

Elena Alfaro*

DATOS, INTELIGENCIA E INNOVACIÓN

Mucho se ha dicho ya sobre los datos como uno de los ingredientes fundamentales para el futuro digital de nuestra economía y en general de la sociedad. Los datos y la inteligencia basada en ellos no solo nos permiten hacer más eficientes los procesos existentes o tomar mejores decisiones. También permiten crear nuevos servicios, productos y soluciones que mejoran la vida de las personas haciéndola más cómoda, conveniente, segura e incluso más interesante y divertida. Pero también se abren incógnitas sobre si ciertos usos pueden hacer vulnerables a aquellos sobre los que la inteligencia se aplica o a los que no pueden acceder a ella.

Palabras clave: big data, inteligencia artificial, ciencia de datos, algoritmos, machine learning.

Clasificación JEL: C81, O31.

1. Introducción

Big data, analítica avanzada, aprendizaje automático, ciencia de datos, inteligencia artificial (IA), *data-driven*... Todos estos son conceptos que hoy por hoy encontramos casi en todos los lugares, desde las universidades y los centros de investigación más avanzados hasta los medios más generalistas, pasando por las reuniones de la alta dirección de las grandes empresas. ¿Por qué todo el mundo habla de *big data*? ¿por qué ahora? ¿realmente se entienden las implicaciones que la revolución de los datos tiene sobre nuestra evolución como sociedad, o incluso como especie? ¿entendemos lo que significan todas estas palabras? En las próximas paginas intentaremos explicar de una forma pragmática y no técnica el porqué de todo esto, y sobre todo por qué es importante entenderlo.

2. Porqué ahora

Hay tres razones principales por las que esta revolución está sucediendo en este momento (cuando decimos este momento nos referimos sobre todo a los últimos cinco años), y son:

— La digitalización de la información. Hoy por hoy todo lo que hacemos deja un rastro digital, es decir, que en el mundo digital existe una representación cada vez más perfecta y fiel de lo que ocurre en el físico. Cuando usamos el móvil para hablar, accedemos a una aplicación, hacemos un pago con tarjeta, leemos la prensa *online*, pasamos junto a una cámara de seguridad o sacamos un billete del transporte público, dejamos datos que dicen qué hemos hecho, dónde y cuándo, y en la mayoría de los casos esto está vinculado directamente a quienes somos. Y no solo eso, sino que cada vez más, suceden cosas en el entorno digital que no tienen por qué tener su origen en el

* Head of Data & Open Innovation, BBVA.

físico, por ejemplo, cuando jugamos *online* o valoramos un hotel en una web de reseñas. Aunque en muchos casos estas interacciones digitales sí que tienen un impacto muy claro en el mundo «real». Todo esto ha hecho que el volumen de datos digitales generados cada día, así como su diversidad de orígenes haya explotado, y lo siga haciendo.

— La ley de Moore y sus implicaciones en la disponibilidad cada vez mayor y con más potencia de capacidad de proceso y almacenamiento de esos datos que se están generando de forma masiva. Quizás una cosa sea consecuencia de la otra, pero la realidad es que cada vez es más barato y más sencillo acceder a sistemas que hasta hace muy poco solo estaban al alcance de las grandes empresas con fuerte capacidad inversora. Hoy por hoy cualquier emprendedor o *startup* puede acceder a infraestructuras potentes y flexibles en la nube, acceder a las últimas tecnologías y herramientas, y lo que es más importante, pagar solo por lo que se usa.

— Las expectativas y demandas crecientes de los clientes y usuarios, que ya no están dispuestos a esperar más de unos segundos para tener acceso a lo que quieren, que buscan experiencias instantáneas, sorprendentes y personalizadas, y que dejan de ser fieles a aquellas compañías y marcas que no consiguen ofrecer el nivel de servicio que sí están recibiendo de empresas digitales en principio mucho menos conocidas para ellos, pero en las que confían en muy poco tiempo. Estamos viendo este proceso por ejemplo en la disrupción que está teniendo lugar en el sector financiero, donde empresas con más de 100 años pierden partes clave de su negocio en favor de pequeñas *fintech* que optimizan una parte de la oferta para hacerla a veces más barata y, en todos los casos, más conveniente y sencilla.

Estas tres tendencias, que evidentemente no son nuevas en sí mismas, tienen unas implicaciones muy importantes en los negocios que se «digitalizan», es decir, que pasan de ser puramente físicos a ser también digitales, porque una vez un negocio o actividad

se convierte en datos, los costes de transacción bajan y también lo hacen las barreras de entrada al mismo. Además, si tampoco existen barreras económicas en cuanto al acceso a los sistemas y plataformas necesarios para poder crear y ofrecer los servicios a bajo coste y de forma personalizada, el resultado es que muchos negocios «tradicionales» caen en la disrupción, y quedan transformados para siempre. Hemos visto casos claros de este proceso en la industria musical, en la de los medios, en la de los viajes... pero también en negocios que en principio estaban basados en activos puramente físicos como son el transporte, los hoteles y la medicina. Realmente, al igual que nada escapa a su representación digital en datos, ningún negocio va a escapar a esta ola de digitalización y disrupción.

3. Las empresas *data-driven*

¿Y quién queda mejor posicionado dada esta situación? Pues aquellas empresas que saben utilizar los datos para crear esas experiencias increíbles, personales e inmediatas, y que además aprenden con cada interacción a ser cada vez mejores. Son empresas que tienen los datos en el centro de su estrategia, donde se busca la manera de obtener de forma permanente nuevos datos, ya que más datos son siempre más oportunidades, y no más problemas, como sucede todavía en contextos aún no digitales. Veamos algunos ejemplos que tienen algo en común: el uso de los datos para la creación de motores de decisión automáticos o algoritmos que, envueltos de una gran experiencia (que además mejora con el uso), han sido capaces de atraer a decenas o incluso centenares de millones de clientes en muy poco tiempo. Estos ejemplos, no por coincidencia, son los mismos que los que se suelen utilizar para hablar de las empresas digitales de éxito, aunque nos vamos a centrar en cómo ponen en valor los datos.

- Amazon: podríamos pensar que el secreto de Amazon está en su logística, en su amplio catálogo, en la posibilidad de usarlo como plataforma para cualquier comerciante, o en sus precios, y esto es cierto,

FIGURA 1

FILTRO COLABORATIVO DE AMAZON PARA PRODUCTOS A CLIENTES SIMILARES



pero lo que realmente hizo a esta tienda *online* empezar a crecer a nivel global y a través de todos los sectores fue la introducción de su algoritmo de recomendación: «gente que compró X también compró Y».

Este es un tipo de algoritmo conocido como filtro colaborativo, que se basa en la historia de compra, de navegación y de *ratings* de producto de millones de clientes para crear similitudes entre ellos, y esto les permite mostrar y ofrecer, con alta probabilidad de acierto, otros productos que nos pueden interesar. Amazon ha llevado a cabo experimentos para ver si su motor de recomendación es mejor o peor que el conocimiento de un grupo de expertos en un dominio específico como puede ser el de la literatura en torno a una temática, haciéndoles competir para ver quién consigue vender más a clientes interesados en ese dominio. El resultado siempre es el mismo: gana el algoritmo al humano experto de forma consistente y además por mucho. Con este algoritmo se personaliza y se facilita la experiencia de compra, y esto lleva a vender más. En el segundo trimestre de 2012, año en el que este tipo de recomendaciones se extendió a toda su web y no solo a los libros, Amazon reportó un incremento de ventas del 29 por 100, y lo imputó en su mayoría a esta mejora algorítmica (Figura 1).

FIGURA 2

LA EXPERIENCIA PERSONALIZADA DE NETFLIX



● Netflix: de nuevo, podríamos pensar que donde ha innovado Netflix ha sido sobre todo en su catálogo, producción propia y modelo de distribución de entretenimiento, y también sería verdad, pero el secreto de su éxito lo tenemos que buscar de nuevo en un algoritmo, o más bien un conjunto de algoritmos. Se trata de su sistema de recomendación, que les permite personalizar en extremo la experiencia que recibe el cliente cuando entra en la aplicación. Lo que hemos visto, lo que hemos calificado, el tiempo que hemos tardado en acabar una serie o si hemos dejado capítulos a medias, incluso si entramos más en contenido cuya carátula o foto tiene tonos más oscuros o más claros, o si aparece en primer plano el protagonista masculino o femenino... Todo esto, multiplicado por sus más de 100.000.000 de clientes de pago en todo el mundo, les permite ofrecer una experiencia única, que además sigue mejorando gracias al análisis de los datos de uso de su plataforma (Figura 2).

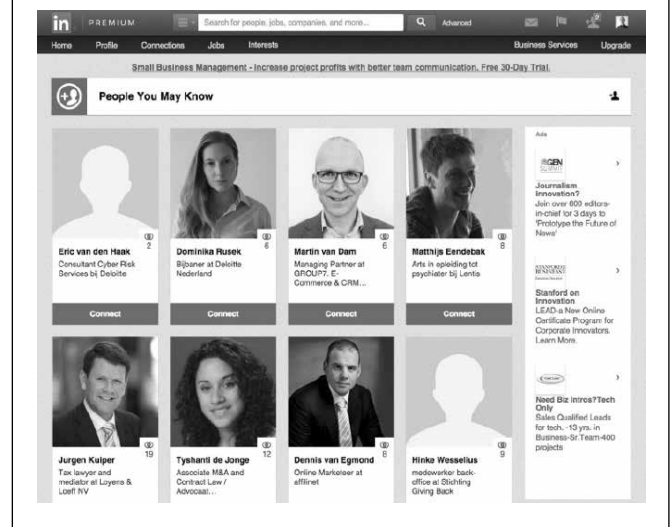
Hoy por hoy, aproximadamente el 85 por 100 de los contenidos que los clientes ven a través de Netflix no están basados en búsquedas directas de los clientes, sino en las recomendaciones que nos hacen y nos muestran a cada uno en la pantalla de inicio. Creemos

que vemos una serie porque la hemos buscado, pero la realidad es que Netflix lo hizo para nosotros, y solo tenemos que decir que sí. Y no solo se usan los datos para ofrecer contenido o personalizar la experiencia, sino que también se utilizan para generar el propio contenido de manera que «enganche» con lo que interesa a un tipo de usuario. Esto sucede ya en muchas de sus series y películas, siendo el caso más famoso el de *House of Cards*, con un guión totalmente *data-driven*. En el último dato al que hemos tenido acceso (cuarto trimestre de 2014), Netflix afirmó que empleaba a 300 personas e invertía 150.000.000 de dólares al año en mantener y mejorar su sistema de recomendación. En 2016 ingresó 9.000 millones de dólares en todo el mundo.

- LinkedIn: esta red social profesional tiene hoy más de 500.000.000 de clientes en 200 países, pero cuando nació en 2003 y durante los siguientes años las cosas no parecían fáciles. De hecho, tardaron casi cuatro años en llegar a 4.000.000 de clientes, y fue en 2006 cuando la cosa empezó a explotar. De nuevo, podemos echar parte de la culpa a un algoritmo. Se trata de la funcionalidad *People you may know* (gente que podrías conocer), la cual, basada en tecnologías de recomendación similares a las de los sistemas de Netflix y Amazon, junto con otras investigaciones en torno a sistemas complejos y teoría de grafos, hizo que el número de personas que invitaban a otras personas creciera de forma exponencial, lo cual no hizo más que alimentar al algoritmo para que descubriera a más y más personas que podríamos conocer. Tres años después de introducir esta funcionalidad, LinkedIn ya contaba con 54.000.000 de clientes. Es decir, 50.000.000 en el mismo tiempo en el que en los tres años anteriores había conseguido 4.000.000. Desde entonces, LinkedIn ha ido poniéndonos cada vez más fácil invitar a gente conocida, como la introducida recientemente, la cual cada vez que entramos nos selecciona a cinco personas que con mucha probabilidad querremos invitar a conectar y les manda un mensaje con solo un clic (Figura 3).

Hay muchos más ejemplos en prácticamente todos los sectores donde se cumple la misma idea: algoritmos

FIGURA 3
«GENTE QUE PODRÍAS CONOCER»
DE LINKEDIN



envueltos en experiencias increíbles que además se hacen mejores a medida que los usamos más. Citamos aquí algunos otros ejemplos que puede resultar interesante investigar también:

- Spotify: radar de novedades, *Discover Weekly*, *Daily Mix*, todos ellos servicios personalizados de descubrimiento de nueva música desde diferentes perspectivas.
- Uber Pool: servicio de compartición de vehículos basado en la predicción de clientes en una zona, dados los datos históricos de uso del servicio en la misma.
- Kayak y su motor predictivo de precios que recomienda si deberíamos comprar un billete ahora o esperar a que baje.
- Trackers o medidores de actividad física como Fitbit, Withings (ahora Nokia), Jawbone, etc., que nos hacen recomendaciones basadas en el comportamiento de otras personas.
- Google: prácticamente todos sus productos están basados en algoritmos y en los datos que recogen de todos ellos, lo cual les permite personalizar y mejorar

cada vez más la experiencia: *places, maps, gmail* (versión gratuita), *ads, translate...*

Ahora bien, ¿son estas empresas las únicas que pueden ser *data-driven*? ¿qué es ser *data-driven* realmente? En esencia, dos cosas: ser capaz de tomar tus decisiones basadas en datos (en lugar de hacerlo solo por intuición, o por jerarquía, por inercia...), y en el extremo ser capaz de automatizar estas decisiones, es decir, de eliminarlas implementando algoritmos inteligentes que no requieren de intervención humana, o que sea mínima. Y a esto puede y debe querer aspirar cualquier compañía, tradicional o no, que quiera competir en el mundo digital, ya que, si no, no podrá ofrecer servicios y experiencias únicas y personalizadas, y además no podrá competir en costes. ¿Pero qué necesita una empresa para dar este paso?

4. Los ingredientes: datos, talento y buenas preguntas

Evidentemente necesitamos datos, y muchos, pero no solo esto. También necesitamos que se puedan usar, que tengan calidad, entender bien lo que nos están diciendo, y si son o no útiles para el problema que queremos resolver, así como saber si es legal y éticamente aceptable usarlos en un contexto concreto. De esto último hablaremos un poco más adelante. Subestimamos muchas veces el trabajo que implica tener datos disponibles y de calidad, aunque se generen dentro de tu propia empresa o institución. Y esto sucede porque, hasta hace poco, nadie pensaba que los datos podrían tener más vidas que el motivo por el que se generan, normalmente para dar un servicio y como mucho para medirlo. En muchos casos, incluso cuando el servicio ya se ha prestado, y a no ser que haya un requerimiento legal, los datos se borran para que no ocupen espacio. Una empresa que no cuida sus datos nunca será *data-driven*, y por tanto se quedará muy lejos del valor que puede extraer de ellos.

Además de datos, necesitamos a personas que los trabajen, que sean capaces de extraer de ellos conocimiento nuevo que nos ayude a abordar los problemas

de manera más eficiente o acertada, o que nos ayude a descubrir oportunidades tanto en los servicios que ofrecemos, como en las decisiones que tomamos o cómo operamos un proceso. Un nuevo ámbito de conocimiento ha nacido derivado de la explosión en la generación de datos, de su diversidad, y de la capacidad que tenemos hoy de procesarlos y trabajar con ellos en toda su granularidad. Se trata de la Ciencia de Datos, que es una disciplina que se encuentra en la intersección entre los dominios más cuantitativos como son las matemáticas y la estadística, y la ingeniería informática principalmente. Hay otras áreas que también se suelen destacar en estos perfiles —llamados científicos de datos o *data scientists*—, como son el conocimiento del dominio de negocio, la comunicación, la visualización de datos..., pero sobre todo se trata de personas que son capaces de entender un problema, buscar su respuesta en conjuntos de datos sin importar su volumen o procedencia, y de programar un algoritmo que automatiza la resolución de ese problema para resolverlo de forma masiva y con capacidad de aprender de su propio resultado. Hasta hace poco no ha existido un camino claro hacia esta carrera, ya que sus conocimientos se enmarcaban en varios títulos universitarios, pero cada vez más están surgiendo grados, cursos y másteres enfocados directamente a esta profesión, que cuenta con una demanda creciente, para la que estamos lejos de contar con la oferta y la experiencia necesarias. Entrando un poco más en algunas de las disciplinas que se trabajan en la Ciencia de Datos, y sin ánimo de ser exhaustivos (sobre todo porque a nivel de herramientas y lenguajes son ámbitos tecnológicos que evolucionan rápido) tenemos:

— Matemáticas y estadística: modelización, optimización y predicción, inferencia bayesiana, aprendizaje automático o *machine learning*, diseño de experimentos...

— Informática y tecnologías de la información: fundamentos de computación, bases de datos y sistemas *big data* (MapReduce, Hadoop, Pig, Spark...), lenguajes de programación (R, Python, Scala...), *cloud computing*...

— Dominio o tipos de problemas a resolver. Por ejemplo, no es lo mismo el dominio y el dato financiero que el de la previsión meteorológica, y es importante entender bien lo que se quiere resolver, sobre todo para poder buscar una solución óptima y creativa. Pero en nuestra experiencia, esto es algo que los buenos científicos de datos no suelen tener problema en aprender y abordar en poco tiempo, al menos en lo que hemos tenido la oportunidad de observar.

— Comunicación y visualización: esto no es imprescindible en todo *data scientist*, pero sí en un equipo que desarrolla esta disciplina, ya que muchas veces el problema es la comprensión humana del *insight* o aprendizaje derivado de los datos, con el fin de llevarlo a la acción, y en muchos casos la visualización actúa como traductor de ese *insight* entre las máquinas y las personas.

El motivo por el que llamamos a esto Ciencia de Datos, es porque realmente es una ciencia, ya que se basa en el establecimiento de hipótesis para su validación o refutación. Es decir, que sigue el método científico de investigación, por el cual trata de responder preguntas concretas, de ahí la importancia de que esas preguntas sean las adecuadas. En general los proyectos tienen una fase de exploración, y otra muy importante de prueba de diferentes datos y modelos hasta llegar al que optimiza la solución a nuestro problema. Se trata por tanto de procesos iterativos, donde, volviendo al problema de la disponibilidad y la calidad de los datos, muchas veces se ocupa la mayoría del tiempo en limpiarlos y prepararlos antes de aplicar las metodologías analíticas.

Cuando estamos hablando de desarrollo de servicios y productos de datos, como los que hemos explicado en el apartado de las empresas *data-driven*, es muy importante que la ciencia de datos se desarrolle muy cerca del diseño de experiencia y servicio, y de nuevo hacerlo de forma iterativa y no secuencial. No se trata de diseñar la experiencia y luego contar simplemente con conectarle un algoritmo, sino que ambos deben desarrollarse juntos, ya que si no, corremos el riesgo, o bien

de creer que un modelo matemático va a funcionar para algo antes de probarlo, o lo contrario, no vamos a contar con todas las posibilidades del aprendizaje automático a la hora de ser creativos con nuestro producto. Por ejemplo, cuando hablamos de predicciones, antes de ver cómo una predicción, ya sea de un precio o de la llegada de un autobús, se muestra a los usuarios, hay que ver con cuánta precisión podemos hacerlo y si funciona en todos los casos ya que, si no, podemos entregar una experiencia decepcionante, cuando podría ser al contrario (un ejemplo simple es que no acertaremos las mismas veces si decimos que un billete de avión va a costar 543 euros que si decimos que va a costar entre 530 euros y 560 euros, y esto puede llevar a un diseño diferente de la interacción y la comunicación).

5. De *big data* a inteligencia artificial

Ya hemos visto algunos de los problemas que resuelve y automatiza la ciencia de datos en los casos de:

- Clasificar: por ejemplo, cuando Amazon junta productos similares.
- Comparar: cuando Facebook nos dice si en una foto hay una misma persona o Google nos dice que un *e-mail* es *spam*.
- Detectar: la detección de un fraude en una compra *online*, o de un *trending topic* en Twitter.
- Optimizar: determinar cuál es el mejor precio para un cliente, algo que hace Amazon en su tienda *online* sin que muchas veces lo sepamos.
- Predecir: Kayak y su predicción de precios de los billetes de avión.
- Recomendar: Netflix, Spotify, «a quién seguir» en Twitter.
- Hacer un *ranking* o *score*: cuando los buscadores nos dan los resultados por relevancia.

Para poder aplicar la ciencia de datos a este tipo de problemas, se pueden utilizar modelos más simples o más complejos, pero lo que vemos en la actualidad es que, cuando disponemos de muchos datos sobre el fenómeno que estamos estudiando, lo que

mejor funciona es que sea una máquina la que infiera la solución óptima a nuestro problema en base a esos datos. A esto se le llama aprendizaje automático, o *machine learning*, y básicamente da la vuelta a la manera en la que se han programado tradicionalmente los sistemas inteligentes o expertos. Hemos pasado de tener un experto que, mediante reglas, enseña o programa a una máquina para que resuelva un problema, a alimentar a una máquina con muchos datos del fenómeno a estudiar, de manera que es capaz de inferir los patrones que le llevan a una solución normalmente mejor que la basada en reglas, ya que puede tener en cuenta mucha más información. Por ejemplo, las primeras máquinas que sabían jugar al ajedrez contenían las reglas que un experto jugador es capaz de enumerar, y que después un programador puede codificar. Se trataría de un gran árbol con muchas ramas y a muchos niveles, del estilo «Si X, entonces Y» siendo X una situación de las fichas, e Y el siguiente mejor movimiento según la experiencia de ese experto. Sería un árbol muy grande, y probablemente no contemplaría todas las posibilidades, pero serviría para jugar una partida (y probablemente para ganar a alguien poco experimentado). Comparemos esto con tener registradas todas las partidas de los mejores jugadores del mundo de los últimos 20 años. ¿Cuánta experiencia es esta comparada con el conocimiento de un experto o de una docena de ellos? Probablemente varios órdenes de magnitud superior. Ahora cojamos todos esos datos de las partidas históricas, con todos sus movimientos y sus resultados finales, y démoselos a una máquina y a un *data scientist* que entiende lo que queremos conseguir y para qué, y que es capaz de probar diferentes modelos o algoritmos sobre esos datos hasta que el sistema converge, es decir, llega a la mejor solución. Para ello habrá probado miles, millones o trillones de posibilidades con no menos variables, hasta dar con el movimiento en cada caso que maximiza las probabilidades de ganar, dada cualquier situación en el tablero. ¿Cómo puede competir una persona contra eso? Pues aquí está la magia del

machine learning cuando cuenta con muchos datos, y de ahí que en los últimos años hayamos visto un salto cualitativo en los ámbitos de la traducción, la visión por computador, o la conducción autónoma, entre muchas otras cosas. Esto no significa que la máquina aprenda sin ayuda, pero sí que dado un objetivo, un modelo o conjunto de modelos (ahí entra la intervención humana) y muchos datos, es capaz de optimizar la solución a un problema, de forma que esta se pueda industrializar, es decir, que generalice en un número muy alto de los casos (esto dependerá del problema).

¿Qué tiene que ver esto con la inteligencia artificial? Pues va quedando claro que mucho. Pero primero pensemos un poco en este término: inteligencia artificial. ¿Qué significa? Pues lamentablemente algo no muy concreto. Primero, este término se acuñó por primera vez en los años cincuenta. El que haya llegado a la actualidad con el mismo halo de misterio, algo que no deja de ser curioso, tiene que ver con el hecho de que, semánticamente, siempre representa algo que está por venir, que es del futuro. Y claro, eso siempre genera incertidumbre. ¿Por qué sucede esto? Pues porque una vez la IA se aplica, entonces cambia de nombre y se empieza a llamar, por ejemplo, traducción (o procesamiento del lenguaje natural), optimización logística, reconocimiento de imágenes, recomendación, *scoring* de riesgos, *pricing* dinámico, Alexa, Netflix o Google. Es decir, estamos rodeados de IA, pero se sigue percibiendo más como imágenes de la película *Terminator* que como las aplicaciones que acabamos de mencionar.

La IA, en origen, quería representar la capacidad de las máquinas para llevar a cabo tareas que se atribuyen normalmente a los humanos. Pero, ¿qué significa esto? ¿es una calculadora que hace cuentas mucho más rápido que nosotros? ¿Lo es SIRI? ¿Cuánto más inteligentes deben ser las máquinas para que digamos que son como los humanos? Estas preguntas no tienen una respuesta clara, ni siquiera es comúnmente aceptado que el test de Turing (cuando una máquina es capaz de hacer creer a un humano que no es una máquina) sea una buena manera de

medirla. Lo que sí sabemos es que la forma más eficiente a día de hoy para generar IA, definiéndola como aquello que nos ayuda a resolver un problema de forma automática, masiva y lo más cercana a lo óptimo, es el aprendizaje automático en base a ejemplos o datos. Y todavía nos queda mucho por ver en este sentido, aunque es difícil prever dónde está el límite, si es que existe.

Por tanto, para generar IA, necesitamos aplicar metodologías de aprendizaje automático o *machine learning*, que a su vez necesita de *big data* para poder funcionar.

6. El uso responsable de los datos y los algoritmos

Hemos visto que gracias a los datos podemos generar inteligencias que nos ayudan a tomar decisiones, o que las eliminan. Esto evidentemente es una herramienta muy potente desde muchos puntos de vista (¿quién no quiere decidir mejor o hacerlo con una predicción sobre lo que pasará en el futuro?), pero como cualquier herramienta se puede usar para lo bueno y para lo malo, ya sea de forma deliberada o no.

Hoy día, los algoritmos ya no solo deciden qué canción escuchamos, qué publicidad vemos, o qué carretera tomamos para llegar a nuestro destino sin tráfico, sino que también abren y cierran la puerta a oportunidades reales. Desde el acceso a servicios públicos a personas o a comunidades, al acceso a una entrevista de trabajo, la concesión de un préstamo o la fijación del precio de un seguro, o incluso a ser o no más «observado» por las fuerzas de seguridad (caso real en algunos estados de EE UU), los algoritmos pueden hacer la vida más fácil, pero también más difícil a las personas. Y el problema de fondo reside precisamente en culpar de ello a los propios algoritmos, es decir, a la tecnología, cuando realmente lo que estamos haciendo es replicar los prejuicios o imitar el planteamiento moral de quienes los poseen y aplican. Ahí es donde hay que buscar la responsabilidad sobre si lo que hace un algoritmo es o no aceptable. La regulación

evidentemente debe jugar y juega un papel clave, pero va a ser complicado que se puedan regular todas las posibilidades presentes y futuras de la IA. Y de hecho no es deseable que así ocurra porque seguramente esto limitaría muchos usos positivos que son sacrificados para evitar que los negativos se den. Limitaría la innovación. ¿Qué podemos hacer entonces? Pues debemos pensar, como empresas, instituciones y como sociedad, que esto va mucho más allá de la regulación, y que tiene que ver con la ética y con la visión que tenemos para el mundo en el que queremos que vivan nuestros hijos y nietos. Y aquí entran en juego cuatro conceptos en los que trabajar desde una visión mucho más amplia que la del cumplimiento normativo: la autorregulación, la transparencia, la educación de ciudadanos y clientes sobre los usos de los datos (*data literacy*), y por último el diseño de servicios centrado en las personas.

Sin ánimo de ser exhaustivos, nos permitimos exponer aquí una serie de preguntas o *checklist* que creemos debería aplicarse cualquiera que esté desarrollando, aplicando o poniendo en valor datos y algoritmos inteligentes o de automatización. Y haríamos bien siendo transparentes en cómo y en qué medida nuestros productos, servicios o sistemas de decisión puntúan en una *checklist* de este tipo:

— Desde el punto de vista más técnico, desafiar siempre nuestras propias matemáticas, no creernos el primer resultado que satisface el objetivo, buscar, como hacen los científicos, el *peer review*, o la revisión por parte de otros expertos para validar la calidad de los modelos.

— Asegurarnos de que los datos de origen no contienen sesgos humanos que pueden desvirtuar el resultado, que matemáticamente podría ser correcto, pero quizás inaceptable desde la ética.

— Preguntarnos ¿qué es lo peor que se podría hacer con este algoritmo? no significa que lo vayamos a hacer, pero solo hacernos la pregunta ya nos da pistas sobre hacia dónde no debería ir nunca su uso, o lo que sucedería si cayera en según qué

manos. También es válida la pregunta: ¿si el funcionamiento y resultados de un algoritmo se hacen públicos, que ocurriría?

— ¿A quién estamos «empoderando» con el algoritmo? ¿quién o quiénes se benefician de su uso? No se trata de ser ingenuo y pensar que las empresas no van a buscar el beneficio, pero también se puede hacer aportando valor a los clientes, o a la sociedad, o al menos no quitándoselo.

— ¿Cómo de accionables son las variables que usa el algoritmo para las personas o comunidades sobre las que se aplica? Por ejemplo, si un algoritmo discrimina por la raza para dar acceso a un recurso, tendría poco sentido, ya que la raza es una variable sobre la que el sujeto tiene poco que hacer.

— ¿Cómo de necesario es comprender el funcionamiento de un algoritmo, y cómo de transparentes somos con ello? Por ejemplo, entender cómo funciona un algoritmo de traducción no es necesario para que su uso sea legítimo y útil (lo importante es que traduzca bien), pero sí es necesario entender por qué a alguien se le deniega un préstamo, tanto para ver si se está discriminando injustamente, como para entender cómo puede mejorar esa persona de cara al futuro.

7. El futuro y las implicaciones de la automatización

Hay otro aspecto relevante a considerar, y que tiene que ver con el potencial de generación de desigualdad de las tecnologías de IA, apuntado ya desde diferentes ámbitos académicos, empresariales, sociales y políticos, aunque no siempre observado desde el mismo rigor y desde luego no siempre con la intención de ser constructivos, que es la única forma de llevar todo esto a buen lugar.

Desde un punto de vista, el origen de la desigualdad se puede basar en el acceso a estas tecnologías desde solo algunos estratos sociales o empresariales. Una buena forma de explicar esto es con la «paradoja

de la predicción», la cual nos dice que cuanto más capacidad se tiene de predecir el futuro, más posibilidades hay de que esa predicción no se cumpla, ya que al conocerla podemos actuar sobre ella a nuestro interés. Pero claro, si solo unos pocos tienen acceso a dicha predicción, la cambiarán en su beneficio mientras el resto seremos meros espectadores y/o afectados por el resultado de ese cambio.

Desde otro punto de vista, el uso masivo de IA y automatización traerá consigo eficiencia y aumento de la productividad gracias en parte a la eliminación de humanos de muchos procesos. Esto, inevitablemente, conllevará un ajuste de los sistemas productivos que puede tener como víctima más o menos permanente el empleo. Esto no es nuevo y ha venido ocurriendo sistemáticamente desde que el hombre empezó a crear herramientas y tecnología, acelerándose en momentos clave como la Revolución Industrial. No sabemos si esta vez será diferente, pero sí que el tipo de trabajos afectados será mucho mayor que antes, como también lo será la inteligencia creciente de los sistemas autónomos. Sea como sea, esto puede ser también un generador de desigualdad, ya que las mejoras en la productividad se pueden concentrar en unas pocas manos, a costa de muchas personas que ya no son útiles como recurso productivo.

Aquí hay muchas opiniones, muchas veces polarizadas. Nosotros somos de la opinión de que el riesgo existe y por eso el impacto debe ser observado muy de cerca, al igual que sucede con la ética en su uso de algoritmos, pero que tras un proceso inevitable de ajuste los humanos volverán a encontrar su lugar. Y esto es así, porque si en algo son buenos los humanos y donde las máquinas van a tener difícil llegar, es a enfrentarse con éxito a situaciones nuevas, que requieren enfoques creativos o laterales, y no repetitivos. Como hemos explicado, los algoritmos, al fin y al cabo, solo pueden aprender de (muchos) datos históricos, con lo que siempre estarán replicando estructuras y patrones anteriores.

En cualquier caso, es difícil prever el impacto en los próximos 15 a 20 años, con lo que de más allá ni hablamos. A lo que sí debemos enfrentarnos inmediatamente es al hecho de que culpando a la tecnología de las futuras desigualdades, ni las va a parar, ni las va a solucionar. Quizás sea el momento de ser creativos (y esto lo tendremos que hacer los humanos), y pensar en modelos de redistribución de la riqueza distintos al actual. Al fin y al cabo, hubo un tiempo en que ideas como el voto universal o el estado del bienestar fueron consideradas utopías, en el mejor de los casos, o una transgresión contra lo establecido, en el peor. Lo que sí es seguro que sucederá es que habrá algoritmos que nos ayudarán en la transición que nos espera.

Referencias bibliográficas

- [1] ANDREW Ng (2017). *What Artificial Intelligence Can and Can't do Right Now*.
- [2] DOMINGOS, P. (2015). *The Master Algorithm*.
- [3] GIRARDIN, F. (2016). *Experience Design in the Machine Learning Era*.
- [4] KELLEHER, J.D; MAC NAMEE, B. y D'ARCY, A. (2015). *Fundamentals of Machine Learning for Predictive Data Analytics: Algorithms, Worked Examples and Case Studies*. MIT Press.
- [5] LATHIA, N. (2014). *Machine Learning for Product Managers*.
- [6] LEVY, S. (2016). *How Google is Remaking Itself as a «Machine Learning First» Company*.
- [7] LINKEDIN (2013). *A Brief History of LinkedIn*. En www.linkedin.com
- [8] O'NEILL, C. (2016). *Weapons of Math Destruction*.