

INTELIGENCIA ARTIFICIAL Y SEGUROS*

ARTIFICIAL INTELLIGENCE AND INSURANCES

FELIPE TABARES CORTES**

Fecha de recepción: 17 de octubre de 2025

Fecha de aceptación: 30 de octubre de 2025

Disponible en línea: 30 de diciembre de 2025

Para citar este artículo/To cite this article:

Tabares Cortes, Felipe. *Inteligencia artificial y seguros*, 63 Rev.Ibero-Latinoam.Seguros, 17-42 (2025). <https://doi.org/10.11144/Javeriana.ris63.iays>

doi:10.11144/Javeriana.ris63.iays

* Artículo de Reflexión.

** Abogado del CSJ de Colombia y de la barra de Francia (Paris), candidato a Doctor, Université Sorbonne-Nouvelle, Especialista en Derecho Médico de la Universidad del Rosario, Maestría en Derecho de Seguros de l'Université Lyon III Jean Moulin, Preparación del Centro de Formación de Abogados de l'Université Paris II Panthéon-Assas, Academy on International Investment Disputes, Kuala Lumpur Regional Centre for Arbitration, Malasia. Ejercicio profesional en Derecho de Seguros y Derecho del Comercio Internacional e Industrial, en firmas de abogados y empresas multinacionales en Francia, Singapur y Colombia. Contacto: felipetaba@gmail.com www.fr.linkedin.com/in/felipetabarescortes.

RESUMEN

La inteligencia artificial (IA), compuesta de tecnologías como Big Data y Machine Learning, ha revolucionado la actividad de las empresas aseguradoras al automatizar y optimizar procesos como la tarificación, la gestión de pólizas y siniestros, y la atención al cliente. Sin embargo, su uso implica riesgos éticos y jurídicos, especialmente en la protección de datos personales y la posibilidad de sesgos discriminatorios en la toma de decisiones automatizadas. En su función reguladora de la Unión, el Parlamento Europeo y el Consejo han establecido una regulación específica (Reglamento UE 2024/1689) que define sistemas de IA de alto riesgo, prohíbe ciertos usos, exige transparencia, supervisión humana y gestión de riesgos, y reconoce derechos ciudadanos frente a decisiones tomadas por IA. El presente artículo pretende desarrollar esta reglamentación que entrara en vigor en agosto de 2026, poniendo en contexto sobre el uso de esta tecnología por las compañías aseguradoras, las perspectivas y los riesgos.

Palabras clave: inteligencia artificial, macrodatos, regulación, seguros, datos personales, sesgo algorítmico.

ABSTRACT

Artificial intelligence (AI), composed of technologies such as Big Data and Machine Learning, has revolutionized the activity of insurance companies by automating and optimizing processes such as pricing, policy and claims handling, and customer service. However, its use involves ethical and legal risks, especially regarding the protection of personal data and the possibility of discriminatory biases in automated decision-making. In its regulatory role within the Union, the European Parliament and Council have established specific regulation (EU Regulation 2024/1689) that defines high-risk AI systems, prohibits certain uses, requires transparency, human supervision, and risk management, and recognizes citizens' rights regarding decisions made by AI. This article aims to develop this regulation, which will come into force in August 2026, providing context on the use of this technology by insurance companies, the perspectives, and the risks.

Keywords: artificial intelligence, big data, regulation, insurance, personal data, algorithmic bias.

SUMARIO:

Introducción. 1. Implicaciones de la inteligencia artificial en los seguros. 1.1. La inteligencia artificial como tecnología disruptiva. 1.1.1. Los fundamentos técnicos de la inteligencia artificial. 1.1.2. Aplicaciones y oportunidades de la inteligencia artificial. 1.2. Los riesgos que implica el uso de la inteligencia artificial. 1.2.1. El uso de datos por las empresas aseguradoras y sus límites. 1.2.2. Los riesgos específicos del uso de la IA para las empresas aseguradoras. 2. La reglamentación europea sobre el uso de la inteligencia artificial. 2.1. La gobernanza de la inteligencia artificial en Europa. 2.1.1. Los objetivos reglamentarios y los sujetos pasivos de la regulación. 2.1.2. La inteligencia artificial de alto riesgo. 2.1.3. Los usos prohibidos de la inteligencia artificial. 2.1.4. Los requisitos de utilización de la IA por las empresas. 2.1.5. Derechos Ciudadanos. 2.2. Régimen jurídico de la Inteligencia artificial en la Unión Europea. 2.2.1. Control jurídico del algoritmo. 2.2.2. Control por supervisión humana y explicabilidad (comprensibilidad) del algoritmo. 2.2.3. Potestades de las Autoridades Estatales de supervisión. 2.2.4. Régimen Sancionatorio. Conclusiones. Bibliografía.

INTRODUCCIÓN

Cuando se habla de inteligencia artificial (IA), es común pensar en el programa digital creado por la empresa OpenAI, pero suele pasarse por alto que la digitalización de todas las actividades humanas ha implicado una utilización constante de algoritmos en la vida cotidiana, abarcando casi la totalidad de las acciones realizadas día a día. Cada vez que una persona instala una aplicación, publica una foto, paga una factura, chatea con sus vecinos, comparte opiniones políticas o da “like” a publicaciones de otros, la inteligencia artificial (o sus componentes Big Data y Machine Learning) está realizando un trabajo constante, soterrado y arbitrario de registro y compilación de todos los contenidos. La digitalización apenas ha comenzado a integrarse en la vida diaria de las personas, las empresas y las instituciones, pero ya sus implicaciones para el sector asegurador son profundas, reorientando la estrategia global de las compañías y estableciendo nuevas formas de identificar, cuantificar y gestionar los riesgos.

La IA ha provocado un cambio profundo al automatizar, optimizar y desarrollar las operaciones (o segmentos de las operaciones) propias de las funciones internas de las compañías de seguros, como la tarificación, la suscripción, la gestión de pólizas y siniestros, y la atención al cliente. Esta tecnología, basada en el análisis avanzado de datos (*Big Data*, *Behavioural Analytics*, *Internet of Things*, *Telematics*), ha acelerado el ritmo de trabajo y las capacidades técnicas de las empresas aseguradoras. Sin embargo, esta aceleración conlleva riesgos, como las dificultades de control, la falta de transparencia en el razonamiento del sistema, los desafíos éticos en el uso de los datos de los clientes y la dependencia absoluta a los sistemas informáticos. En este texto, nos proponemos abordar algunas implicaciones generales de la inteligencia artificial en los seguros (1), para luego analizar la reglamentación europea (Reglamento (UE) 2024/1689 que establece reglas armonizadas sobre la inteligencia artificial) sobre el uso de estas tecnologías (2).

1. IMPLICACIONES DE LA INTELIGENCIA ARTIFICIAL EN LOS SEGUROS

En este texto pretendemos enunciar la manera en que la IA ha cambiado profundamente el aspecto operacional de las compañías de seguros específicamente en los temas de suscripción y emisión de pólizas, redacción de contratos, gestión de siniestros, reclamaciones y litigios. En este artículo será abordado el tema de la revolución de la IA que ha cambiado la dinámica del sector asegurador aportando grandes beneficios, pero habiendo también introducido riesgos, específicamente jurídicos y éticos. Antes de entrar de lleno en el tema debemos comenzar por recordar el contexto general del tema, las particularidades técnicas de la inteligencia artificial y su uso en el mundo de los seguros (1.1). Posteriormente mencionaremos el impacto del uso de datos, los riesgos y desafíos éticos, así como las exigencias de política pública al respecto (1.2).

1.1. La inteligencia artificial como tecnología disruptiva

Esta tecnología se considera disruptiva ya que transforma la manera en que piensan y se comportan los individuos, las empresas y los gobiernos, creando nuevas estrategias de acción para los actores, pero también abriendo nuevos riesgos de tipo técnico, jurídico y deontológico. Como veremos, el término inteligencia artificial engloba varios mecanismos tecnológicos (*big data*, *machine learning*, *deep learning*, *natural text processing*, *internet of things*). Su característica principal es la capacidad para procesar un volumen extremadamente amplio de datos, analizarlos, razonar y tomar decisiones de manera similar a los humanos. La IA posee una notable capacidad para asumir en espacios de tiempo mucho más breves y a costos significativamente menores funciones que tradicionalmente han sido humanas. La automatización de funciones empresariales e institucionales, que es el resultado directo de la creación de la IA, ha comenzado a transformar de manera radical la forma, velocidad y organización de nuestro trabajo. Por ello, en este acápite nos centraremos en mencionar algunos elementos sobre el funcionamiento de la IA (1.1.1), así como en sus aplicaciones y oportunidades específicas para la industria aseguradora (1.1.2).

1.1.1. Los fundamentos técnicos de la inteligencia artificial

La IA tiene como objetivo imitar diversas funciones cognitivas humanas, como la percepción, la memoria, el razonamiento y el aprendizaje, así como reproducir competencias tales como la organización, la descripción y el tratamiento de la información. Además, abarca funciones más específicas como el reconocimiento de voz, la toma de decisiones y la traducción de idiomas. Se trata de un conjunto de tecnologías informáticas que simulan el funcionamiento del cerebro humano de manera autónoma, gracias al perfeccionamiento en el uso de algoritmos. Los algoritmos se emplean en la computación porque establecen estructuras condicionales para la toma de decisiones (Jean, 2020), y son los componentes que conforman la IA, permitiéndole realizar tareas propias del intelecto humano.

La información base que analiza la inteligencia artificial se denomina *Big Data* y proviene de la recopilación de la totalidad, o casi totalidad, de la información que

se produce, procesa y almacena en internet. Esta información es entonces utilizada y procesada por la inteligencia artificial mediante algoritmos. Los algoritmos que permiten a la máquina aprender se encargan de identificar relaciones entre los datos y generan modelos predictivos de manera autónoma. El aprendizaje profundo (*deep learning*) es un campo específico del aprendizaje automático (*machine learning*) cuyos algoritmos son especialmente eficaces en el procesamiento de grandes volúmenes de información, incluso cuando los datos son no estructurados, como las imágenes o la voz (Fliche & Yans, 2018).

La IA abarca así todos los instrumentos informáticos que permiten a la máquina aprender (*machine learning*), adaptarse y operar en entornos dinámicos, utilizando algoritmos avanzados que aprenden con cada nuevo registro de datos, ajustando y mejorando continuamente sus predicciones (Kivisaari et al., 2021). Por otro lado, el aprendizaje profundo (*deep machine learning*) es una rama del aprendizaje automático que se basa en modelos estadísticos complejos y en múltiples capas de procesamiento paralelo que “modelan de manera aproximada el funcionamiento del cerebro biológico”. Aquí entran en juego las redes neuronales (*neural networks*), que son algoritmos programados como unidades de aprendizaje que procesan y transforman los datos mediante funciones matemáticas. Estas redes neuronales pueden aprender a realizar tareas a partir de ejemplos, generalmente sin estar programadas con reglas específicas para cada tarea (Kivisaari et al., 2021). Es de esta manera que la máquina puede aprehender representaciones complejas de datos y aumentar el nivel de abstracción, lo que facilita la complejización de la lógica utilizada por el sistema (Jean, 2020).

El *machine learning* utiliza dos instrumentos algorítmicos principales: los árboles de decisión (*decision trees*) y las redes bayesianas (*Bayesian networks* -un tipo de modelo matemático probabilístico predictivo-), lo que permite un seguimiento del procesamiento intelectual por parte de la máquina. Hasta este punto, nos encontramos ante capacidades de automatización derivadas de la matemática, logradas gracias al desarrollo de procesadores informáticos muy potentes, capaces de tratar enormes cantidades de datos. Sin embargo, otras áreas del aprendizaje de máquinas se acercan al funcionamiento biológico, o al menos así se asimila para poder explicar los procedimientos avanzados en el funcionamiento de los algoritmos (Kivisaari et al., 2021).

En conclusión, el desarrollo tecnológico ha permitido la creación de diversas herramientas informáticas que se incluyen dentro de la categoría de IA (*big data*, *machine learning*, *deep learning*, *natural text processing*, *internet of things*), pero que, para efectos de simplificación comprensiva, pueden caracterizarse por otorgar a las máquinas la capacidad de imitar funciones cognitivas humanas y realizar tareas complejas de manera autónoma.

1.1.2. Aplicaciones y oportunidades de la inteligencia artificial

La industria aseguradora ha sido una de las más favorablemente impactadas con el advenimiento de la inteligencia artificial. El marketing, la evaluación de riesgos, la tarificación, la suscripción, la gestión de pólizas, siniestros, reclamaciones y litigios

son funciones propias de las compañías de seguros que han comenzado a transformar sus actividades gracias a la IA. Esto es completamente comprensible: las tecnologías de procesamiento de texto utilizan la información de las bases de datos (pero no únicamente) de las empresas aseguradoras para mejorar, proponer, identificar, adaptar y ejecutar todo tipo de funciones relacionadas con la actividad básica de procesamiento de texto y cifras.

El soporte material del seguro es el contrato, pero también la póliza, que resulta del análisis principalmente matemático sobre el riesgo. Todos estos soportes materiales emplean caracteres numéricos y alfanuméricos que son procesados por sistemas informáticos e interpretados en forma de bits. El derecho, expresado en el contrato y la póliza, es lenguaje, y el lenguaje es la materia principal de procesamiento por la inteligencia artificial; por ello, los sectores industriales basados en ‘letras’ son los primeros en ser transformados por la utilización de la inteligencia artificial.

Como se expondrá en detalle, la inteligencia artificial no utiliza únicamente las fuentes de información clásicas empleadas por las compañías de seguros, sino que también se nutre de otras fuentes muy variadas y extremadamente potentes, provenientes de los datos de comportamiento (*Behavioural Data*), el internet de las cosas (*Internet of Things, IoT*), la telemática (*Telematics*), los rastreadores personales (*Personal Tracker Data* y *GPS*), lo que enriquece el análisis de comportamiento realizado por las empresas de seguros (Kivisaari et al., 2021). Estas nuevas fuentes de información permiten a las aseguradoras tomar decisiones más precisas y predictivas, creando perfiles por microsegmentos de riesgos, lo cual beneficia además a los clientes, quienes pueden recibir ofertas de pólizas más adaptadas a su exposición real (Kivisaari et al., 2021).

Los avances son espectaculares en lo que respecta a la relación con el cliente: la IA permite la optimización de precios mediante tarificación personalizada basada en datos de comportamiento individual que van más allá del simple análisis del riesgo, como la elasticidad de precios (*Price Elasticity*), el valor de vida del cliente (*Lifetime Value*) y la propensión a cambiar de proveedor (*Propensity to Churn*), así como el análisis de las ofertas de la competencia (Kivisaari et al., 2021).

En cuanto al procedimiento de ventas, distribución y marketing, las técnicas digitales basadas en el análisis dinámico del comportamiento de búsqueda en línea de los clientes (*Online Search Behaviour*) permiten captar la atención del consumidor, incrementando la posibilidad de realizar ventas online directamente a través de los sitios web de seguros, así como de los buscadores o comparadores de precios. Los asistentes virtuales (*Virtual Assistant*) y los *Chatbots* (cada vez con mejor procesamiento de lenguaje natural) automatizan la recopilación de información, mejorando la experiencia online y reduciendo costos (Kivisaari et al., 2021). La gestión de relaciones con el cliente mediante IA posibilita una comunicación proactiva, además de mejorar las oportunidades de nuevos negocios mediante la venta cruzada de servicios relacionados (seguros afinitarios), optimizando el servicio y la retención (Pravina Ladva, 2023). Asimismo, el análisis del comportamiento de los clientes en ramos como salud y automóviles permite una mejor comprensión del riesgo, su tarificación y su mitigación (Kivisaari et al., 2021).

La IA también desempeña un papel relevante en la reducción de los costos internos de las empresas de seguros. Por un lado, la automatización robótica de procesos y la gestión electrónica de documentos (pólizas, boletines de cobertura, certificados, contratos) han mejorado la eficiencia de los sistemas, pero también de los grupos de trabajo humanos, reduciendo costos administrativos y liberando a los empleados de tareas rutinarias y de bajo valor intelectual. La resolución automatizada de quejas también reduce el tiempo de respuesta y agiliza la comunicación con los clientes (Kivisaari et al., 2021).

Otro sector con grandes posibilidades gracias a la llegada de la IA es la gestión de siniestros. Estos representan una parte significativa de la carga administrativa en las compañías de seguros, ya que requieren procesos complejos (muchas veces manuales) de evaluación y provisión de daños, así como de control de pagos (Hufeld, 2022). La IA ha automatizado la notificación inicial de pérdida (*FNOL*) y ha acelerado las tareas de verificación de siniestros mediante el seguimiento telemático o predictivo con sensores, *IoT* y *GPS*, así como la recopilación de información sobre accidentes, que algunas empresas gestionan a través de tecnologías de reconocimiento de imágenes que facilitan la estimación de perjuicios y la reparación en ciertos ramos (agrícola, habitación y automóviles). Existen además proveedores que emplean drones para capturar imágenes de alta resolución y aplicar IA en la evaluación directa de daños (Hufeld, 2022).

La IA proporciona a los asegurados alertas meteorológicas y recomendaciones específicas para proteger sus bienes, lo que puede influir en la prima si no se siguieron las advertencias (Hufeld, 2022). En seguros de inmuebles, el análisis de datos meteorológicos y de dispositivos conectados (*IoT*) permite detectar riesgos como incendios por fallos eléctricos y alertar automáticamente a la aseguradora sobre la ocurrencia de daños. De este modo, los siniestros pueden ser identificados en tiempo real, permitiendo iniciar medidas como avisar a la policía y comenzar el proceso de liquidación de siniestros incluso si cuando el asegurado no está presente (Hufeld, 2022).

Además, algunas empresas han comenzado a utilizar la IA y, más concretamente, las funciones de *machine learning* para estimar las reservas técnicas (*Loss Reserving*), los riesgos judiciales y, en general, las previsiones financieras de los siniestros (Kivisaari et al., 2021). Así, algunas compañías han empezado a desarrollar el uso de la IA para automatizar verificaciones sobre la plausibilidad de los hechos relatados por los clientes, así como sobre el alcance y extensión de los daños, utilizando fuentes de datos adicionales, tanto jurídicas como técnicas del sector en cuestión. Se prevé que en el futuro cercano ciertos ramos serán completamente gestionados por la IA, dependiendo sin embargo de la cuantía y la complejidad del caso (Hufeld, 2022).

Por último, la prevención de fraude y la lucha contra el lavado de dinero y la financiación del terrorismo (*LCB-FT*), funciones propias de los departamentos jurídicos de las compañías de seguros, emplean cada vez más la inteligencia artificial en sus labores de control. Las técnicas de IA se utilizan especialmente para el reconocimiento, análisis y validación de los documentos presentados, así como para cruzar información, identificar partes, controlar el tránsito bancario y detectar transacciones fraudulentas (Fliche & Yans, 2018). Así, algunas empresas han comenzado a filtrar automáticamente

los siniestros según patrones identificados con *Big Data*, y posteriormente una unidad investigadora revisa manualmente los casos sospechosos para evitar fraudes o errores (Hufeld, 2022).

1.2. Los riesgos que implica el uso de la inteligencia artificial

Las cuestiones relativas a la ética en el uso de los datos y las tecnologías digitales requieren un enfoque más amplio que el simple cumplimiento de las normas legales, de manera que se deben también considerar los riesgos específicos para la actividad de las empresas aseguradoras. Por ello conviene de presentar en un primer momento cuales es el uso de la *Big Data* por las aseguradoras (1.2.1) para luego mencionar los riesgos específicos que implica para este sector la utilización de la inteligencia artificial (1.2.2).

1.2.1. El uso de datos por las empresas aseguradoras y sus límites

La industria de seguros siempre ha requerido el uso de datos como medio para identificar los riesgos que busca garantizar a través del contrato de seguro. La recuperación de datos solía generar inconvenientes en algunos ramos del seguro, especialmente tratándose de información de muy diverso orden (por ejemplo, sobre catástrofes naturales, sobre riesgos de endeudamiento del asegurado, sobre los cambios jurisprudenciales en un país dado). A partir de los datos recopilados, la ciencia actuarial analiza la posibilidad concreta de cubrir el número de asegurados siniestrados con las reservas técnicas establecidas para ese tipo de riesgo. Sin embargo, gracias a la digitalización, el análisis de riesgos realizado por las empresas de seguros ha sido drásticamente transformado con la llegada de la inteligencia artificial, y más específicamente con uno de sus componentes principales: la *Big Data* (Macrodatos o Datos Masivos).

El desarrollo industrial y la eficiencia empresarial han sido profundamente impulsados por el uso de la tecnología digital. En la actualidad, todas las empresas, instituciones, personas y, en general, todo tipo de organizaciones tienen presencia virtual. Esta presencia virtual aporta una cantidad de información (data) muy detallada sobre esas personas, empresas e instituciones. La inteligencia artificial tiene la capacidad para analizar esta información en beneficio primero de la sociedad y, posteriormente, de los intereses privados. Así, cuando una persona utiliza redes sociales, cuando una empresa realiza pagos, cuando una institución pública publica información en línea, o cuando una alcaldía transfiere impuestos a las rentas nacionales, en todas estas situaciones existe un flujo de información a través de los sistemas informáticos, flujo que constituye la *Big Data*, así como también el *Internet of Things*, la información sobre uso transferida subrepticamente por todos los dispositivos conectados a las empresas propietarias de dichos medios de comunicación.

La *Big Data* proviene entonces de datos propios de las empresas (datos de clientes, como datos personales o de productos, por ejemplo), así como de datos obtenidos a través de terceros, como la información recopilada por los motores de búsqueda, las transacciones crediticias, los detalles de identificación de las personas de las páginas web, su estado de salud (*Fitbit* y *Apple Watch*), la siniestralidad pasada, las cotizaciones pagadas, el uso de sitios web, los detalles de uso de redes sociales (Facebook, Twitter, Instagram, TikTok, LinkedIn) y los datos de dispositivos conectados, como el uso de celulares (ubicación, cercanía a otros dispositivos, desplazamiento, hora de uso, información biométrica del utilizador) y de objetos telemáticos en automóviles, drones y bicicletas (Kivisaari et al., 2021).

Toda esta información se denomina *Big Data* y no necesariamente se refiere a información específica de personas naturales, ya que los datos suelen estar anonimizados o semianonimizados, o por lo menos es lo que se afirma... Mientras que la *Big Data* se refiere principalmente a la información obtenida por todos los medios, la IA es el proceso de aprendizaje automático (*machine learning*) que utiliza esos metadatos y macrodatos para explotarlos y obtener información específica (Kivisaari et al., 2021). La *Big Data* es, entonces, la forma de obtener información sobre el funcionamiento de todos los estamentos sociales, mientras que la inteligencia artificial es el método de procesamiento de esa información para fines determinados por el Estado y las empresas.

Lo particular es que las empresas de seguros sí tienen la posibilidad de utilizar esta información –*Big Data*– a través de la IA (incluso información personal y/o sensible de personas determinadas) bajo ciertas condiciones establecidas en la legislación europea: primero, el Reglamento (UE) 2016/679 de 27 de abril de 2016 relativo a la protección de las personas físicas en lo que respecta al tratamiento de datos personales y a la libre circulación de estos datos, y segundo, el Reglamento (UE) 2024/1689 de 13 de junio de 2024, por el que se establecen normas armonizadas en materia de inteligencia artificial.

Las facultades de las empresas aseguradoras han estado delimitadas por consideraciones jurídicas sobre el trato justo del grupo asegurado y de cada tomador de pólizas. Este análisis jurídico se extiende incluso a la obligación de la empresa aseguradora de gestionar con extrema prudencia la mutualidad del seguro, al menos en lo que respecta a los intereses de personas físicas e intereses particulares. Del mismo modo, es sabido que la compañía debe evitar que un solo individuo beneficie indebidamente de la mutualidad. Por ello, el análisis de las cuestiones éticas no es, ni debe ser, confiado exclusivamente a la inteligencia artificial, pues se trata de un análisis de carácter jurídico, y por ello se considera que existe una obligación de control humano sobre las decisiones más delicadas, ya que en el fondo se terminan tomando decisiones sobre los derechos fundamentales de las personas (físicas y jurídicas).

Este control se eleva así en una obligación fundamentar para evitar, por ejemplo, que el sistema no interprete erróneamente el reparto de riesgos según la Ley de Grandes Números y evitar entonces que se discrimine ciertos riesgos que, tras un análisis de datos y de eficacia económica, se considere que no deberían ser cubiertos (Kivisaari et al., 2021).

Aquí, nuevamente, interviene el Derecho para evitar situaciones de discriminación y exclusión de personas. El principio de trato justo, de utilización de la información, la equidad contractual y, en general, el Estado de Derecho prohíben los actos discriminatorios y reconocen la necesidad de medidas afirmativas en favor de la protección de grupos vulnerables (Kivisaari et al., 2021). Las instituciones, entre ellas las compañías de seguros, adquieren gran poder frente a las personas naturales y algunas jurídicas, ya que pueden obtener datos concretos de las personas con las que interactúan, a través de redes sociales (Facebook, Instagram, Twitter, TikTok, LinkedIn), además de la *metadata* y la *macrodata* que obtienen de otras instituciones como el sistema bancario o los sistemas de seguridad del Estado.

Esta capacidad de obtener información con tal grado de precisión es lo que se denomina *Granularidad*, y aquí reside uno de los objetivos regulatorios de la Unión Europea con el Reglamento de 2016 sobre Datos Personales (*GDPR*) y el Reglamento de 2024 sobre Inteligencia Artificial (*IA Act*). Estas consideraciones jurídicas son las consecuencias esperables de la aplicación de principios constitucionales del Estado de Derecho (concepción occidental) a la utilización de la *Big Data* y la inteligencia artificial por parte de entidades bancarias y aseguradoras.

Otro aspecto regulado respecto a la *Big Data* es la protección de la información de salud y financiera de las personas físicas. El Reglamento General de Protección de Datos (*GDPR*) de la Unión Europea otorga ciertos derechos a las personas, como la posibilidad de exigir la eliminación de sus datos de las bases de datos de las compañías, más conocido como el “derecho al olvido” (Kivisaari et al., 2021) que ha tomado mucha importancia respecto de los pacientes con cáncer. Además, este reglamento exige que la información privada se borre sin demora cuando ya no sea necesaria para el propósito originalmente buscado. También obliga a las empresas a mantener un registro preciso del consentimiento del titular de los datos para su uso en fines primarios y secundarios, sin el cual no pueden utilizar dicha información (Kivisaari et al., 2021).

Sea pertinente precisar que el uso de datos personales no es en todas las circunstancias negativo; por el contrario, la política pública sobre el tema puede tener una influencia benéfica para la sociedad. Cubrir una mutualidad de riesgos implica disponer de datos específicos sobre todos los individuos que integran esa mutualidad. El aumento de la cantidad de datos gracias a la digitalización y al procesamiento de los macrodatos mejora la analítica predictiva (*Predictive Analytics*) del proceso de tarificación de las coberturas, permitiendo así precios mejor adaptados para cada tipo de riesgo y, por lo tanto, más ajustados (granulares) al comportamiento del tomador de póliza.

La *Big Data* permite además al Estado determinar en realidad cuales son los sectores o grupos en los cuales la cobertura es posible, a pesar de lo que afirmen las compañías de seguros frente a grupos inasegurables como los médicos obstetras, los anestesiólogos, los cirujanos plásticos, los urgenciólogos, los pilotos de aviación, los controladores aéreos, los ingenieros estructurales, los mineros, entre otros. La IA va permitir a las entidades estatales efectuar un análisis (al menos estimativo) sobre la asegurabilidad real de estos riesgos y establecer entonces obligaciones de cobertura.

Por ello, se afirma que el procesamiento de los macrodatos informáticos tiene el potencial de transformar todo el proceso de producción de seguros (Kivisaari et al., 2021). La industria aseguradora utiliza tradicionalmente herramientas estadísticas como los modelos lineales generalizados (*Generalised Linear Models*) que se utilizan para evaluar y tarificar riesgos. Sin embargo, algunos sectores, como el de las catástrofes naturales, presentan mayor dificultad, ya que las estimaciones históricas no proporcionan una justificación estadística suficiente. Este problema se resuelve en gran medida con la utilización de la Big Data, que permite una tarificación basada en el riesgo (Kivisaari et al., 2021).

1.2.2. Los riesgos específicos del uso de la IA para las empresas aseguradoras

La automatización de los procesos en la industria del seguro aporta, como se ha mencionado, avances significativos; sin embargo, la garantía de control humano sobre el los sistemas IA sigue siendo una de las preocupaciones fundamentales en Europa (EIOPA). Las gerencias jurídicas de las compañías de seguro parecen concentrarse en los riesgos que la IA generativa puede implicar respecto al uso de datos protegidos por derechos de autor y la posible infracción de derechos de propiedad intelectual (Pravina Ladva, 2023). Este riesgo existe, aunque es menor si se tiene en cuenta que la IA efectúa una anonimización de los datos, sin identificar así a las personas individualmente consideradas. No obstante, el uso de la Big Data de clientes directos para entrenar la IA continúa siendo motivo de inquietud, ya que no se conoce con certeza en qué medida las GAFAM (*Google, Amazon, Facebook, Apple y Microsoft*) emplean la información proporcionada por las empresas. Estas empresas han invertido masivamente en la IA y los *datacenters*, a pesar de manejar gran opacidad sobre el funcionamiento real de estas tecnologías y los objetivos que ellos persiguen, provocando entonces la oposición europea que habla entonces de un *tecnofeudalismo*¹.

Pero contrariamente a la creencia de los departamentos de conformidad (*compliance*) de las compañías de seguros, el principal riesgo radica en la dificultad para identificar los análisis probabilísticos realizados por los algoritmos. Explicar las razones por las cuales la IA llega a una conclusión determinada no es sencillo, dado que el procedimiento analítico llevado a cabo por la IA no es fácilmente discernible. Esto complica la gobernanza y el control de esta tecnología, considerando que puede perpetuar la discriminación, ser vulnerable a ataques maliciosos o causar perjuicios. Los desafíos de gestión, por tanto, no son menores, de manera que se habla de fallos algorítmicos (*algorithmic failures*) cuando la IA puede generar análisis inexactos o establecer criterios de tarificación injustos (Pravina Ladva, 2023), llegando incluso a excluir grupos poblacionales para los cuales los criterios comerciales y actuariales de la mutualidad no se cumplen. Esto se conoce como sesgo algorítmico (*algorithmic bias*), que resulta de datos limitados o no representativos (Kivisaari et al., 2021), generando una selección de riesgos incorrecta (injusta o ilegal).

¹ Morozov, E. (2025, August 1). ¿Nos está devolviendo la tecnología digital a la Edad Media? *Le Monde Diplomatique, En Español*.

Otro reto en cuanto a la conformidad jurídica del sistema de IA se encuentra en el tipo de datos utilizados. Cuando se trata de datos provenientes de redes sociales (*Social Media Data*) o de navegación (*Clickstream Data*), las empresas deben respetar los criterios establecidos principalmente en el Reglamento (UE) 2024/1689 y el Reglamento (UE) 2016/679 para poder utilizar dichos datos. Por ejemplo, si la aseguradora debe abstenerse de emitir una póliza debido a información negativa encontrada en internet sobre una persona en particular que busca ser cubierta por un seguro obligatorio, el potencial cliente puede impugnar la decisión, pero en ese caso, la aseguradora no podrá justificarla si tomo los datos de redes sociales o de la navegación web de esta persona. Los principios de equidad y transparencia se ven así afectados, aunque desde una perspectiva estrictamente comercial, la compañía aseguradora esté justificada para no proceder en este caso con el contrato.

De esta manera, frente a las personas físicas el riesgo principal es que la IA sea sesgada en su análisis. Los sesgos algorítmicos (*algorithmic bias*) pueden estar presentes en variables como el género, raza, religión, orientación sexual y ser implícitos, difíciles de detectar y requerir análisis de un experto (*Pravina Ladva, 2023*). El algoritmo puede reforzar entonces discriminaciones y generar tratamientos injustos, como segregar por departamento/lugar de residencia o respecto de clientes jóvenes con poco historial. Estos efectos adversos pueden dificultar la inclusión financiera, ya que los consumidores vulnerables, como personas mayores, con bajo nivel educativo o con ingresos bajos, pueden verse afectados por estrategias de comercialización centradas en los sectores más lucrativos (Kivisaari et al., 2021, Fliche & Yans, 2018). Estos errores de funcionamiento de la IA pueden surgir de un entrenamiento de la IA con datos insuficientes (*flawed algorithm design*) o en una interpretación incorrecta de resultados provenientes de datos parciales y no objetivos (*Pravina Ladva, 2023*).

Estas situaciones permiten concluir que una forma esencial de evitar sesgos en los resultados de la IA es verificar las fuentes, la pertinencia y la exhaustividad de los datos, asegurando que sean representativos para evitar exclusión injustificada o discriminación. Esto es lo que se entiende por *calidad* de los datos, que debe ser garantizada en primer lugar por los *data scientists* (Fliche & Yans, 2018), profesionales que emplean programas informáticos especializados para identificar patrones, tendencias y relaciones. Como veremos más adelante, el Reglamento (UE) 2024/1689 busca mitigar estos riesgos al establecer la obligación de asegurar la supervisión humana y la explicabilidad de los resultados de la IA, así como ciertos derechos de protección para las personas físicas que interactúan con estas plataformas.

Las posibilidades son inmensas, pero los riesgos también son significativos. El análisis de comportamiento (*Behavioural Analysis*) puede evolucionar hacia la vigilancia de masa, afectando la privacidad y autonomía ciudadana, y alterando las relaciones de poder entre consumidores y empresas (Kivisaari et al., 2021). Así, el *Big Data* no está lejos de convertirse en el *Big Brother*, y en tiempos de gobiernos autocráticos y represivos como el nuestro, sumado al desarrollo de la tecnología óptica y la proliferación de drones, nuestras sociedades no están lejos de una distopía totalitaria que ya asoma sus tentáculos en los regímenes de China, Rusia y Venezuela. No es difícil tampoco identificar esta tendencia en occidente al observar la persecución

intensa y violenta del presidente Trump contra sus opositores políticos o la represión policial contra los movimientos sociales en Francia.

Estos regímenes no dudan en utilizar a las grandes empresas tecnológicas - *Big Tech* o *GAFAM* (*Google, Amazon, Facebook, Apple y Microsoft*)—para ejercer control social, perseguir opositores, emplear terrorismo de Estado y silenciar voces disidentes. La sociedad civil debe, por tanto, ser extremadamente cautelosa, y la protección jurídica establecida en el Reglamento de 2016 sobre Datos Personales y el Reglamento de 2024 sobre Inteligencia Artificial representan un bastión de resistencia del Estado de Derecho frente al poder de los oligopolios revitalizados por el tecnofeudalismo.

2. LA REGLAMENTACION EUROPEA SOBRE LE USO DE LA INTELIGENCIA ARTIFICIAL

El Reglamento (UE) 2024/1689 que establece reglas armonizadas sobre la inteligencia artificial fue adoptado por el Parlamento y el Consejo de la Unión Europea el 13 de junio de 2024 estableciendo normas claras para el uso seguro y responsable de sistemas de IA, y buscando principalmente proteger la salud, la seguridad y los derechos fundamentales de los habitantes de la Unión. Esta reglamentación definió qué sistemas son de alto riesgo y estableció las reglas que exigen una gestión de riesgos, la supervisión humana, la transparencia y la realización de evaluaciones de impacto. El Reglamento fue innovador en el sentido de haber otorgado a las autoridades la potestad de supervisar toda actividad relacionada con la inteligencia artificial, además de haber reconocido el derecho ciudadano de reclamar o pedir explicaciones sobre decisiones tomadas con base en estas tecnologías. El Reglamento además estableció sanciones estrictas para garantizar el cumplimiento y la confianza en la IA en Europa. Valga la pena precisar que este texto no aborda el Reglamento (EU) 2016/679 sobre la Protección de Datos Personales (RGPD), puesto que, a pesar de estar estrechamente vinculado al tema de la Inteligencia Artificial, la extensión del tema nos lo impide. Debemos entonces concentrarnos en el Reglamento (UE) 2024/1689 para presentar en un primer momento cuales son las disposiciones de gobernanza tecnológica sobre el uso de la inteligencia artificial (2.1) para luego señalar cual es el régimen de control y de sanción establecido en este tema por la Unión Europea (2.2).

2.1. La gobernanza de la inteligencia artificial en Europa

La Unión Europea ha establecido unos principios de gobernanza de la inteligencia artificial: transparencia, explicabilidad, equidad y supervisión humana (Kivisaari et al., 2021). Dichos principios fueron desarrollados en la regulación sobre inteligencia artificial que establece los objetivos reglamentarios y los sujetos pasivos de la normativa (2.1.1), además de definir lo que se entiende por inteligencia artificial de alto riesgo (2.1.2). La gobernanza también delimita la actividad de la IA al prohibir ciertos usos que puedan afectar derechos fundamentales (2.1.3), imponiendo además requisitos específicos para su utilización por parte de las empresas (2.1.4). La reglamentación menciona además cuales son los derechos ciudadanos frente a los riesgos derivados de estas tecnologías (2.1.5).

2.1.1. Los objetivos reglamentarios y los sujetos pasivos de la regulación

El Artículo 1, apartados 1 y 2 establece que el objetivo del Reglamento es mejorar el funcionamiento del mercado interior y promover una inteligencia artificial (IA) centrada en el ser humano, garantizando un alto nivel de protección de la salud, la seguridad y los derechos fundamentales, incluyendo la democracia, el Estado de Derecho y la protección del medio ambiente, frente a los efectos perjudiciales de los sistemas de IA en la Unión, así como apoyar la innovación.

Para ello, el Reglamento fija normas armonizadas para la introducción en el mercado, la puesta en servicio y la utilización de sistemas de IA. Particularmente, el Reglamento prohíbe determinadas prácticas de IA establece requisitos específicos para los sistemas de IA de alto riesgo y establece obligaciones para los operadores de dichos sistemas definiendo las reglas relacionadas con la transparencia.

El Artículo 2, apartado 1, literales b) y g) determina que el Reglamento se aplicará a los responsables del despliegue de sistemas de IA que estén establecidos o ubicados en la Unión, así como a las personas afectadas que estén ubicadas en la Unión. Este artículo hace entonces aplicable el Reglamento a las empresas aseguradoras que se consideran “responsables del despliegue” es decir, las personas físicas o jurídicas que utilicen un sistema de IA bajo su propia autoridad (Artículo 3, numerales 4), 49), 61) y 65).

2.1.2. La inteligencia artificial de alto riesgo

El artículo 6, apartados 2 y 3 regula la clasificación de los sistemas de IA de alto riesgo. El apartado 2 establece que se considerarán de alto riesgo los sistemas de IA contemplados en el anexo III del Reglamento.

Según este anexo, los sistemas de IA de alto riesgo son aquellos que operan en ámbitos sensibles como la biometría, infraestructuras críticas, educación, empleo, acceso a servicios esenciales, cumplimiento de requisitos jurídicos, migración y administración de justicia. En el área de biometría, se consideran de alto riesgo los sistemas de identificación remota, categorización basada en atributos sensibles y reconocimiento de emociones, excluyendo los sistemas de verificación biométrica simples.

En infraestructuras críticas, los sistemas de IA de alto riesgo son aquellos que se utilizan como componentes de seguridad en la gestión de infraestructuras digitales, tráfico, y servicios básicos como agua, gas, calefacción y electricidad. En el ámbito educativo y profesional, se incluyen los sistemas que determinan el acceso, evalúan resultados, asignan niveles educativos y detectan comportamientos prohibidos durante exámenes.

Respecto al empleo, se consideran de alto riesgo los sistemas de IA que intervienen en la contratación, selección, evaluación de candidatos, toma de decisiones laborales y supervisión del rendimiento.

Por último, en relación con el tema que nos ocupa, este Anexo III indica que son de alto riesgo los sistemas IA que evalúan la admisibilidad y solvencia para servicios públicos y privados esenciales, fijan precios de seguros, gestionan llamadas de emergencia y priorizan servicios de intervención.

Por cierto, el apartado 3 de artículo 6 indica que, no obstante, un sistema de IA del anexo III no se considerará de alto riesgo si no plantea un riesgo importante para la salud, la seguridad o los derechos fundamentales, especialmente si no influye sustancialmente en la toma de decisiones. Esto se aplica si el sistema realiza tareas de procedimiento limitadas, mejora actividades humanas, detecta patrones sin sustituir la valoración humana, o realiza tareas preparatorias para evaluaciones relevantes. Sin embargo, los sistemas que elaboren perfiles de personas físicas siempre serán considerados de alto riesgo.

Por último, el Artículo 7, apartado 1, letras h) y j) otorga a la Comisión Europea la facultad de modificar el anexo III mediante actos delegados, añadiendo o modificando casos de uso de sistemas de IA de alto riesgo cuando exista un desequilibrio de poder o vulnerabilidad de las personas respecto al responsable del despliegue, y cuando el despliegue del sistema de IA resulte beneficioso para personas, colectivos o la sociedad en general, incluyendo mejoras en la seguridad de los productos.

2.1.3. Los usos prohibidos de la inteligencia artificial

El artículo 5 establece una serie de usos prohibidos de la IA en la Unión Europea, principalmente aquellas que emplean técnicas subliminales, manipuladoras o engañosas utilizadas para alterar el comportamiento de las personas, así como sistemas que explotan vulnerabilidades derivadas de la edad, discapacidad o situación social y económica. También se prohíbe el uso de sistemas de IA para la evaluación o clasificación de personas según su comportamiento social o características personales, cuando esto pueda resultar en tratos perjudiciales o desproporcionados.

Asimismo, se prohíbe la utilización de sistemas de IA para la predicción de delitos basada únicamente en perfiles personales, la creación de bases de datos de reconocimiento facial mediante la extracción masiva de imágenes de internet o cámaras, y la inferencia de emociones en lugares de trabajo y centros educativos, salvo excepciones médicas o de seguridad. Igualmente, se restringe la categorización biométrica para deducir la raza, las opiniones políticas, la afiliación sindical, las creencias religiosas, la vida o la orientación sexuales, excepto en contextos legales específicos.

Por otro lado, el mismo artículo del Reglamento indica que el uso de sistemas de identificación biométrica remota “en tiempo real” en espacios públicos solo se permite bajo circunstancias estrictamente necesarias, como la búsqueda de víctimas, prevención de amenazas graves o identificación de sospechosos de delitos graves. Este uso requiere autorización previa de una autoridad judicial o administrativa independiente, salvo en casos de urgencia debidamente justificados.

2.1.4. Los requisitos de utilización de la IA por las empresas

El artículo 9 del Reglamento indica que las empresas que utilizan la IA deben implementar un sistema de gestión de riesgos para sistemas de IA de alto riesgo, el cual debe identificar, analizar y evaluar peligros previsible para la salud, la seguridad y los derechos fundamentales, tanto en el uso previsto como en usos indebidos razonablemente previsible, incluyendo riesgos adicionales detectados tras la comercialización. Las medidas de gestión deben minimizar los riesgos de manera eficaz, mediante diseño adecuado, controles de mitigación, información técnica y formación, considerando el contexto y los usuarios finales.

El artículo 27 del Reglamento indica que debe ser realizada una evaluación de impacto sobre los derechos fundamentales antes de desplegar sistemas de IA de alto riesgo, especialmente cuando se trata de entidades privadas que presten servicios públicos, como puede ser el caso de algunas empresas de seguro.

Un requisito adicional es mencionado por el artículo 50 del Reglamento que indica que, en aplicación del principio de transparencia, los responsables del despliegue (aseguradoras) que ponen a disposición de sus clientes personas físicas sistemas IA, deben informarles que están interactuando con la inteligencia artificial. Además, se debe informar a las personas físicas cuando estos están sometidos a sistemas de reconocimiento de emociones o de categorización biométrica. También deben alertar públicamente sobre cualquier riesgo de “ultrasuplantación”.

2.1.5. Derechos Ciudadanos

El artículo 85 del Reglamento reconoce el derecho de cualquier persona física o jurídica a presentar una reclamación ante la autoridad de vigilancia si considera que se ha infringido el Reglamento, sin perjuicio de otras vías administrativas o judiciales. Las reclamaciones serán gestionadas conforme a los procedimientos específicos y al Reglamento (UE) 2019/1020 sobre vigilancia del mercado. Además de ello, el artículo 86 otorga a las personas afectadas por decisiones basadas en sistemas de IA de alto riesgo el derecho a recibir explicaciones claras y significativas sobre el papel de la IA en el proceso de toma de decisiones, incluyendo los principales elementos de la decisión, especialmente si tiene efectos jurídicos o impacto significativo en salud, seguridad o derechos fundamentales.

2.2. Régimen jurídico de la Inteligencia artificial en la Unión Europea

Una gobernanza digital efectiva implica la realización de un control jurídico sobre el funcionamiento de los algoritmos (2.1.4), indicando además cual es el margen de supervisión humana de estos sistemas (2.1.5). Este control indica además que la supervisión debe ser otorga a través de las potestades propias de las autoridades estatales (2.1.6), además de que se establece un régimen sancionatorio en caso de incumplimiento de las reglas de gobernanza de la IA (2.1.8).

2.2.1. Control jurídico del algoritmo

Un punto importante es mencionado en el artículo 10, apartados 1 a 5 que establece que los sistemas de IA de alto riesgo que emplean técnicas de entrenamiento de modelos deben desarrollarse a partir de conjuntos de datos de entrenamiento, validación y prueba que cumplan criterios de calidad (por ejemplo, que la muestra de datos sea suficientemente representativa). Con este artículo busca el legislador europeo de evitar que el sistema contenga sesgos o alucinaciones que puedan llegar a tener efectos sobre el disfrute de los derechos a la seguridad, los derechos fundamentales o crear discriminaciones. Por ello este artículo indica que los datos deben ser pertinentes, representativos, completos y con propiedades estadísticas adecuadas, considerando el entorno geográfico, contextual, conductual o funcional. Es además mencionado en el mismo artículo que el tratamiento excepcional de categorías especiales de datos personales solo se permite si es estrictamente necesario para detectar y corregir sesgos, bajo garantías técnicas y de seguridad, y cumpliendo condiciones estrictas de confidencialidad, acceso, eliminación y documentación.

2.2.2. Control por supervisión humana y explicabilidad (comprensibilidad) del algoritmo

Un punto importante de control es mencionado en el artículo 14 que exige que los sistemas de IA de alto riesgo sean diseñados y desarrollados para permitir una supervisión humana efectiva durante su uso, incorporando herramientas de interfaz humano-máquina adecuadas. El mismo artículo menciona que las medidas de supervisión pueden ser integradas por el responsable del despliegue. Esto quiere decir que las compañías de seguros pueden exigir a los proveedores de las plataformas de inteligencia artificial la creación o adaptación de herramientas de supervisión humana para realizar un control de los hallazgos o conclusiones de los sistemas.

El Reglamento establece además que los supervisores deben poder estar en la capacidad de identificar las limitaciones del sistema, vigilar su funcionamiento, detectar anomalías y resolver problemas, evitando el sesgo de automatización (confianza excesiva en los sistemas automatizados). El control humano debe entonces tener la posibilidad de inspeccionar el funcionamiento de los algoritmos y esto se relaciona con el tema de la explicabilidad (mejor empleado en castellano “comprensibilidad”).

La justificación de la necesidad de la intervención humana es que se debe verificar la coherencia de los resultados generados por los algoritmos, especialmente en ámbitos sensibles desde el punto de vista regulatorio y comercial, como en la información y el asesoramiento proporcionados a los clientes, así como en los temas de fraude y de prevención de lavado de activos y financiación del terrorismo (Fliche & Yans, 2018). Esta precaución es comprensible, sobre todo porque las técnicas de IA están todavía en una etapa embrionaria de perfeccionamiento de estas tecnologías. Es bien sabido que la IA puede detectar errores evidentes, gracias al aprendizaje y entrenamiento efectuado con los algoritmos, pero no es seguro que esté igualmente preparada para identificar otros sesgos menos visibles, pero potencialmente problemáticos (Fliche & Yans, 2018).

Cuando una persona interviene es más sencillo justificar el resultado arrojado por la IA sobre todo para garantizar la detección de debilidades estructurales en los datos o algoritmos utilizados, así como para evitar problemas de transparencia ante los clientes y los entes de control. La explicabilidad de la IA es fundamental para respetar el cumplimiento de los fundamentos jurídicos, técnicos y estadísticos establecidos en la materia por la ley. La explicabilidad es además un requisito que resulta de las reglas de registro y control de las profesiones reglamentadas. Así, en el asesoramiento personalizado en seguros o en la evaluación de la solvencia en créditos, el profesional debe demostrar la pertinencia de la diligencia realizada en relación con la información proporcionada por el cliente, y parte de esta explicación debe presentarse al cliente para aclarar la propuesta (Fliche & Yans, 2018), de manera que la explicación del funcionamiento de la IA deberá ser garantizada.

Por último, el artículo 26 del Reglamento establece que los responsables del despliegue de sistemas de IA de alto riesgo deben adoptar medidas técnicas y organizativas para garantizar el uso conforme a las instrucciones del proveedor, asignando la supervisión humana a personas competentes y formadas. Dicho artículo establece además que se debe asegurar que los datos de entrada sean pertinentes y representativos, vigilar el funcionamiento del sistema y, en caso de riesgos o incidentes graves, informar de inmediato al proveedor, distribuidor y autoridad de vigilancia del mercado. Por cierto, cuando se trata de los sistemas de identificación biométrica remota, el Reglamento establece una condición inmanente: las decisiones solo podrán tomarse si al menos dos personas verifican y confirman la identificación.

2.2.3. Potestades de las Autoridades Estatales de supervisión

El artículo 77 del Reglamento otorga a las autoridades nacionales responsables de proteger los derechos fundamentales la facultad de solicitar y acceder a cualquier documentación creada o conservada en virtud del Reglamento, especialmente sobre sistemas de IA de alto riesgo. La documentación debe ser proporcionada en lenguaje y formato accesibles. Si la documentación es insuficiente para determinar el cumplimiento o incumplimiento por una empresa o entidad que despliegue la IA, este artículo indica que las Autoridades pueden ordenar la realización de pruebas técnicas sobre el sistema.

2.2.4. Régimen Sancionatorio

En último lugar el Reglamento (UE) 2024/1689 establece en su artículo 99 el régimen sancionatorio aplicable a los sistemas de IA. Este artículo establece que la utilización de prácticas prohibidas de la (indicadas en el artículo 5) será sancionada con multas de hasta 35 millones de euros o el 7% del volumen de negocios mundial, aplicándose la cantidad mayor. El incumplimiento de otras disposiciones relacionadas con operadores u organismos notificados, distintas del artículo 5, estará sujeto a multas de hasta 15 millones de euros o el 3% del volumen de negocios mundial. La presentación de información inexacta, incompleta o engañosa a organismos notificados o autoridades

nacionales competentes se sancionará con multas de hasta 7,5 millones de euros o el 1% del volumen de negocios mundial. Estas sanciones son proporcionales a la gravedad de la infracción y buscan garantizar el cumplimiento del Reglamento y la protección de los derechos fundamentales.

Vale la pena mencionar que el numeral 49) del artículo 3º del Reglamento define “incidente grave” como incidente que, directa o indirectamente, cause el fallecimiento de una persona o un perjuicio grave para su salud, una alteración grave e irreversible de infraestructuras críticas, el incumplimiento de obligaciones de derechos fundamentales, o daños graves a la propiedad o al medio ambiente. El numeral 61) del mismo artículo describe que una “infracción generalizada” es todo acto u omisión contrario al Derecho de la Unión que perjudique intereses colectivos de personas en al menos dos Estados miembros distintos, o que tenga características comunes y sea cometido simultáneamente por el mismo operador en al menos tres Estados miembros. Finalmente, el numeral 65) del mismo artículo establece que es “riesgo sistémico” aquel específico de los modelos de IA de uso general con gran impacto, que puede tener repercusiones considerables en el mercado de la Unión por su alcance o efectos negativos previsibles en la salud pública, la seguridad, los derechos fundamentales o la sociedad en su conjunto, propagándose a gran escala a lo largo de la cadena de valor.

CONCLUSIONES

Los avances tecnológicos siempre han impuesto dilemas éticos que no han sido siempre resueltos. Los beneficios de la IA son inmensos y estamos a penas en la aurora de esta tecnología. Pero dicha aurora puede convertirse en pesadilla si el Estado de Derecho no es protegido con unas políticas públicas de gobernanza jurídica reforzada de dichas tecnologías disruptivas. Las Big-Tech promocionan y publicitan las ventajas de las herramientas de IA como unas tecnologías al servicio de la sociedad, que promueven la libertad y el comercio, pero su comportamiento indica lo contrario. El uso de datos empresariales de la industria literaria y periodística estadounidense ha sido el máximo ejemplo de una piratería subterfugia que rindió sus frutos durante el entrenamiento de los sistemas algorítmicos de OpenAI. Varios procedimientos judiciales o extrajudiciales fueron entonces iniciados por empresas como News Corp², Wiley y Harper Collins para obtener indemnización por cada obra cuya utilización ilícita por OpenAI fue demostrada³.

Ante la constatación de esta piratería uno entiende porque algunos sectores son reticentes a utilizar datos confidenciales en las plataformas de IA. El uso de datos personales es aún más patente. Todos hemos notado que basta con hablar de un producto para que éste sea deslizado como oferta de publicidad entre dos publicaciones de un feed de LinkedIn, Facebook o Twitter. Pero los reproches deontológicos son mucho

² “News Corp and OpenAI Sign Landmark Multi-Year Global Partnership | News Corporation.” *News Corporation*, 2024, investors.newscorp.com/news-releases/news-release-details/news-corp-and-openai-sign-landmark-multi-year-global-partnership.

³ Morozov, Evgeny. “Le Numérique Nous Ramène-t-Il Au Moyen Âge ?” *Le Monde Diplomatique*, 31 July 2025, www.monde-diplomatique.fr/2025/08/MOROZOV/68672. Accessed 17 Oct. 2025.

mas profundos. Estas plataformas pueden captar una cantidad impresionante de informaciones de cada usuario⁴: tipo de sistema y operador utilizado, geolocalización, lugares de acceso a wifi, contactos, fotografías y archivos, micrófono, cámara, historial de interacciones, datos de búsqueda y navegación, “impresión digital” (información del usuario aun cuando este no esta conectado), y muchas otras mas detalladas sobre los usos del dispositivo, informaciones bancarias, origen demográfico, edad, genero, preferencias sexuales⁵.

Vemos como la ética no es la primera preocupación de las grandes empresas tecnológicas (Google, Amazon, Facebook, Apple y Microsoft) y la ausencia de gobernanza jurídica llega hoy a niveles extremos. *TikTok* permite al algoritmo mostrar contenido pornográfico a menores de edad⁶ y consiente la difusión de peligrosos *trends* de automutilación⁷, *Instagram* manipula adolescentes para mantenerlos conectados, con problemas de autoestima y con propensión al consumismo⁸, el algoritmo de *Linkedin* busca interacciones hostiles entre los usuarios⁹, *Twitter X* es mundialmente conocido por reproducir el odio¹⁰, la mentira y el masculinismo, *Whatsapp* utiliza herramientas visuales para crear ansiedad en los usuarios, *Youtube* se ha convertido en una caverna de complotistas y antifeministas¹¹, grupos armados utilizan *Facebook* para amedrentar a la población y reclutar menores de edad¹², e incluso aplicaciones tan pueriles como *Candycrush* impactan las capacidades de discernimientos de las personas mayores de edad¹³. La lista continua y la casuística es variada en las metodologías de afectación de los derechos fundamentales de los usuarios.

⁴ Gruhier, Camille. “Données Personnelles—Comment Empêcher Facebook et Instagram d’Utiliser Vos Données Pour Son IA—Actualité.” *Quechoisir.org*, UFC-Que Choisir, 23 Apr. 2025, www.quechoisir.org/actualite-donnees-personnelles-comment-empêcher-facebook-et-instagram-d-utiliser-vos-donnees-pour-son-ia-n166168/. Accessed 17 Oct. 2025.

⁵ Commission Nationale de l’Informatique et des Libertés. “IA : Meta Entraînera Ses Systèmes d’IA Avec Les Données Des Utilisateurs Européens Dès Fin Mai 2025.” *Cnil.fr*, 2025, www.cnil.fr/fr/meta-entrainement-ia-donnees-utilisateurs. Accessed 17 Oct. 2025.

⁶ Global Witness. “TikTok’s Algorithm Directs 13-Year-Olds to Porn.” *Global Witness*, 2025, globalwitness.org/en/campaigns/digital-threats/tiktok-directs-13-year-olds-to-porn/.

⁷ Amnesty International. “Monde. Le Fil “Pour Toi” de TikTok Risque de Pousser Des Enfants et Des Jeunes Vers Du Contenu Dangereux En Lien Avec La Santé Mentale.” *Amnesty International*, 7 Nov. 2023, www.amnesty.org/fr/latest/news/2023/11/tiktok-risks-pushing-children-towards-harmful-content/.

⁸ Le Monde. “Instagram Peut Avoir Des Effets Néfastes Sur Les Adolescents, Selon Une Étude Menée Par Facebook.” *Le Monde.fr*, 14 Sept. 2021, www.lemonde.fr/pixels/article/2021/09/14/instagram-peut-avoir-des-effets-nefastes-sur-les-adolescents-selon-une-etude-menee-par-facebook_6094640_4408996.html.

⁹ Corera, Gordon. “MI5 Warns of Spies Using LinkedIn to Trick Staff into Spilling Secrets.” *BBC*, 20 Apr. 2021, www.bbc.com/news/technology-56812746.

¹⁰ Mercier, Arnaud, and Laura Amigo. “Tweets Injurieux et Haineux Contre Les Journalistes et Les “Merdias.”” *Mots*, vol. 1, no. 125, 4 Mar. 2021, pp. 73–91, <https://doi.org/10.4000/mots.28043>. Accessed 22 June 2022.

¹¹ Mari, Elsa, and Christine Mateus. ““Arrêtez de Vous Laisser Émasculer” : Sur Les Réseaux Sociaux, La Fabrique Des Masculinistes.” *Leparisien.fr*, *Le Parisien*, 22 May 2024, www.leparisien.fr/societe/arretez-de-vous-laisser-emasculer-sur-les-reseaux-sociaux-la-fabrique-des-masculinistes-22-05-2024-PPKOJ55F2ZBTROBWUGTLLPLIZM.php. Accessed 17 Oct. 2025.

¹² ONU Colombia. “Los Influencers Del Fusil.” *Noticias ONU*, 2 Aug. 2025, news.un.org/es/story/2025/08/1540278. Accessed 17 Oct. 2025.

¹³ Smith, Dana. “This Is What Candy Crush Saga Does to Your Brain | Dana Smith.” *The Guardian*, *The Guardian*, Apr. 2014, www.theguardian.com/science/blog/2014/apr/01/candy-crush-saga-app-brain.

Además de ello, las democracias occidentales han tenido consecuencias concretas en el funcionamiento de los cuerpos políticos. Las redes sociales perturbaron la decisión ciudadana en decisiones como el Brexit, la guerra de Ucrania, la reelección de Netanyahu en Israel, la utilización de armas en Estados Unidos, las políticas migratorias en Italia, las redes de prostitución en Francia, el proceso de paz en Colombia. Vale entonces la pena preguntarse hasta qué punto el Reglamento de la Unión Europea 2024/1689 que entra en vigor en agosto 2026 tendrá un impacto concreto en la protección del Estado de Derecho global. Lo cierto es que los abusos son innumerables y el legislador europeo ha dado un paso positivo en la regulación del comportamiento de las Big Tech y el uso de la IA.

BIBLIOGRAFÍA

- Agamben, G. (2014). *Qu'est-ce qu'un dispositif?* (M. Rueff, Trans.), Rivages (Poche, Vol. 1, p. 80).
- Amnesty International. "Monde. Le Fil "Pour Toi" de TikTok Risque de Pousser Des Enfants et Des Jeunes Vers Du Contenu Dangereux En Lien Avec La Santé Mentale." *Amnesty International*, 7 Nov. 2023, www.amnesty.org/fr/latest/news/2023/11/tiktok-risks-pushing-children-towards-harmful-content/.
- Commission Nationale de l'Informatique et des Libertés. "IA : Meta Entraînera Ses Systèmes d'IA Avec Les Données Des Utilisateurs Européens Dès Fin Mai 2025." *Cnil.fr*, 2025, www.cnil.fr/fr/meta-entrainement-ia-donnees-utilisateurs. Accessed 17 Oct. 2025.
- Corera, Gordon. "MI5 Warns of Spies Using LinkedIn to Trick Staff into Spilling Secrets." *BBC*, 20 Apr. 2021, www.bbc.com/news/technology-56812746.
- Fliche, O., & Yans, S. (2018). *Intelligence artificielle : enjeux pour le secteur financier*. ACPR, Banque de France.
- Foucault, M. (2025). *Théories et institutions pénales: cours au Collège de France. 1971-1972*. Seuil.
- Global Witness. "TikTok's Algorithm Directs 13-Year-Olds to Porn." *Global Witness*, 2025, globalwitness.org/en/campaigns/digital-threats/tiktok-directs-13-year-olds-to-porn/.
- Gruhier, Camille. "Données Personnelles—Comment Empêcher Facebook et Instagram d'Utiliser Vos Données Pour Son IA—Actualité." *Quechoisir.org*, UFC-Que Choisir, 23 Apr. 2025, www.quechoisir.org/actualite-donnees-personnelles-comment-empecher-facebook-et-instagram-d-utiliser-vos-donnees-pour-son-ia-n166168/. Accessed 17 Oct. 2025.
- Hufeld, F. (2022). *Big data meets artificial intelligence : Challenges and implications for the supervision and regulation of financial services*. BaFin, Federal Financial Supervisory Authority.

- Jean, A. (2020). A brief history of artificial intelligence. *Medecine Sciences: M/S*, 36(11), 1059–1067.
- Kivisaari, E., Wilhelmy, L., & Ecija Serrano, P. (2021). *Artificial Intelligence governance principles, towards ethical and trustworthy Artificial Intelligence in the European insurance sector : a report from EIOPA's Consultative Expert Group on Digital Ethics in insurance*. Publications Office of the European Union.
- Le Monde. “Instagram Peut Avoir Des Effets Néfastes Sur Les Adolescents, Selon Une Étude Menée Par Facebook.” *Le Monde.fr*, 14 Sept. 2021, www.lemonde.fr/pixels/article/2021/09/14/instagram-peut-avoir-des-effets-nefastes-sur-les-adolescents-selon-une-etude-menee-par-facebook_6094640_4408996.html.
- Lewis, P. (2025). Big Tech deploys Orwellian doublespeak to mask its democratic corrosion. *The Guardian*.
- Mari, Elsa, and Christine Mateus. ““Arrêtez de Vous Laisser Émasculer” : Sur Les Réseaux Sociaux, La Fabrique Des Masculinistes.” *Leparisien.fr*, Le Parisien, 22 May 2024, www.leparisien.fr/societe/arretez-de-vous-laisser-emasculer-sur-les-reseaux-sociaux-la-fabrique-des-masculinistes-22-05-2024-PPKOJ55F2ZB-TROBWUGTLLPLIZM.php. Accessed 17 Oct. 2025.
- Mercier, Arnaud, and Laura Amigo. “Tweets Injurieux et Haineux Contre Les Journalistes et Les “Merdias.”” *Mots*, vol. 1, no. 125, 4 Mar. 2021, pp. 73–91, <https://doi.org/10.4000/mots.28043>. Accessed 22 June 2022.
- Morozov, E. (2025, August 1). ¿Nos está devolviendo la tecnología digital a la Edad Media? *Le Monde Diplomatique, En Espanol*.
- Morozov, Evgeny. “Le Numérique Nous Ramène-t-Il Au Moyen Âge ?” *Le Monde Diplomatique*, 31 July 2025, www.monde-diplomatique.fr/2025/08/MOROV/68672. Accessed 17 Oct. 2025.
- “News Corp and OpenAI Sign Landmark Multi-Year Global Partnership | News Corporation.” *News Corporation*, 2024, investors.newscorp.com/news-releases/news-release-details/news-corp-and-openai-sign-landmark-multi-year-global-partnership.
- ONU Colombia. “Los Influencers Del Fusil.” Noticias ONU, 2 Aug. 2025, news.un.org/es/story/2025/08/1540278. Accessed 17 Oct. 2025.
- Pravina Ladva, Group Chief Digital & Technology Officer & Antonio Grasso,. “The Future of AI in Insurance.” <https://www.swissre.com/Risk-Knowledge/Advancing-Societal-Benefits-Digitalisation/Future-Ai-Insurance.html>, Swiss Re Group, 12 July 2023, www.swissre.com/risk-knowledge/advancing-societal-benefits-digitalisation/future-ai-insurance.html.

- Rosa, H. (2010). *Alienation and acceleration: Towards a critical theory of late-modern temporality* (Vol. 3). Aarhus University Press.
- Smith, Dana. "This Is What Candy Crush Saga Does to Your Brain | Dana Smith." *The Guardian*, The Guardian, Apr. 2014, www.theguardian.com/science/blog/2014/apr/01/candy-crush-saga-app-brain.
- Vivant, M., & Warusfel, B. (2024). *Le Lamy droit du numérique* (M. Vivant, Ed.; 1st ed.). Lamy Liaisons.