93] Gert Cauwenberghs. A Learning Analog Neural Network Chip with Continuous-Time Recurrent Dynamics. In Jack D. Cowan, Gerald Tesauro, and Joshua Alspector, editors, *NIPS*, pages 858–865. Morgan Kaufmann, 1993.

- [CCK⁺00] Pete Chapman, Julian Clinton, Randy Kerber, Thomas Khabaza, Thomas Reinartz, Colin Shearer, and Rudiger Wirth. CRISP-DM 1.0 Step-by-step data mining guide. Technical report, The CRISP-DM consortium, August 2000.
- [CdB93] Carlos Serrano Cinca and Bonifacio Martín del Brío. Predicción de la quiebra bancaria mediante el empleo de redes neuronales artificiales. *Revista española de financiación y contabilidad*, pages 153–176, 1993.
- [CDP97] Enrique Cervera and Angel P Del Pobil. Multiple self-organizing maps: a hybrid learning scheme. *Neurocomputing*, 16(4):309–318, 1997.
- [CE97] M. Cox and D. Ellsworth. Managing Big Data for Scientific Visualization. pages 5–1–5–17. ACM Siggraph, 1997.
- [CF99] G. Castellano and A.M. Fanelli. Feature selection: a neural approach. In *Neural Networks, 1999. IJCNN '99. International Joint Conference on*, volume 5, pages 3156–3160 vol.5, 1999.
- [CF00] Giovanna Castellano and Anna Maria Fanelli. Variable selection using neural-network models. *Neurocomputing*, 31(1-4):1–13, 2000.
- [CFM⁺07] Barbara Cannas, A Fanni, A Montisci, G Murgia, P Sonato, and Maria Katiusia Zedda. Dynamic neural networks for prediction of disruptions in tokamaks. In *Proceedings of the 10th International Conference on Engineering Applications of Neural Networks (EANN)*, 2007.
- [CFPS99] Philip K Chan, Wei Fan, Andreas L Prodromidis, and Salvatore J Stolfo. Distributed data mining in credit card fraud detection. *Intelligent Systems and their Applications, IEEE*, 14(6):67–74, 1999.

- [Che85] Peter Cheeseman. In Defense of Probability. In Aravind K. Joshi, editor, *IJCAI*, pages 1002–1009. Morgan Kaufmann, 1985.
- [CLPS02] Michael H Cahill, Diane Lambert, José C Pinheiro, and Don X Sun. Detecting fraud in the real world. In *Handbook of massive data sets*, pages 911–929. Springer, 2002.
- [CM85] Eugene Charniak and Drew McDermott. *Introduction to Artificial Intelligence*. Addison Wesley, 1985.
- [CTYR12] Huanhuan Chen, Peter Tino, Xin Yao, and Ali Rodan. Learning in the Model Space for Fault Diagnosis. *CoRR*, abs/1210.8291, November 2012.
- [dAS03] Jorge de Andrés Sánchez. Dos aplicaciones empíricas de las redes neuronales artificiales a la clasificación y la predicción financiera en el mercado español. *RAE: Revista Asturiana de Economía*, (28):61–87, 2003.
- [DC07] Sanjiv R Das and Mike Y Chen. Yahoo! for amazon: Sentiment extraction from small talk on the web. *Management Science*, 53(9):1375–1388, 2007.
- [DG04] Jeffrey Dean and Sanjay Ghemawat. MapReduce: simplified data processing on large clusters. In *In OSDI’04: Proceedings of the 6th conference on Symposium on Operating Systems Design & Implementation*. USENIX Association, 2004.
- [DGS00] Wlodzislaw Duch, Karol Grudzinski, and G Stawski. Symbolic features in neural networks. In *In Proceedings of the 5th Conference on Neural Networks and Their Applications*. Citeseer, 2000.
- [DGSC97] José R. Dorronsoro, Francisco Ginel, Carmen Sanchez, and Carlos Santa Cruz. Neural fraud detection in credit card ope-

- rations. *IEEE Transactions on Neural Networks*, 8(4):827–834, 1997.
- [DH08] MIRELA Danubianu and V Hapenciuc. Improving customer relationship management in hotel industry by data mining techniques. *Competitiveness and Stability in the Knowledge-Based Economy.-CD.-2008.-Craiova, Romania*, pages 2444–2452, 2008.
- [DHK07] BMM Dissananayake, CH Hendarahewa, and AS Karunananda. Artificial neural network approach to credit risk assessment. In *Industrial and Information Systems, 2007. ICIIS 2007. International Conference on*, pages 301–306. IEEE, 2007.
- [DPG02] Carlotta Domeniconi, Jing Peng, and Dimitrios Gunopulos. Locally Adaptive Metric Nearest-Neighbor Classification. *IEEE Trans. Pattern Anal. Mach. Intell.*, 24(9):1281–1285, 2002.
- [DS14] Ragu Gurumurthy David Schatsky, Craig Muraskin. Desmystifying artificial intelligence what business leaders need to know about cognitive technologies. Technical report, Deloitte University Press, 2014.
- [DVN⁺13] Kristina Dervojeda, Diederik Verzijl, Fabian Nagtegaal, Mark Lengton, and Elco Rouwmaat. Big Data Artificial Intelligence, September 2013.
- [EDT89] S. Eberhardt, T. Duong, and A. Thakoor. Design of parallel hardware neural network systems from custom analog VLSI ‘building block’ chips. In *Neural Networks, 1989. IJCNN., International Joint Conference on*, pages 183–190 vol.2, 1989.
- [EHKW99] Oya Ekin, Peter L Hammer, Alexander Kogan, and Pawel Winter. Distance-based classification methods. *INFOR-OTTAWA-*, 37:337–352, 1999.

- [Ein08] Arnar Ingi Einarsson. *Credit risk modeling*. PhD thesis, Technical University of Denmark, DTU, DK-2800 Kgs. Lyngby, Denmark, 2008.
- [FH95] John A. Flanagan and Martin Hasler. Self-Organising Artificial Neural Networks. In José Mira and Francisco Sandoval Hernández, editors, *IWANN*, volume 930 of *Lecture Notes in Computer Science*, pages 322–329. Springer, 1995.
- [FI08] Fernando Fernández and Pedro Isasi. Local feature weighting in nearest prototype classification. *Neural Networks, IEEE Transactions on*, 19(1):40–53, 2008.
- [FP97] Tom Fawcett and Foster Provost. Adaptive fraud detection. *Data mining and knowledge discovery*, 1(3):291–316, 1997.
- [FP15] Sebastián Dormido Canto Fernando Pavón, J. Vega. Som and feature weights based method for dimensionality reduction in large gauss linear models. volume SLDS 2015, LNAI 9047, pp. 376–385, 2015. A. Gammerman et al. (Eds.), Abril 2015.
- [FPSS96] U. Fayyad, G. Piatetsky-Shapiro, and P. Smyth. From Data Mining to Knowledge Discovery in Databases. *AI Magazine*, pages 37–54, 1996.
- [Fyf93] C. Fyfe. A simple, homogeneous parallel PCA network. In *Neural Networks, 1993. IJCNN '93-Nagoya. Proceedings of 1993 International Joint Conference on*, volume 3, pages 2496–2499 vol.3, Oct 1993.
- [Gar13] Emilia García. Big Data. El reto de tratar de forma efectiva una ingente cantidad de información . *Boletic 65*, April 2013.
- [Gau13] Priyanka Gaur. Neural networks in data mining. *International Journal of Electronics and Computer Science Engineering*, 1(3), 2013.

- [GCBMH09] Salvador García, Jose-Ramon Cano, Ester Bernadó-Mansilla, and Francisco Herrera. Diagnose effective evolutionary prototype selection using an overlapping measure. *International Journal of Pattern Recognition and Artificial Intelligence*, 23(08):1527–1548, 2009.
- [GDCH12] Salvador Garcia, Joaquín Derrac, José Ramón Cano, and Francisco Herrera. Prototype selection for nearest neighbor classification: Taxonomy and empirical study. *Pattern Analysis and Machine Intelligence, IEEE Transactions on*, 34(3):417–435, 2012.
- [GDL03] Muriel Gevrey, Ioannis Dimopoulos, and Sovan Lek. Review and comparison of methods to study the contribution of variables in artificial neural network models. *Ecological Modelling*, 160(3):249–264, 2003.
- [GG07] M Gouvea and Eric Bacconi Gonçalves. Credit risk analysis applying logistic regression, neural networks and genetic algorithms models. In *POMS 18th Annual Conference*, 2007.
- [GP07] Ben Goertzel and Cassio Pennachin, editors. *Artificial General Intelligence*. Cognitive Technologies. Springer, 2007.
- [GR12] John Gantz and David Reintzel. The Digital Universe in 2020: Big Data, Bigger Digital Shadows, and Biggest Growth in the Far Est. Technical report, IDC, December 2012.
- [GrMH13] Alex Graves, Abdel rahman Mohamed, and Geoffrey E. Hinton. Speech recognition with deep recurrent neural networks. In *ICASSP*, pages 6645–6649. IEEE, 2013.
- [Hai02] Zhang Xiuling Li Haibin. A digital recognition method based on rbf neural network [j]. *Chinese Journal of Scientific Instrument*, 3:011, 2002.

- [Hal99] Mark A Hall. *Correlation-based feature selection for machine learning*. PhD thesis, The University of Waikato, 1999.
- [Han99] David J. Hand. Statistics and Data Mining: Intersecting Disciplines. *SIGKDD Explorations*, 1(1):16–19, 1999.
- [Hau85] John Haugeland. *Artificial Intelligence The Very Idea*. MIT Press, 1985.
- [Hay99] S. Haykin. *Neural Networks: A Comprehensive Foundation*. Prentice Hall, 1999.
- [HCC⁺14] OA Hurricane, DA Callahan, DT Casey, PM Celliers, C Cerjan, EL Dewald, TR Dittrich, T Döppner, DE Hinkel, LF Berzak Hopkins, et al. Fuel gain exceeding unity in an inertially confined fusion implosion. *Nature*, 506(7488):343–348, 2014.
- [HE07] James Hays and Alexei A Efros. Scene Completion Using Millions of Photographs. *ACM Transactions on Graphics (SIGGRAPH 2007)*, 26(3), 2007.
- [Heb49] Donald O. Hebb. *The Organization of Behavior*. John Wiley, New York, 1949.
- [HHM⁺07] Jerry Wayne Hughes, Amanda E Hubbard, Dmitri A Mossessian, Brian LaBombard, Theodore Matthias Biewer, Robert Seth Granetz, Martin Greenwald, Ian H Hutchinson, James H Irby, Yijun Lin, et al. H-mode pedestal and lh transition studies on alcator c-mod. *Fusion science and technology*, 51(3):317–341, 2007.
- [HO00] A Hyvärinen and E Oja. Independent component analysis: algorithms and applications. *Neural Networks: The Official Journal of the International Neural Network Society*, 13(4-5):411–430, Jun 2000. PMID: 10946390.

- [Hop82] J. J. Hopfield. Neural Networks and Physical Systems with Emergent Collective Computational Abilities. *Proceedings of the National Academy of Sciences*, 79:2554–2558, 1982.
- [Hop04] J. J. Hopfield. Encoding for computation: Recognizing brief dynamical patterns by exploiting effects of weak rhythms on action-potential timing. *Proceedings of the National Academy of Sciences*, 101(16):6255–6260, 2004.
- [HOT06] G. E. Hinton, S. Osindero, and Y. W. Teh. A Fast Learning Algorithm for Deep Belief Nets. *Neural Computation*, 18:1527–1554, 2006.
- [HSW89] K. Hornik, M. Stinchcombe, and H. White. Multi-layer feedforward networks are universal approximators. *Neural Networks*, 2:359–366, 1989.
- [Hu95] Xiaohua Hu. *Knowledge discovery in databases: an attribute-oriented rough set approach*. PhD thesis, Citeseer, 1995.
- [ID05] Timo MJ Ikonen and Olgierd Dumbrajs. Search for deterministic chaos in elm time series of asdex upgrade tokamak. *Plasma Science, IEEE Transactions on*, 33(3):1115–1122, 2005.
- [JGS⁺09] Rafael Jiménez, Elena Gervilla, Albert Sesé, Juan José Montañó, Berta Cajal, and Alfonso Palmer. Dimensionality reduction in data mining using artificial neural networks. *Methodology*, 5(1):26–34, 2009.
- [JK92] Biing-Hwang Juang and Shigeru Katagiri. Discriminative learning for minimum error classification [pattern recognition]. *Signal Processing, IEEE Transactions on*, 40(12):3043–3054, 1992.

- [JK01] Geurt Jongbloed and Ger Koole. Managing uncertainty in call centres using poisson mixtures. *Applied Stochastic Models in Business and Industry*, 17(4):307–318, 2001.
- [JP04] Ruth A Judson and Richard D Porter. Currency demand by federal reserve cash office: what do we know? *Journal of Economics and Business*, 56(4):273–285, 2004.
- [JZS⁺14] Weikuan Jia, Dean Zhao, Tian Shen, Chunyang Su, Chanli Hu, and Yuyan Zhao. A new optimized ga-rbf neural network algorithm. *Computational intelligence and neuroscience*, 2014:44, 2014.
- [K⁺95] Ron Kohavi et al. A study of cross-validation and bootstrap for accuracy estimation and model selection. In *Ijcai*, volume 14, pages 1137–1145, 1995.
- [KB98] K Kiviluoto and P Bergius. Exploring corporate bankruptcy with two-level self-organizing map. In *Decision Technologies for Computational Finance*, pages 373–380. Springer, 1998.
- [KBC88] Teuvo Kohonen, György Barna, and Ronald Chrisley. Statistical pattern recognition with neural networks: Benchmarking studies. In *Neural Networks, 1988., IEEE International Conference on*, pages 61–68. IEEE, 1988.
- [KKLT92] Teuvo Kohonen, Jari Kangas, Jorma Laaksonen, and Kari Torkkola. Lvq pak: A program package for the correct application of learning vector quantization algorithms. In *Proc. IJCNN*, volume 92, pages 725–730, 1992.
- [KNF91] V. Kadirkamanathan, M. Niranjan, and F. Fallside. Sequential adaptation of radial basis function neural networks and its application to time-series prediction. *Advances in Neural Information Processing Systems*, 3:721–727, 1991.

- [Koh82] Teuvo Kohonen. Self-Organized Formation of Topologically Correct Feature Maps. *Biological Cybernetics*, 43:59–69, 1982.
- [Koh86] T. Kohonen. Learning vector quantization for pattern recognition. Technical Report TKK-F-A601, Helsinki University of Technology, Laboratory of Computer and Information Science, 1986.
- [Koh90] Teuvo Kohonen. Improved versions of learning vector quantization. In *Neural Networks, 1990., 1990 IJCNN International Joint Conference on*, pages 545–550. IEEE, 1990.
- [Koh95] Ron Kohavi. A Study of Cross-Validation and Bootstrap for Accuracy Estimation and Model Selection. In *Proceedings of the International Joint Conference on Artificial Intelligence (IJCAI), 1995*, pages 1137–1145, 1995.
- [Koh01] Teuvo Kohonen. *Self-Organizing Maps*. Springer Berlin Heidelberg, Berlin, Heidelberg, 2001.
- [Koh13] Teuvo Kohonen. Essentials of the self-organizing map. *Neural Networks*, 37:52–65, 2013.
- [KOW⁺97] Juha Karhunen, Erkki Oja, Liuyue Wang, Ricardo Vigário, and Jyrki Joutsensalo. A class of neural networks for independent component analysis. *IEEE Transactions on Neural Networks*, 8(3):486–504, 1997.
- [KSM07] Efstathios Kirkos, Charalambos Spathis, and Yannis Manolopoulos. Data mining techniques for the detection of fraudulent financial statements. *Expert Systems with Applications*, 32(4):995–1003, 2007.
- [Kur92] Raymond Kurzweil. *Age of intelligent machines*. MIT Press, 1992.

- [KV] James Max Kanter and Kalyan Veeramachaneni. Deep feature synthesis: Towards automating data science endeavors.
- [KX96] Adam Krzyzak and Ley Xu. Optimal Radial Basis Function nets with Applications to Nonlinear Function Learning and Classification. 1996.
- [LBFL93] Robert K. Lindsay, Bruce G. Buchanan, Edward A. Feigenbaum, and Joshua Lederberg. DENDRAL: A Case Study of the First Expert System for Scientific Hypothesis Formation. *Artif. Intell.*, 61(2):209–261, 1993.
- [LC09] Ying Li and Bo Cheng. An improved k-nearest neighbor algorithm and its application to high resolution remote sensing image classification. In *Geoinformatics, 2009 17th International Conference on*, pages 1–4. Ieee, 2009.
- [LEK97] Cheng-Lin Liu, In-Jung Eim, and Jin H Kim. High accuracy handwritten chinese character recognition by improved feature matching method. In *Document Analysis and Recognition, 1997., Proceedings of the Fourth International Conference on*, volume 2, pages 1033–1037. IEEE, 1997.
- [LGT98] Steve Lawrence, C Lee Giles, and Ah Chung Tsoi. What size neural network gives optimal generalization? convergence properties of backpropagation. 1998.
- [LMD⁺11] Quoc V. Le, Rajat Monga, Matthieu Devin, Greg Corrado, Kai Chen, Marc’Aurelio Ranzato, Jeffrey Dean, and Andrew Y. Ng. Building high-level features using large scale unsupervised learning. *CoRR*, abs/1112.6209, 2011.
- [LN01] Cheng-Lin Liu and Masaki Nakagawa. Evaluation of prototype learning algorithms for nearest-neighbor classifier in application to handwritten character recognition. *Pattern Recognition*, 34(3):601–615, 2001.

- [LNR87] J. E. Laird, A. Newell, and P. S. Rosenbloom. Soar: An architecture for general intelligence. *Artificial Intelligence*, 33(1):1–64, 1987.
- [LPX02] Rong-Zhou Li, Su-Lin Pang, and Jian-Min Xu. Neural network credit-risk evaluation model based on back-propagation algorithm. In *Machine Learning and Cybernetics, 2002. Proceedings. 2002 International Conference on*, volume 4, pages 1702–1706. IEEE, 2002.
- [LS94] Seong-Whan Lee and Hee-Heon Song. Optimal design of reference models for large-set handwritten character recognition. *Pattern recognition*, 27(9):1267–1274, 1994.
- [LSL96] Hongjun Lu, Rudy Setiono, and Huan Liu. Effective data mining using neural networks. *Knowledge and Data Engineering, IEEE Transactions on*, 8(6):957–961, 1996.
- [LZ11] Yuxuan Li and Xiuzhen Zhang. Improving k nearest neighbor with exemplar generalization for imbalanced classification. *Advances in knowledge discovery and data mining*, pages 321–332, 2011.
- [Man13] Priyanka Manchanda. Quantum Artificial Intelligence : A Survey of Application of Quantum Physics in Artificial Intelligence. *CoRR*, abs/1309.7173, September 2013.
- [MCB⁺11] James Manyika, Michael Chui, Brad Brown, Jacques Bughin, Richard Dobbs, Charles Roxburgh, and Angela Hung Byers. Big data: The next frontier for innovation, competition, and productivity. Technical report, June 2011.
- [McD82] John McDermott. R1: A rule-based configurator of computer systems. *Artificial Intelligence*, 19(1):39–88, 1982.

-
- [McN01] James McNames. A fast nearest-neighbor algorithm based on a principal axis search tree. *Pattern Analysis and Machine Intelligence, IEEE Transactions on*, 23(9):964–976, 2001.
- [MCR14] Sérgio Moro, Paulo Cortez, and Paulo Rita. A data-driven approach to predict the success of bank telemarketing. *Decision Support Systems*, 62:22–31, 2014.
- [MD89] John Moody and Christian J Darken. Fast learning in networks of locally-tuned processing units. *Neural computation*, 1(2):281–294, 1989.
- [MGS15] Egidijus Merkevičius, Gintautas Garšva, and Rimvydas Simutis. Forecasting of credit classes with the self-organizing maps. *Information technology and control*, 33(4), 2015.
- [Mit06] Tom Mitchell. The discipline of machine learning. Technical Report CMU ML-06 108, 2006.
- [MMDF08] Robert J May, Holger R Maier, Graeme C Dandy, and TMK Gayani Fernando. Non-linear variable selection for artificial neural networks using partial mutual information. *Environmental Modelling & Software*, 23(10):1312–1326, 2008.
- [MMP02] Pabitra Mitra, C. A. Murthy, and Sankar K. Pal. Density-Based Multiscale Data Condensation. *IEEE Trans. Pattern Anal. Mach. Intell.*, 24(6):734–747, 2002.
- [MMRS55] J. McCarthy, M. L. Minsky, N. Rochester, and C.E. Shannon. A PROPOSAL FOR THE DARTMOUTH SUMMER RESEARCH PROJECT ON ARTIFICIAL INTELLIGENCE. <http://www-formal.stanford.edu/jmc/history/dartmouth/dartmouth.html>, 1955.

- [MMWB06] James Malone, Kenneth McGarry, Stefan Wermter, and Chris Bowerman. Data mining using rule extraction from kohonen self-organising maps. *Neural Computing & Applications*, 15(1):9–17, 2006.
- [MP43] W. S. McCulloch and W. Pitts. A logical calculus of the ideas immanent in nervous activity. *Bulletin of Mathematical Biophysics*, 5:115–133, 1943.
- [MP69] M. Minsky and S. Papert. *Perceptrons*. MIT Press, Cambridge, MA, 1969.
- [MPM02] Sushmita Mitra, Sankar K. Pal, and Pabitra Mitra. Data mining in soft computing framework: a survey. *IEEE Transactions on Neural Networks*, 13(1):3–14, 2002.
- [MSS04] Marvin Minsky, Push Singh, and Aaron Sloman. The St. Thomas Common Sense Symposium: Designing Architectures for Human-Level Intelligence. *AI Magazine*, 25(2):113–124, 2004.
- [MU94] John Moody and Joachim Utans. Architecture selection strategies for neural networks: Application to corporate bond rating prediction. In *Neural networks in the capital markets*, pages 277–300. Citeseer, 1994.
- [MVV97] Yves Moreau, Herman Verrelst, and Joos Vandewalle. Detection of mobile phone fraud using supervised neural networks: A first prototype. In *Artificial Neural Networks—ICANN’97*, pages 1065–1070. Springer, 1997.
- [NB13] Kalamkas Nurlybayeva and Gulnar Balakayeva. Algorithmic scoring models. *Applied Mathematical Sciences*, 7(12):571–586, 2013.
- [Nil98] Nils J. Nilsson. *Artificial Intelligence: A New Synthesis*. Morgan Kaufmann Publishers, San Mateo, CA, 1998.

- [NON12] Laura L. Nathans, Frederick L. Oswald, and Kim Nimon. Interpreting Multiple Linear Regression: A Guidebook of Variable Importance. *Practical Assessment, Research & Evaluation*, 17(9), 2012.
- [NP11] Aliakbar Niknafs and Soodabeh Parsa. A neural network approach for updating ranked association rules, based on data envelopment analysis. *J. Artif. Intell.*, 4:279–287, 2011.
- [NW89] N. Nguyen and B. Widrow. The truck backer-upper: An example of self learning in neural networks. In *Proceedings of the International Joint Conference on Neural Networks*, pages 357–363. IEEE Press, 1989.
- [OGO01] Paul O’Dea, Josephine Griffith, and Colm O’Riordan. Combining feature selection and neural networks for solving classification problems. In *Proc. 12th Irish Conf. Artificial Intell. Cognitive Sci.*, pages 157–166. Citeseer, 2001.
- [Oja89] E. Oja. Neural Networks, Principal Components, and Subspaces. *International Journal of Neural Systems*, 1(1):61–68, 1989.
- [Oja91] E. Oja. Data Compression, Feature Extraction, and Auto-association in Feedforward Neural Networks. In T. Kohonen, K. Mäkisara, O. Simula, and J. Kangas, editors, *Artificial Neural Networks*, volume 1, pages 737–745. Elsevier Science Publishers B.V., North-Holland, 1991.
- [Oja92] Erkki Oja. Principal Components, Minor Components, and Linear Neural Networks. *Neural Networks*, 5:927–935, 1992.
- [OJK12] Olanrewaju A Oludolapo, Adisa A Jimoh, and Pule A Kholopane. Comparing performance of mlp and rbf neural network models for predicting south africa’s energy consumption. *Journal of Energy in Southern Africa*, 23(3):41, 2012.

- [O'L13] D.E. O'Leary. Artificial Intelligence and Big Data. *Intelligent Systems, IEEE*, 28(2):96–99, 2013.
- [Omo08] S. Omohundro. The Basic AI Drives. 2008.
- [PA⁺11] Vincenzo Pacelli, Michele Azzollini, et al. An artificial neural network approach for credit risk management. *Journal of Intelligent Learning Systems and Applications*, 3(02):103, 2011.
- [Pan07] Sulin Pang. Credit risk evaluation model based on self-organizing competitive network. In *Control and Automation, 2007. ICCA 2007. IEEE International Conference on*, pages 2725–2728. IEEE, 2007.
- [Pea88] J. Pearl. *Probabilistic Reasoning in Intelligent Systems*. Morgan Kaufmann, 1988.
- [PG90] T. Poggio and F. Girosi. A Theory of Networks for Approximation and Learning,. *Proceedings of IEEE*, 78:1481–1497, 1990.
- [Pla91a] John Platt. A Resource-Allocating Network for Function Interpolation. *Neural Computation*, 3(2):213–225, 1991.
- [Pla91b] John C. Platt. Learning by Combining Memorization and Gradient Descent. *Advances in Neural Information Processing Systems*, 3:715–720, 1991.
- [PMG98] David Poole, Alan K. Mackworth, and Randy Goebel. *Computational intelligence - a logical approach*. Oxford University Press, 1998.
- [PRFC07] Fredy Ocaris Pérez Ramírez and Horacio Fernández Castaño. Las redes neuronales y la evaluación del riesgo de crédito. *Revista Ingenierías Universidad de Medellín*, 6(10):77–91, 2007.

-
- [PS91] J. Park and I. W. Sandberg. Universal Approximation Using Radial-Basis-Function Networks. *Neural Computation*, 3:246–257, 1991.
- [PS93] J. Park and I. W. Sandberg. Approximation and Radial-Basis-Function Networks. *Neural Computation*, 5:305–316, 1993.
- [PS11] Tuomas A. Peltonen Peter Sarlin. Mapping the state of financial stability. Technical report, European Central Bank, 2011.
- [PSS00] Michael Pinedo, Sridhar Seshadri, and J George Shanthikumar. Call centers in financial services: strategies, technologies, and operations. In *Creating Value in Financial Services*, pages 357–388. Springer, 2000.
- [Pun11] Joseph King-Fung Pun. *Improving Credit Card Fraud Detection using a Meta-Learning Strategy*. PhD thesis, University of Toronto, 2011.
- [PVM⁺15] Augusto Pereira, Jesús Vega, Raúl Moreno, Sebastián Dormido-Canto, Giuseppe A Rattá, Fernando Pavón, and JET EFDA Contributors. Feature selection for disruption prediction from scratch in jet by using genetic algorithms and probabilistic predictors. *Fusion Engineering and Design*, 2015.
- [Qui87] J. Ross Quinlan. Simplifying decision trees. *International journal of man-machine studies*, 27(3):221–234, 1987.
- [RHW86] David E. Rumelhart, Geoff E. Hinton, and R. J. Wilson. Learning representations by back-propagating errors. *Nature*, 323:533–536, 1986.
- [RK91] Elaine Rich and Kevin Knight. *Artificial intelligence (2. ed.)*. McGraw-Hill, 1991.

- [RM86] D.E. Rumelhart and J.L. McClelland. *Parallel Distributed Processing*, volume I+II. MIT Press, 1986.
- [RN10] Stuart Russell and Peter Norvig. *Artificial Intelligence: A Modern Approach*. Prentice Hall, 3 edition, 2010.
- [Ros58] Frank Rosenblatt. The Perceptron: A Probabilistic Model for Information Storage and Organization in the Brain. *Psychological Review*, 65(6):386–408, 1958.
- [RRK13] H Rasouli, C Rasouli, and A Koohi. Identification and control of plasma vertical position using neural network in damavand tokamak. *Review of Scientific Instruments*, 84(2):023504, 2013.
- [SA11] Ravinder Singh and Rinkle Rani Aggarwal. Comparative evaluation of predictive modeling techniques on credit card data. *International Journal of Computer Theory and Engineering*, 3(5):1–6, 2011.
- [Sam03] Ulrich Samm. Controlled thermonuclear fusion at the beginning of a new era. *Contemporary Physics*, 44(3):203–217, 2003.
- [SB98] R. S. Sutton and A. G. Barto. *Reinforcement Learning: An Introduction*. MIT Press, Cambridge, MA, 1998.
- [SBB⁺92] Eduard Säckinger, Bernhard E. Boser, Jane Bromley, Yann LeCun, and Lawrence D. Jackel. Application of the ANNA neural network chip to high-speed character recognition. *IEEE Transactions on Neural Networks*, 3(3):498–505, 1992.
- [Sch14] Jürgen Schmidhuber. Deep Learning in Neural Networks: An Overview. Technical report, IDSIA-03-14, 2014.
- [SdCC⁺12] Antonio Manuel Rubio Serrano, João Paulo Carvalho Lustosa da Costa, Carlos Henrique Cardonha, Ararigleno Almeida Fernandes, and Rafael Timóteo de Sousa Júnior. Neural network

- predictor for fraud detection: A study case for the federal patrimony department. In *THE SEVENTH INTERNATIONAL CONFERENCE ON FORENSIC COMPUTER SCIENCE - ICoFCS 2012*, 2012.
- [SK89] Avijit Saha and James D. Keeler. Algorithms for Better Representation and Faster Learning in Radial Basis Function Networks. In David S. Touretzky, editor, *NIPS*, pages 482–489. Morgan Kaufmann, 1989.
- [Ska97] David B Skalak. *Prototype selection for composite nearest neighbor classifiers*. PhD thesis, University of Massachusetts Amherst, 1997.
- [SM07] Bruno Silva and Nuno Marques. A Hybrid Parallel SOM Algorithm for Large Maps in Data-Mining. 2007.
- [SM13] E. Schikuta and E. Mann. N2Sky — Neural networks as services in the clouds. In *Neural Networks (IJCNN), The 2013 International Joint Conference on*, pages 1–8, Aug 2013.
- [SMO11] Ali Sadr, Najmeh Mohsenifar, and Raziye Sadat Okhovat. Comparison of mlp and rbf neural networks for prediction of ecg signals. *Int. J. Comput. Sci. Netw. Secur*, 11(11):124–128, 2011.
- [SNS⁺06] Maria Teresinha Arns Steiner, Pedro José Steiner Neto, Nei Yoshihiro Soma, Tamio Shimizu, and Júlio Cesar Nievola. Using neural network rule extraction for credit-risk evaluation. *International Journal of Computer Science and Network Security*, 6:6–16, 2006.
- [SP11] Peter Sarlin and Tuomas A Peltonen. Mapping the state of financial stability. 2011.

- [SP13] Peter Sarlin and Tuomas A Peltonen. Mapping the state of financial stability. *Journal of International Financial Markets, Institutions and Money*, 26:46–76, 2013.
- [SS92] Robert M Sanner and Jean-Jacques E Slotine. Gaussian networks for direct adaptive control. *Neural Networks, IEEE Transactions on*, 3(6):837–863, 1992.
- [Sze84] David Y Sze. Or practice—a queueing model for telephone operator staffing. *Operations Research*, 32(2):229–249, 1984.
- [SZMY08] Siti Mariyam Shamsuddin, Anazida Zainal, and Norfadzila Mohd Yusof. Multilevel kohonen network learning for clustering problems. *Journal of ICT*, 7:1–25, 2008.
- [TCD] Luigi Troiano, Gerardo Canfora, and Vincenzo D’Alessandro. Soft computing in the basel ii framework.
- [Tec13] Demystifying Big Data, A practical guide to transforming the business of government. Technical report, TechAmerica Foundation, 601 Pennsylvania Avenue, N.W. North Building, Suite 600, 2013.
- [Tur50] Alan M. Turing. Computing Machinery and Intelligence. *Mind*, 59:433–460, 1950.
- [TWK03] Nicolas Tsapatsoulis, Manolis Wallace, and Stathis Kasderidis. Improving the performance of resource allocation networks through hierarchical clustering of high-dimensional data. In *Artificial Neural Networks and Neural Information Processing—ICANN/ICONIP 2003*, pages 867–874. Springer, 2003.
- [VCG09] Carlo Vallini, Francesco Ciampi, and N Gordini. Using artificial neural networks analysis for small enterprise default prediction modeling: Statistical evidence from italian firms.

- In *2009 Oxford Business & Economics Conference Proceedings, Association for Business and Economics Research (ABER)*, pages 1–26, 2009.
- [Vel00] Alfredo Vellido. *A methodology for the characterization of business-to-consumer E-commerce*. PhD thesis, Liverpool John Moores University, 2000.
- [VGS05] Vladimir Vovk, Alex Gammerman, and Glenn Shafer. *Algorithmic learning in a random world*. Springer Science & Business Media, 2005.
- [VL15] Oriol Vinyals and Quoc Le. A neural conversational model. *arXiv preprint arXiv:1506.05869*, 2015.
- [VNB⁺13] Dan Vesset, Ashish Nadkarni, Rob Brothers, Christian A. Christiansen, Steve Conway, et al. Worldwide Big Data Technology and Services 2013–2017 Forecast. Technical report, December 2013.
- [VSN03] Vladimir Vovk, Glenn Shafer, and Ilia Nourtdinov. Self-calibrating probability forecasting. In *Advances in Neural Information Processing Systems*, page None, 2003.
- [WD91] Dietrich Wettschereck and Thomas Dietterich. Improving the performance of radial basis function networks by learning center locations. In *NIPS*, volume 4, pages 1133–1140. Citeseer, 1991.
- [Wer74] P. Werbos. *Beyond Regression: New Tools for Prediction and Analysis in the Behavioral Sciences*. PhD thesis, Harvard University, 1974.
- [WF97] Weijian Wan and Donald Fraser. An msom framework for multi-source fusion and spatio-temporal classification. In

- Geoscience and Remote Sensing, 1997. IGARSS'97. Remote Sensing-A Scientific Vision for Sustainable Development., 1997 IEEE International*, volume 4, pages 1657–1659. IEEE, 1997.
- [WH60] B. Widrow and M. E. Hoff. Adaptive switching circuits. *Institute of Radio Engineers, Western Electronics Show and Convention*, Part 4:96–104, 1960.
- [Wid62] B. Widrow. Generalization and information storage in networks of adaline neurons. *Self-Organizing Systems*, pages 435–461, 1962.
- [Wil88] R. J. Williams. Towards a theory of reinforcement-learning connectionist systems. Technical report, Northeastern University, 1988.
- [Win92] Patrick Henry Winston. *Artificial intelligence (3. ed.)*. Addison-Wesley, 1992.
- [WL14] Jiunn-Lin Wu and I-Jing Li. The Improved SOM-Based Dimensionality Reducton Method for KNN Classifier Using Weighted Euclidean Metric. *International Journal of Computer, Consumer and Control (IJ3C)*, 3(1), 2014.
- [WW84] Bernard Widrow and E. Walach. On the statistical efficiency of the LMS algorithm with nonstationary inputs. *IEEE Transactions on Information Theory*, 30(2):211–221, 1984.
- [WZ89] R. J. Williams and D. Zipser. A Learning Algorithm for Continually Running Fully Recurrent Neural Networks. *Neural Computation*, 1(2):270–280, 1989.
- [WZF⁺03] Ke Wang, Senqiang Zhou, Ada Wai-Chee Fu, Jeffrey Xu Yu, Fu Jeffrey, and Xu Yu. Mining changes of classification by correspondence tracing. In *SDM*, pages 95–106. SIAM, 2003.

- [Yan12] Weizhong Yan. Toward automatic time-series forecasting using neural networks. *IEEE transactions on neural networks and learning systems*, 23(7):1028–1039, Jul 2012.
- [Yar95] David Yarowsky. Unsupervised Word Sense Disambiguation Rivaling Supervised Methods. In *In Proceedings of the 33rd Annual Meeting of the Association for Computational Linguistics*, pages 189–196, 1995.
- [yC09] S. Ramón y Cajal. *Histología del sistema nervioso del hombre y de los vertebrados*. A. Maloine, 1909.
- [YF98] Jen-Lun Yuan and T.L. Fine. Neural-network design for small training sets of high dimension. *Neural Networks, IEEE Transactions on*, 9(2):266–280, Mar 1998.
- [Yin02] Hujun Yin. Visom-a novel method for multivariate data projection and structure visualization. *Neural Networks, IEEE Transactions on*, 13(1):237–243, 2002.
- [Yiu12] Chris Yiu. The Big Data Opportunity. Technical report, Police Exchange, 2012.
- [Yud08] E. Yudkowsky. Artificial Intelligence as a Positive and Negative Factor in Global Risk. . *Oxford University Press*, 2008.
- [YXL15] Hui Yuan, Wei Xu, and Raymond YK Lau. Topic sentiment mining for product sales performance prediction. 2015.
- [Zad65] Lotfi Zadeh. Fuzzy logic and its applications. *New York, NY, USA*, 1965.
- [ZBRJ05] Xiong Zhi-Bin and Li Rong-Jun. Credit risk evaluation with fuzzy neural networks on listed corporations of china. In *VLSI Design and Video Technology, 2005. Proceedings of 2005 IEEE International Workshop on*, pages 397–402. IEEE, 2005.

- [ZdP⁺12] P. Zikopoulos, D. deRoos, K. Parasuraman, T. Deutsch, J. Giles, and D. Corrigan. *Harness the Power of Big Data – The IBM Big Data Platform*. McGraw-Hill, 2012.
- [ZKA12] Masoumeh Zareapoor, Seeja KR, and M Afshar Alam. Analysis on credit card fraud detection techniques: Based on certain design criteria. *International Journal of Computer Applications*, 52(3), 2012.