

Inteligencia Artificial: Hacia la singularidad



iStock/Getty Images.

Que el mundo de la Inteligencia Artificial avanza a una velocidad frenética es algo innegable. Se están logrando año tras año avances que hasta hace muy poco creíamos imposibles. Estos hitos imposibles quedaban en manos del cine y de la literatura de ciencia ficción. ¿Es posible que pronto veamos en el mundo real cosas únicamente vistas en películas y libros? Este artículo ofrece un repaso al panorama actual de la Inteligencia Artificial.

José Manuel Rodríguez Madrid | Consultor del área de Soluciones Digitales de Afi

Durante la última década hemos asistido al [«mejor verano de la Inteligencia Artificial»](#) con un gran desarrollo de **algoritmos de Machine Learning en el ámbito de las redes neuronales**, conocidos con el nombre de Deep Learning. A diferencia del resto de algoritmos, éstos permiten inferir patrones de los datos con una mínima intervención humana a la hora de realizar el proceso de extracción de características (*features extraction*) de los mismos. Estas redes neuronales con innumerables capas aprenden aspectos de los datos, y a medida que éstos pasan a través de la red, ésta puede «aprender» conceptos más abstractos. Para todo ello necesitan gran cantidad de datos y capacidades de computación masivas. Dado que estamos

en la era del **Big Data**, asistimos a un momento ideal para el crecimiento de este tipo de algoritmos.

Presenciamos la creación de **sistemas de procesamiento de lenguaje natural** (PLN) como son [GPT-3](#) (inteligencia artificial creada por OpenAI) o [Mega-tron-Turin](#) (inteligencia artificial creada por Microsoft y NVIDIA), con redes neuronales de miles de millones de parámetros que son capaces de generar textos creíbles y realistas, de responder preguntas sobre cualquier tema y de una larga lista de capacidades. Empresas como Deepmind (del grupo Alphabet) crean inteligencias artificiales como [AlphaZero](#), capaz de aprender en unas pocas horas a jugar como el mejor de los maestros de juegos milenarios (p.e. Go) o

como [AlphaFold2](#), que resuelve un problema de alta complejidad como es el plegamiento de proteínas. La empresa [Boston Dynamics](#) crea robots autónomos de aspecto humanoide que nos recuerdan a ciertas películas de ciencia ficción. Y sistemas como [Google Duplex](#), capaces de mantener una conversación con una persona sin que ésta pueda apreciar que está hablando con una máquina, desafían el **test de Turing**. Técnicas como el **deep fake**, la cual puede cambiar el rostro de una persona por el de otra y es capaz de lo mejor en el [mundo cinematográfico](#) y de lo peor si se aplica con un fin no tan noble como la falsificación [del video de un político](#).

Ante este abanico de inteligencias artificiales, ¿podríamos encontrarnos ante un escenario apocalíptico en el que se produzca la llamada **singularidad** (momento en el que aparezca una inteligencia artificial capaz de automejorarse y a su vez crear otras inteligencias artificiales)? Aunque algunos expertos la sitúan cercana en el tiempo (Raymond Kurzweil en su libro «La singularidad está cerca» señala que «ocurrirá alrededor del año 2045», y Vernor Vinge la predice incluso 15 años antes, en el 2030), actualmente este escenario se encuentra más cercano a la ciencia ficción que a la realidad. Sin embargo, sí que debemos tener muy presentes ciertas cuestiones en el futuro cercano.

La primera de ellas es la **desinformación**. En los últimos tiempos es cada vez más frecuente la proliferación de artículos poco rigurosos y con poca base científica, en muchos casos tratando de agitar conspiraciones y persiguiendo el sensacionalismo, así como el fenómeno conocido como *clickbait*. Hace varios años se viralizó una noticia en la que se informaba de que [Facebook tuvo que desconectar una inteligencia artificial de la que perdió el control](#). Nada más lejos de la realidad. Esta inteligencia artificial se desconectó porque no servía para el propósito que había sido creada.

La segunda es el **sesgo** presente en los datos y que puede trasladarse al aprendizaje de las inteligencias artificiales. ¿Queremos trasladar estos errores y fallos a nuestros algoritmos? Rotundamente no. Recientemente apareció en medios de comunicación la noticia sobre una [inteligencia artificial que habiendo](#)

[sido entrenada para emitir juicios morales, éstos resultaban misóginos y racistas](#). El problema radicaba en la base de datos que se había utilizado para el entrenamiento, con juicios éticos de ciudadanos estadounidenses, y en algunos casos, opiniones provenientes de Reddit. Además, en este caso particular se nos plantea el [dilema moral y ético de usar inteligencias artificiales para estos fines](#).

La tercera es la **responsabilidad** que se les está otorgando a los algoritmos. Las inteligencias artificiales son, en muchos casos, sistemas muy complejos, con algoritmos que producen salidas no deterministas ya que están basados en probabilidades. ¿Queremos dejar en manos de inteligencias artificiales cualquier situación, o por el contrario se debería regular su uso? Un ejemplo: [en el año 2010 se produjo el llamado Flash Crash](#), donde algoritmos de *High Frequency Trading* provocaron una caída del 9% en el Dow Jones, la cual se recuperó en escasos minutos.

Y la última cuestión, nos recuerda la **incertidumbre que generan estos cambios tan disruptivos**. Estamos viendo cómo desde ciertos sectores, los puestos de trabajo se ven amenazados por la proliferación de inteligencias artificiales y sistemas basados en algoritmos y datos. Por ejemplo, [Amazon Go, una tienda automatizada donde no existen los cajeros](#), o los sistemas de conducción autónoma como el [Tesla autopilot](#), que puede revolucionar en un futuro próximo los transportes por carretera de personas o mercancías. Debemos ver estas inteligencias artificiales como una oportunidad de transformación de ciertos sectores, en los que, mediante el uso de éstas, se puedan generar nuevos puestos de trabajo cualificados. Paradójicamente, uno de estos sectores es el tecnológico, donde pronto harán falta [menos programadores y más «entrenadores» de inteligencias artificiales](#).

¿Hasta qué punto son peligrosas las inteligencias artificiales? Hoy el peligro realmente reside en el mal uso que el ser humano pueda hacer de ellas. Pronto podríamos ver cómo aparecen regulaciones que marquen el uso de los algoritmos. Uno de los mayores defensores de esta regulación es [Elon Musk, el cual ha abogado en numerosas ocasiones por la regulación](#) ::