

Inteligencia Artificial y Big Data

Impacto en el sector asegurador

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

Anexos

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

Anexos

Introducción

Transformación tecnológica en el seguro (I)

- ❑ Las dos últimas décadas han sido, sin duda, las de mayor **transformación** para **todos los sectores industriales**. El sector asegurador no ha sido una excepción.
- ❑ Las empresas han tenido que realizar **tremendos esfuerzos** para **adaptar** sus **propuestas de valor** a las nuevas exigencias de sus clientes, aprovechando las capacidades que ofrece la tecnología.
- ❑ Todos esos ejemplos tienen algo en común: **estar donde está el cliente (omni-canalidad)** en los múltiples puntos de contacto con la empresa que le presta el servicio.
- ❑ Los **planes de digitalización** que han lanzado, sin excepción, todas las compañías aseguradoras han tratado de lograr los siguientes **objetivos**, entre otros:
 - **Mejora de experiencia del cliente.**
 - **Potenciar y optimizar la gestión de venta online:** Modelos de sensibilidad al precio, comparadores o marketing digital, entre otros.
 - **Mejorar el servicio al cliente con la ayuda de la Inteligencia Artificial:** Asistentes virtuales, asesoramiento proactivo, gestión inteligente de siniestros, etc.
 - **Uso eficiente de los datos (entidades *data driven*):** Empleo eficiente de los datos para **extraer su valor** y **mejorar la toma de decisiones** estratégicas utilizando herramientas como el **Big Data**.
 - **Personalizar los productos (seguros *peer to peer*):** Pólizas basadas en el uso que se haga del servicio, comportamiento del cliente, hábitos, consumo, salud o actividad.
 - **Simplificar los procesos:** RPA (*Robotics Process Automation*), procesamiento automático de datos (estructurados y no estructurados), *chatbots*.

Introducción

Transformación tecnológica en el seguro (II)

❑ **Tecnologías potencialmente utilizables por el sector asegurador** para lograr los anteriores objetivos:

	Beneficios	Riesgos
Big Data	- Almacenamiento de grandes cantidades de datos. - Información precisa y oportuna.	- Privacidad de datos. - Uso ilícito de información.
Blockchain	- Mecanismo automatizado de confianza. - Múltiples validadores. - Contratos inteligentes.	- Falta de regulación legal. - Anonimato (blanqueo de capitales y financiación del terrorismo). - Traslado del lenguaje legal al mundo digital.
Cloud	- Fácil acceso. - Consistencia y centralización de recursos. - Fiabilidad de los recursos. - Reducción de costes.	- Dependencia del acceso a internet. - Dependencia del proveedor del servicio. - Ciber-ataques. - Privacidad de datos del cliente.
IA	- Optimiza la suscripción y gestión de riesgos. - Reducción de costes de transacción. - Ventaja competitiva.	- Privacidad de datos. - Ciber-ataques. - Uso ilícito de información.
IoT	- Acceso a datos del comportamiento de los cliente. - Mayor personalización de los productos.	- Privacidad de datos. - Uso ilícito de la información.

Introducción

¿Qué es Big Data?

❑ Big Data a través de las 3 Vs:

- **Volumen (Alto):** Se refiere al tamaño de los datos. En el caso del Big Data podemos hablar de terabytes e incluso petabytes (1 millón de gigabytes).
- **Variedad (Alta):** Heterogeneidad de los datos. Los datos estructurados representan el 15-20% del total de los datos disponibles hoy en día. Texto, imágenes, audio, video, RRSS son algunos ejemplos de datos no estructurados.
- **Velocidad (Alta):** Hace referencia a la tasa de generación de datos, así como el tiempo que tarda en analizar los datos. La proliferación de dispositivos digitales (smartphones, tablets, etc) se ha traducido en una ingente generación de datos.

❑ En 2016, la FCA (*Financial Conduct Authority*) definió el Big Data como:

- El uso de nuevos y ampliados conjuntos de datos, incluidas fuentes de datos no convencionales como las RRSS.
- Adopción de nuevas tecnologías para generar, recopilar y guardar esas nuevas formas de datos (estructurados y no estructurados).
- Uso de técnicas avanzadas de procesamiento de datos.
- Técnicas analíticas sofisticadas como el Deep Learning.
- Aplicación del conocimiento generado en la toma de decisiones y actividades de las empresas.

Introducción

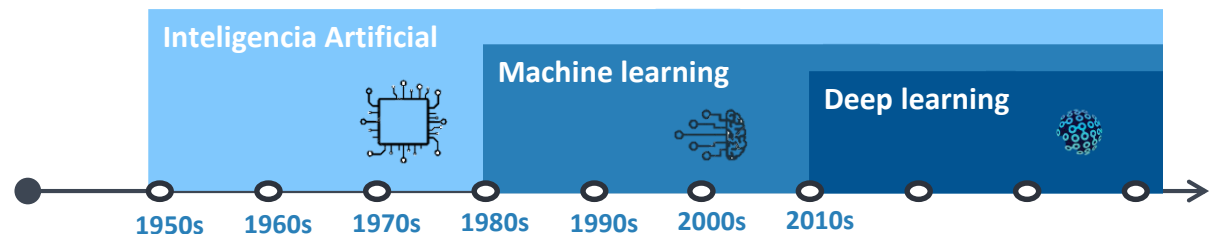
¿Qué es la Inteligencia Artificial (IA)?

❑ La IA comprende principalmente:

- ❑ **Big Data Analytics (BDA) o analítica de macro-datos:** Tecnología utilizada para analizar una enorme cantidad de datos estructurados y no estructurados que son transformados en información útil para la toma de decisiones, conocer tendencias de mercado y comportamientos.
- ❑ **Machine Learning (ML) o aprendizaje automático:** Una de las principales formas en que se aplica la IA con algoritmos que pueden aprender de ejemplos y pueden mejorar su rendimiento a lo largo del tiempo.
- ❑ **Deep Learning (DL):** Rama del aprendizaje automático que se basa en algoritmos y modelos estadísticos complejos con múltiples capas de procesamiento paralelo que tratan de modelar la forma en que funciona el cerebro humano.

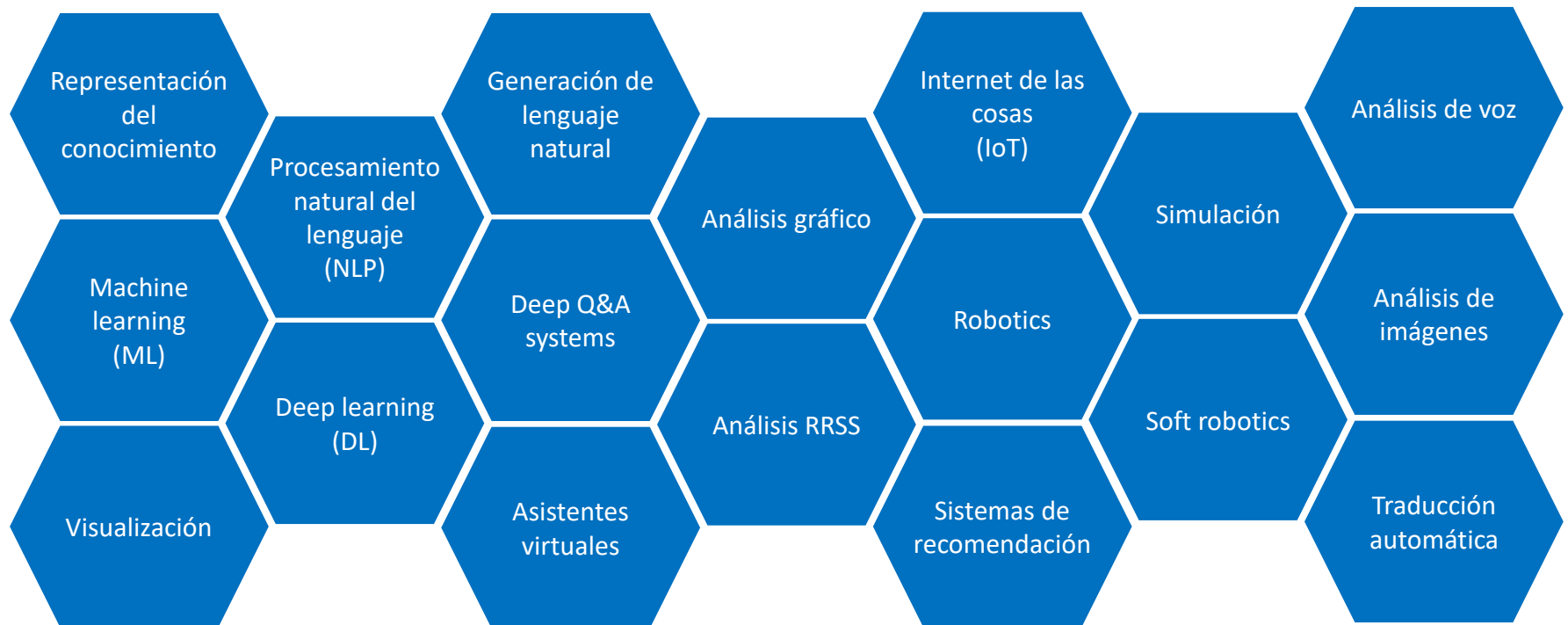
❑ La IA permite:

- ❑ **Identificar patrones** que no serían posibles de identificar por una persona.
- ❑ **Desarrollar nuevos modelos de negocio** que se ajusten a las demandas cambiantes de la sociedad.
- ❑ **Automatizar ciertos procesos básicos y rudimentarios** (centrando los recursos en tareas de mayor valor añadido para el cliente).



Introducción

Áreas de desarrollo de la IA

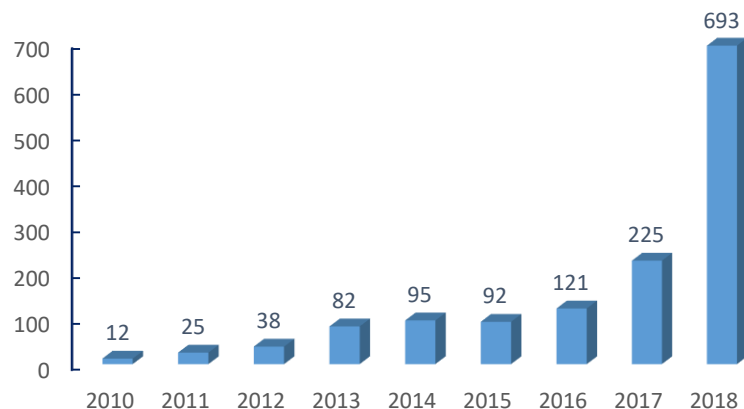


Introducción

La Inteligencia Artificial en números

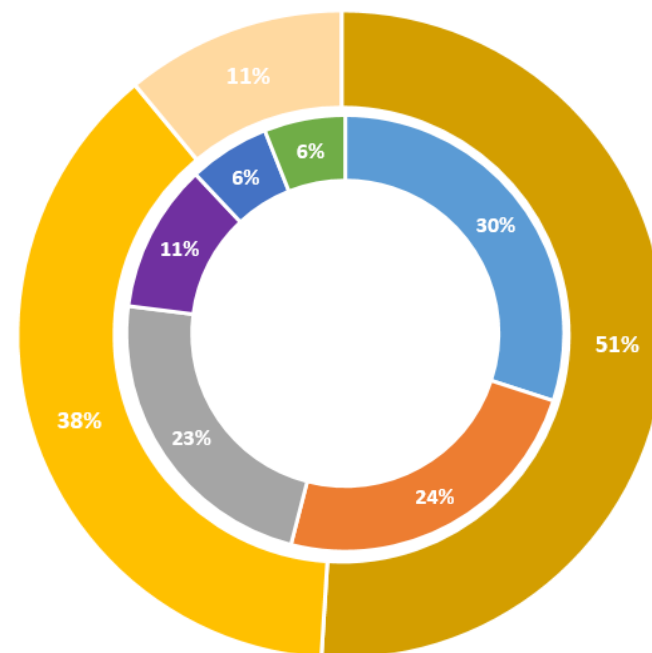
Solicitudes de patente

(Fuente: BD Patentes de Google, Swiss Re)



Distribución de patentes

(Fuente: BD Patentes de Google, Swiss Re)



Solicitudes de patentes de IA*

1.863

EEUU

1.085

China

1.074

UE



400%
incremento de las solicitudes de patentes de IA publicadas en la última década

IA como parte de la transición verde y digital



60 millones de nuevos puestos de trabajo

podrían crearse en el mundo para 2025 gracias a la IA y a la robótica

La UE ayuda a generar beneficios para las personas y las empresas al financiar proyectos de IA a través de la Agenda de Capacidades, el Fondo de Transición Justa, el programa Horizonte, etc.

*en el período 1960-2019

Fuentes: Comisión Europea (2019), IPOL (2020)



europarl.eu

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

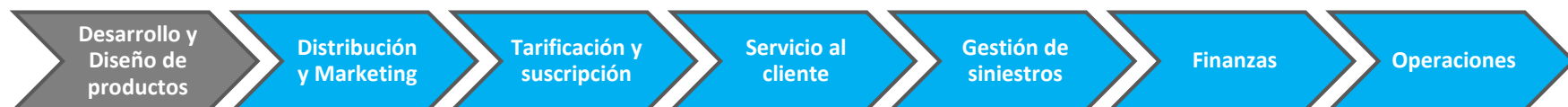
5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

Anexos

La Inteligencia Artificial en la cadena de valor del seguro

Casos de uso de la Inteligencia Artificial (I)

❑ Ejemplos de aplicaciones de IA en la cadena de valor de las entidades aseguradoras:



Desarrollo y Diseño de productos

Análisis de textos

Estructurar las cláusulas de las pólizas existentes para desarrollar productos más rápida y eficientemente

Detección de patrones y anomalías

Identificar patrones basados en el análisis del feedback de los clientes para diseñar nuevos productos

Recomendación

Identificar hábitos de consumo de los clientes para mejorar la oferta de productos ofrecidos al cliente

Soluciones conversacionales

Uso de los datos tomados de las soluciones que interactúan con los clientes (asistentes virtuales o *chatbots*) para mejorar los productos

Reconocimiento de voz

Identificación de los “puntos de dolor” del cliente con los productos contratados a través de la analítica de sentimientos

Análisis de video e imagen / Reconocimiento facial

Análisis de emociones y reacciones de los clientes basados en videos de los mismos para mejorar los productos o diseñar nuevos

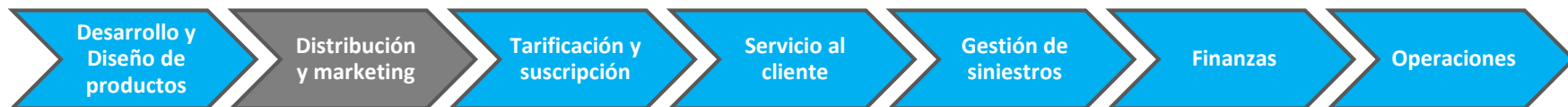
Generación de lenguaje natural

Generación de nuevas cláusulas de pólizas basadas en experiencias previas con el fin de mejorar la eficiencia operativa

La Inteligencia Artificial en la cadena de valor del seguro

Casos de uso de la Inteligencia Artificial (II)

❑ Ejemplos de aplicaciones de IA en la cadena de valor de las entidades aseguradoras:



Distribución y marketing

Análisis de textos

Análisis de las llamadas y los post en RRSS de los clientes para desarrollar nuevas campañas de marketing

Detección de patrones y anomalías

Segmentación de clientes en grupos específicos

Recomendación

Proporcionar nuevas oportunidades de ventas para ganar nuevos clientes

Soluciones conversacionales

Uso de asistentes virtuales en el proceso de distribución para incrementar las ventas

Reconocimiento de voz

Análisis de la voz de los clientes, incluidas emociones/sentimientos, para incrementar la personalización de la oferta al cliente

Análisis de video e imagen / Reconocimiento facial

Creación y mejora de campañas de marketing basadas en el análisis de videos públicos

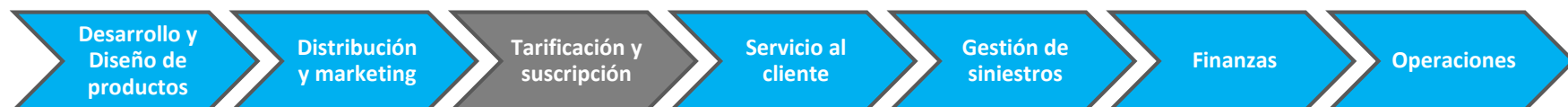
Generación de lenguaje natural

Diseño de campañas de marketing con diferentes estilos (*A/B testing*)

La Inteligencia Artificial en la cadena de valor del seguro

Casos de uso de la Inteligencia Artificial (III)

❑ Potenciales aplicaciones de IA en la cadena de valor de las entidades aseguradoras:



Tarificación y suscripción

Análisis de textos

Identificación de ambigüedades en los datos de suscripción para detectar fraude

Detección de patrones y anomalías

Personalización de primas en base a la evaluación de riesgos históricos

Recomendación

Segmentación de clientes por riesgo

Soluciones conversacionales

Uso de *chatbots* y análisis de sus conversaciones con los clientes para obtener características específicas de los mismos

Reconocimiento de voz

Detección de fraude basado en el análisis de la voz

Análisis de video e imagen / Reconocimiento facial

Análisis de video/imagen para mejorar la personalización de la oferta de la póliza

Generación de lenguaje natural

Generación de informes basados en datos históricos para mejorar la eficiencia operativa

La Inteligencia Artificial en la cadena de valor del seguro

Casos de uso de la Inteligencia Artificial (IV)

❑ Ejemplos de aplicaciones de IA en la cadena de valor de las entidades aseguradoras:

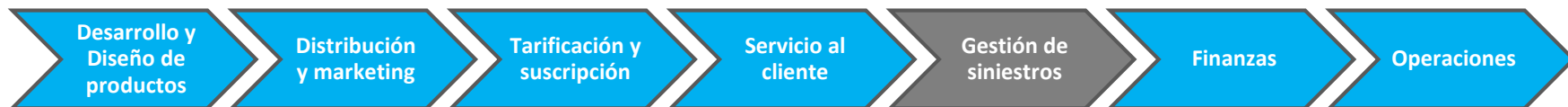


Servicio al cliente	
Análisis de textos	Generación de informes basados en solicitudes/consultas de los clientes para reducir tiempo y costes
Detección de patrones y anomalías	Proveer respuestas automáticas a las consultas de los clientes basados en cuestiones y respuestas históricas
Recomendación	Resolución de incidencias basadas en el análisis de datos históricos para reducir tiempo
Soluciones conversacionales	Integración de asistentes virtuales en las plataformas de servicio al cliente para mejorar la experiencia del cliente
Reconocimiento de voz	Mejora de la comunicación con el cliente gracias al análisis de emociones/sentimientos para mejorar el servicio al cliente
Análisis de video e imagen / Reconocimiento facial	Análisis de emociones/sentimientos durante las video-llamadas para mejorar el servicio prestado
Generación de lenguaje natural	Generación automática de emails y comunicaciones al cliente para optimizar la carga de trabajo de los empleados

La Inteligencia Artificial en la cadena de valor del seguro

Casos de uso de la Inteligencia Artificial (V)

❑ Ejemplos de aplicaciones de IA en la cadena de valor de las entidades aseguradoras:

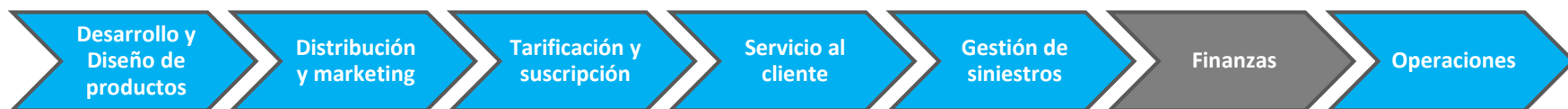


Gestión de siniestros	
Análisis de textos	Generación de datos a partir de la información reportada de siniestros para reducir su tiempo de gestión
Detección de patrones y anomalías	Verificación de siniestros usando fuentes de datos externas para ahorrar costes
Recomendación	Generación de recomendaciones basadas en datos de siniestros históricos
Soluciones conversacionales	Uso de <i>chatbots</i> en las relaciones con el cliente sobre un siniestro
Reconocimiento de voz	Generación automática de texto en base a siniestros comunicados oralmente y análisis de las emociones de los clientes
Análisis de video e imagen / Reconocimiento facial	Determinación de la severidad del siniestro en base al análisis de fotos tomadas por los clientes/peritos para mejorar el tiempo de respuesta
Generación de lenguaje natural	Generación automática de informes sobre siniestros para reducir la carga de trabajo de los empleados

La Inteligencia Artificial en la cadena de valor del seguro

Casos de uso de la Inteligencia Artificial (VI)

❑ Ejemplos de aplicaciones de IA en la cadena de valor de las entidades aseguradoras:

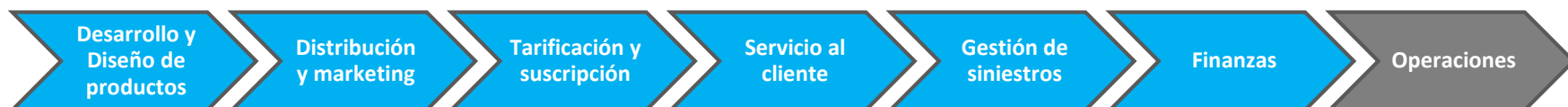


Finanzas	
Análisis de textos	Detección de oportunidades de inversión en base al análisis de noticias y de las RRSS, reduciendo el tiempo de reacción
Detección de patrones y anomalías	Generación de recomendaciones automáticas vinculadas a la gestión del portfolio de inversiones en base al análisis de anomalías del mercado
Recomendación	Generación de recomendaciones automáticas vinculadas a la gestión del portfolio de inversiones en base al análisis del mercado
Soluciones conversacionales	Uso de asistentes virtuales por los empleados para analizar datos financieros y tomar decisiones
Reconocimiento de voz	Análisis de llamadas de grandes inversores para identificar anticipadamente potenciales oportunidades de inversión
Análisis de video e imagen / Reconocimiento facial	Predicción de movimientos en los mercados basada en el análisis de noticias públicas
Generación de lenguaje natural	Generación de informes financieros automáticos

La Inteligencia Artificial en la cadena de valor del seguro

Casos de uso de la Inteligencia Artificial (VII)

❑ Ejemplos de aplicaciones de IA en la cadena de valor de las entidades aseguradoras:



Operaciones	
Análisis de textos	Transformación de texto en datos estructurados para su posterior análisis
Detección de patrones y anomalías	Predicción de la carga de trabajo basada en datos históricos para prevenir cuellos de botella
Recomendación	Recomendación de candidatos en base al análisis de RRSS, mejorando el proceso de reclutamiento de personal
Soluciones conversacionales	Uso de asistentes virtuales para la planificación de reuniones o la reserva de salas de reunión
Reconocimiento de voz	Generación automática de informes y actas de reuniones
Análisis de video e imagen / Reconocimiento facial	Aviso a los empleados de incidencias relacionadas con el sitio de trabajo en base al análisis de las imágenes de las cámaras de seguridad
Generación de lenguaje natural	Generación de informes internos y entradas en el sitio web de la entidad para aumentar la transparencia

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

Anexos

Aplicación de la Inteligencia Artificial en el seguro

Beneficios y ética (I)

	Proceso de negocio	Sistema IA	Beneficios (Entidad)	Beneficios (Consumidor)	Ética
Diseño y desarrollo del producto	Investigación de mercado	Histórico de datos de clientes y encuestas como insumo en el desarrollo de nuevos productos	Innovación para incrementar las ventas y la retención de clientes	Acceso a nuevos productos y más competencia	• Uso de datos de diferentes fuentes donde el procesamiento no se ha explicado claramente al cliente.
	Desarrollo de producto	Nuevos productos basados en el uso y comportamiento de los clientes	Mayor crecimiento y acceso a nuevos mercados	Productos de mayor valor, más personalizados	• La información de los consumidores no digitales no se tiene en cuenta. • Algunos datos de comportamiento pueden estar correlacionados con características protegidas.
	Identificación digital	• Identificación facial. • Lectura automatizada de documentos oficiales (DNI, pasaporte).	• Mayor satisfacción de los clientes (no necesario el uso de claves de acceso). • Reducción de incidencias relacionadas con el acceso.	Autenticación rápida y sencilla más	Necesario el acceso a herramientas digitales (internet, smartphones, tablets, etc.)

Aplicación de la Inteligencia Artificial en el seguro

Beneficios y ética (II)

	Proceso de negocio	Sistema IA	Beneficios (Entidad)	Beneficios (Consumidor)	Ética
Suscripción y tarificación	Evaluación del riesgo	Nuevas formas de macro-datos (datos de comportamiento, IoT, telemática, GPS)	Nuevas fuentes de información para evaluar los riesgos y tarificar con mayor detalle y precisión	La segmentación de clientes puede beneficiar a aquellos con menores riesgos individuales que la media, pero puede identificar subgrupos de mayor riesgo	- Sesgo en los datos de entrenamiento* de los modelos de IA. - Uso de datos no vinculados al riesgo. - Riesgo de exclusión de pequeños grupos de alto riesgo. - Problemas de privacidad. - Problemas de ética relacionados con la venta de datos personales.
	Optimización del precio	- Personalización de tarifas basadas en comportamiento. - Análisis de la competencia.	- Aumento de la rentabilidad. - Mayores tasas de retención.	- Precios justos en base a la pertenencia a un determinado segmento. - Precios dinámicos basados en el comportamiento del asegurado.	- Baja explicabilidad de los modelos de "caja negra"**. - Uso de datos de comportamiento no vinculados al riesgo. - Impacto en consumidores vulnerables.

Aplicación de la Inteligencia Artificial en el seguro

Beneficios y ética (III)

	Proceso de negocio	Sistema IA	Beneficios (Entidad)	Beneficios (Consumidor)	Ética
Distribución	Ventas y marketing	Técnicas de marketing digital basadas en el análisis del comportamiento del cliente en la red	Marketing digital adaptado a captar la atención de clientes y aumentar las ventas online	- Facilidad de uso. - Menores tiempos en la búsqueda y contratación de seguros.	- Exclusión de grupos vulnerables sin acceso a la red. - Suscripción innecesaria de seguros por falta de conocimiento digital.
	Nuevos clientes	Asistentes virtuales y chatbots que usan NLP	- Facilitar la experiencia del cliente. - Reducción de costes. - Automatización del proceso KYC (<i>Know Your Client</i>)	Soporte online inmediato con disponibilidad 24x7	- Recomendación no adecuada. - Riesgo de selección sesgada de clientes.
	Retención de clientes	Modelos de caída	Defensa de la cuota de mercado y predicción de caídas	Mejores precios para los consumidores dispuestos a negociar condiciones	Explotación de características de los clientes no vinculadas al riesgo para reducir la competencia
	CRM (Customer Relationship Management)	Gestión proactiva del cliente y venta cruzada	Mejor servicio al cliente en todos los canales, ofertas diferenciadas y mayores ingresos por cliente	Comunicación eficaz y facilidad de acceso a los servicios	- Reducción de la propensión de los clientes a buscar mejores ofertas. - Asesoramiento sesgado.
	Análisis del comportamiento	Análisis del comportamiento en salud y auto	Mayor conocimiento de los riesgos y mitigación del riesgo mejorando la conducta de los asegurados	Precios competitivos y recompensas por comportamientos de menor riesgo	- Comportamiento forzado de los asegurados (hacer ejercicio estando enfermo). - Seguimiento excesivo de la vida de los clientes (problemas de privacidad).

Aplicación de la Inteligencia Artificial en el seguro

Beneficios y ética (IV)

	Proceso de negocio	Sistema IA	Beneficios (Entidad)	Beneficios (Consumidor)	Ética
Gestión de siniestros	Provisiones	Uso de algoritmos de ML para estimar provisiones, en particular, en los siniestros de alta frecuencia	Predicciones de siniestralidad más precisas	Mejora de las estimaciones de liquidación de siniestros	
	Gestión del siniestro	Reconocimiento de imágenes para estimar costes de reparación en hogar o auto	- Reducción de costes de expertos. - Reducción de tiempos.	- Reducción de tiempos. - Mejora de la calidad de la reparación del siniestro.	Los sistemas automatizados son más opacos para el asegurado
	Detección de fraude	Detección de anomalías y búsqueda de patrones de fraude	Reducción de pagos por siniestros fraudulentos	Reducción de primas como consecuencia de la reducción del fraude	Falsa acusación como consecuencia del sesgo algorítmico
	Liquidación automática	Sistemas de pago integrados en la gestión de siniestros	Reducción de costes administrativos	Liquidaciones más rápidas	

Aplicación de la Inteligencia Artificial en el seguro

Beneficios y ética (V)

	Proceso de negocio	Sistema IA	Beneficios (Entidad)	Beneficios (Consumidor)	Ética
Servicio al cliente	Gestión de pólizas	Robotic Process Automation (RPA)	- Mejora del rendimiento de los sistemas <i>legacy</i>. - Reducción de costes administrativos para actividades rutinarias y repetitivas.	- Reducción de tarifas. - Mejora del servicio.	Podrían producirse sesgos en el procesamiento de pólizas
	Autoservicio	NLP, reconocimiento de voz, asistentes virtuales	- Reducción del coste del servicio al cliente. - Mejora de la comunicación con el cliente.	Servicios de soporte y asesoramiento al cliente disponibles 24x7	- Discriminación a los clientes con escasos conocimientos tecnológicos. - Asesoramiento inadecuado.
	Reclamaciones	Gestión automatizada de reclamaciones y quejas de clientes	Reducción de costes administrativos	Reducción de tiempos	Riesgo de sesgo en la resolución automática de quejas o reclamaciones en detrimento de determinados clientes

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

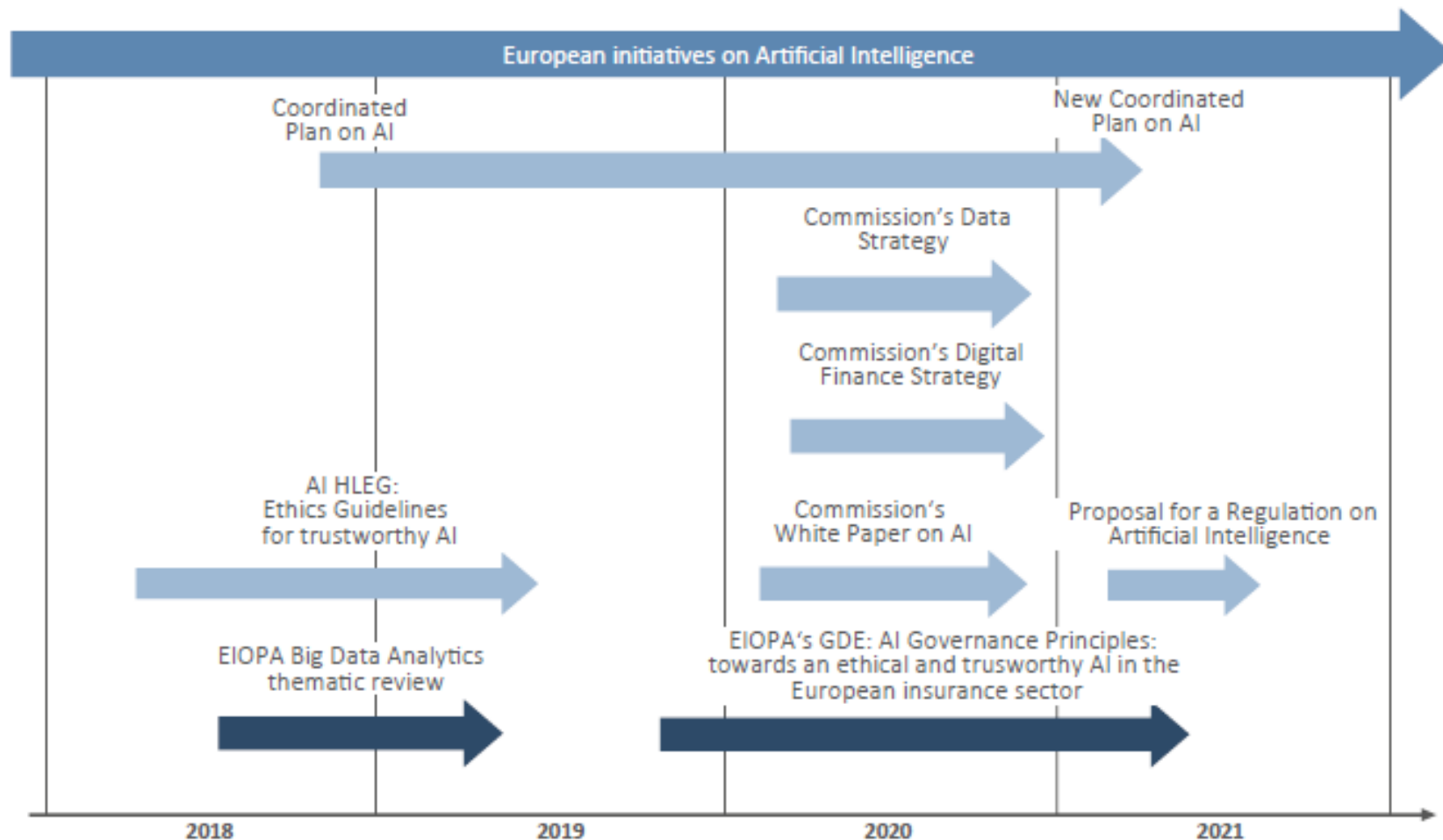
5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

Anexos

Iniciativas reguladoras de la Inteligencia Artificial

Principales Iniciativas en la UE (I)

- ❑ Principales iniciativas en la UE sobre la IA durante los últimos años:



Source: EIOPA Consultative Expert Group on Digital Ethics in insurance

Iniciativas reguladoras de la Inteligencia Artificial

Principales Iniciativas en la UE (II)

❑ Plan Coordinado sobre Inteligencia Artificial (Revisión de 2021):

- En su estrategia sobre la IA para Europa, la Comisión Europea propuso trabajar con los Estados miembros en un plan coordinado sobre IA para finales de 2018, con el objetivo de:
 - Maximizar las inversiones a nivel nacional y de la UE.
 - Fomentar las sinergias y la cooperación en toda la UE.
 - Intercambiar las mejores prácticas.
 - Definir colectivamente el camino a seguir para garantizar que la UE en su conjunto pueda competir globalmente.
- La actualización de 2021 ha ampliado y estabilizado su perspectiva propugnando inversiones de 20.000 millones al año. El foco de atención gira ahora en torno a un marco regulado que sigue teniendo como objetivo permitir la IA, pero que examina de cerca los riesgos de las aplicaciones de la IA desde la perspectiva de:
 - La seguridad.
 - La ética.
 - El uso público.
 - La innovación.
 - El clima y el medio ambiente.

Iniciativas reguladoras de la Inteligencia Artificial

Principales Iniciativas en la UE (III)

❑ Estrategia Europea de Datos (I):

- La **estrategia europea de datos** busca convertir la UE en líder de una **sociedad impulsada por los datos**. La creación de un **mercado único de datos** permitirá que estos fluyan libremente por la UE y entre sectores, en beneficio de las empresas, los investigadores y las administraciones públicas.
- La Comisión Europea propone una **nueva forma de gobernanza** de los datos para **facilitar su intercambio**, generando riqueza, proporcionando control a los ciudadanos y confianza a las empresas (Propuesta de Reglamento relativo a la Gobernanza Europea de Datos de noviembre de 2020).
- La UE creará un **mercado único de datos** en el que:
 - Los datos podrán circular por toda la UE y entre sectores, en beneficio de todos.
 - Se **respeten** plenamente las **normas europeas**, en particular en materia de **privacidad y protección de datos**, así como la legislación sobre competencia
 - Las **normas** para el **acceso** a los **datos** y su utilización sean **justas, prácticas y claras**.
- Para transformarse en una **economía de los datos atractiva, segura y dinámica**, la UE:
 - Establecerá **normas claras y justas** sobre el acceso a los datos y su reutilización.
 - **Invertirá en herramientas e infraestructuras** de próxima generación para almacenar y tratar los datos.
 - Creará **espacios de datos comunes e interoperables**.
 - Dotará a los usuarios de **derechos, herramientas y capacidades** para mantener el **pleno control** de sus datos.

Iniciativas reguladoras de la Inteligencia Artificial

Principales Iniciativas en la UE (IV)

❑ Estrategia Europea de Datos (II):

Cifras previstas para 2025



530 % –
incremento del
volumen global de
datos

De 33 zetabytes en
2018 a 175
zetabytes



**829.000
millones
de euros**

—
valor de la
economía de los
datos en la EU27

Frente a 301.000
millones de euros
(2,4 % del PIB de
la UE) en 2018



**10,9
millones
de**

profesionales de
los datos en la
EU27

Frente a 5,7
millones en 2018



65 % –
porcentaje de
población de la
UE con
competencias
digitales básicas

Frente al 57 % en
2018

Iniciativas reguladoras de la Inteligencia Artificial

Principales Iniciativas en la UE (V)

❑ Paquete de medidas sobre Finanzas Digitales (I):

- **Estrategia de Finanza Digitales:** Sus objetivos se centran:
 - **Hacer que los servicios financieros europeos sean más favorables a la digitalización y estimular la innovación responsable** y la competencia entre los proveedores de servicios financieros de la UE.
 - **Reducir la fragmentación del mercado único digital**, de modo que los consumidores puedan tener acceso a productos financieros transfronterizos y las empresas emergentes de tecnología financiera se expandan y crezcan.
 - **Garantizar que las normas de la UE en materia de servicios financieros se adapten a la era digital**, a aplicaciones tales como la IA y *blockchain*.
- **Estrategia de Pagos Minoritas:** Tiene como objetivos:
 - Ofrecer servicios de pago seguros, rápidos y fiables a las empresas y los ciudadanos europeos.
 - Facilitar a los consumidores el pago en las tiendas.
 - Hacer que las transacciones de comercio electrónico sean seguras y prácticas.

Iniciativas reguladoras de la Inteligencia Artificial

Principales Iniciativas en la UE (VI)

❑ Paquete de medidas sobre Finanzas Digitales (II):

- **Cripto-activos:**

- Por primera vez, la Comisión ha propuesto una nueva legislación en materia de cripto-activos (una representación digital de valores o derechos que pueden almacenarse y negociarse electrónicamente), con las que espera impulsar la innovación, preservando al mismo tiempo la estabilidad financiera y protegiendo a los inversores de los riesgos.
- El reglamento en el que se está trabajando permitirá a los operadores autorizados en un Estado miembro prestar sus servicios en toda la UE mediante un “régimen de pasaporte”.
- Las salvaguardias incluyen requisitos de capital, custodia de activos, un procedimiento obligatorio de reclamación a disposición de los inversores y derechos del inversor frente al emisor. Los emisores de cripto-activos respaldados por activos significativos (las denominadas “cripto-monedas estables” de nivel mundial) estarían sujetos a requisitos más estrictos (por ejemplo, en términos de capital, derechos de los inversores y supervisión).

- ***Digital Operational Resilience Act (DORA):***

- Las empresas tecnológicas tienen cada vez más peso en el ámbito financiero, tanto como proveedores de IT a las empresas financieras como en calidad de proveedores de servicios financieros.
- El Reglamento de Resiliencia Operativa Digital aspira a garantizar que todos los participantes en el sistema financiero cuenten con las garantías necesarias para paliar ciberataques y otros riesgos.
- El Reglamento exigirá a todas las empresas que se aseguren de que pueden resistir a cualquier tipo de perturbaciones y amenazas relacionadas con las tecnologías de la información y la comunicación (TIC).

Iniciativas reguladoras de la Inteligencia Artificial

Principales Iniciativas en la UE (VII)

❑ Hacia un régimen europeo de control de la IA (I):

- La Comisión Europea ha propuesto una ambiciosa legislación a nivel de la UE para defender los valores europeos en una tecnología, o grupo de tecnologías, la Inteligencia Artificial (IA).
- Enfocada, sobre todo, al control de los riesgos que plantea la IA, para la seguridad de consumidores y usuarios y del respeto a los derechos fundamentales, y poder sacar más **provecho con confianza (objetivo principal)** de las ventajas que aporta.
- La Comisión plantea algunos **riesgos inadmisibles**: La IA que contradiga los valores de la UE será prohibida.
- En el mundo hay a grandes rasgos 5 regímenes de gobernanza de la digitalización, que incluye a la IA:
 - **El estadounidense o “capitalismo de vigilancia”**: Un pequeño número de proveedores de servicios digitales disfruta de un poder de mercado preponderante y de ventajas informativas digitales, especialmente de control, sobre los datos de los usuarios y sobre la publicidad.
 - **El chino o “Estado de vigilancia”**: El Estado goza de ventajas informativas preponderantes, apoyado por un pequeño número de proveedores de servicios digitales. El Estado centraliza los datos porque son un instrumento de control político.
 - **El europeo o “humanista”**: Regulación de la UE que se centra en la protección de los valores y derechos de los usuarios digitales.
 - El de un conjunto de democracias (UK, Australia y Japón) preocupadas por el poder de las grandes plataformas estadounidenses y chinas, pero que carecen del poder de negociación de la UE.
 - El del resto del mundo, compuesto principalmente por economías emergentes y poco poder de mercado.

Iniciativas reguladoras de la Inteligencia Artificial

Principales Iniciativas en la UE (VIII)

❑ Hacia un régimen europeo de control de la IA (II):

- **Directrices sobre la IA de Abril de 2019:**

- Una IA confiable debe ser **legal** (respetar todas las leyes y regulaciones aplicables), **ética** (respetar los principios y valores éticos), y **robusta** (tanto desde una perspectiva técnica como teniendo en cuenta su entorno social).
- Presentaban un conjunto de **6 requisitos clave** que los sistemas de IA deben cumplir para ser considerados confiables: (1) agencia y supervisión humana; (2) robustez técnica y seguridad; (3) privacidad y gobernanza de datos; (4) diversidad, no discriminación y equidad; (5) bienestar social y ambiental; y (6) auditabilidad (la IA y sus resultados deben rendir cuentas ante auditores externos e internos).

- Las Conclusiones sobre la **Carta de Derechos Fundamentales en el contexto de la IA** y el cambio digital, de octubre de 2020, que proporcionaron orientaciones sobre la dignidad, las libertades, la igualdad, la solidaridad, los derechos de los ciudadanos y la justicia ante los retos de la IA.
- La Comisión Europea presentó a finales de 2020 dos propuestas de gran calado:
 - **Reglamento de Servicios Digitales**, esencialmente para responsabilizar a las plataformas sobre contenidos digitales (con la amenaza de cuantiosas multas).
 - **Reglamento de Mercados Digitales** para limitar las actividades de algunas de estas empresas (*gatekeepers*, porque otras compañías tienen que usar sus servicios y son capaces de dictar cómo han de funcionar los mercados).

Iniciativas reguladoras de la Inteligencia Artificial

Principales Iniciativas en la UE (IX)

❑ Hacia un régimen europeo de control de la IA (III):

- **Propuesta de *Artificial Intelligence Act* (AIA), Reglamento sobre el uso de la Inteligencia Artificial:** Presentada el 21 de abril de 2021, concreta y avanza “nuevas reglas y acciones para la excelencia y la confianza en la IA” y contiene algunas complejas medidas preventivas en el desarrollo de aplicaciones de IA.
- **Libro Blanco de la IA, presentado en febrero de 2020:** Planteó sólo 2 niveles de riesgo (alto y bajo). Se veían como un binomio muy limitado y que no representaba la realidad en su totalidad. Además, los “requisitos de determinación del nivel de riesgo” eran realmente flojos y abarcaban mucho. Daban un elevado margen de maniobra para etiquetar con cierta laxitud a una aplicación como de riesgo bajo. Tampoco quedaba claro quién determinaría dicho nivel.

Iniciativas reguladoras de la Inteligencia Artificial

Propuesta Reglamento sobre el uso de la Inteligencia Artificial (I)

- ❑ Las aplicaciones de la IA aportarán muchos beneficios, pero pueden conllevar riesgos que minarán la confianza de los usuarios/ciudadanos.
- ❑ Por ello, es necesario regular esos riesgos, diferenciándolos. Ahora bien, de esta legislación están exentas todas las aplicaciones militares de la IA, y deja muchas excepciones para las policiales.
- ❑ **Tipologías de riesgos:**
 - **Riesgo prohibido o inadmisibile:** Algunos riesgos son inadmisibles. La IA que contradice los valores de la UE será prohibida, lo que incluye los sistemas de IA que se consideren una clara amenaza para la seguridad, los medios de subsistencia y los derechos de las personas. Por ejemplo, los sistemas de IA que manipulan el comportamiento humano para eludir la voluntad de los usuarios.
 - **Alto riesgo:** Permitido, pero sujeto al cumplimiento de los requisitos de IA y la evaluación de la conformidad *ex ante*. Por ejemplo, tecnologías empleadas en infraestructuras críticas, la formación educativa o profesional, los componentes de seguridad de los productos o el reclutamiento de empleados, entre otros.
 - **Riesgo limitado:** Sistemas de IA con obligaciones específicas de información/transparencia. Por ejemplo, los usuarios deberán ser conscientes de que están interactuando con una máquina (*chatbots*) para poder tomar una decisión informada de continuar o no.
 - **Riesgo mínimo o nulo:** Resto de sistemas de IA. Por ejemplo, IA aplicada a videojuegos o *spam*.

Iniciativas reguladoras de la Inteligencia Artificial

Propuesta Reglamento sobre el uso de la Inteligencia Artificial (II)

❑ Transparencia:

- Se exige un alto grado de entendibilidad o transparencia de lo que ocurre durante el funcionamiento de los algoritmos.
- La IA debe de ser transparente, lo que supone poder reconstruir cómo y por qué se comporta de una determinada manera.
- Los usuarios que interactúen con los sistemas de IA deben de saber que se trata de IA así como qué personas son sus responsables. Los sistemas deben notificar a los seres humanos, por ejemplo, que:
 - Están interactuando con un sistema de IA, a menos que sea evidente.
 - Se les aplican reconocimiento emocional o sistemas de categorización biométrica.

❑ Sujetos obligados:

- Los **prestadores** que introducen en el mercado o ponen en servicio sistemas de IA en la UE.
- Los **usuarios** de sistemas de IA localizados en la UE.
- Los **prestadores** y **usuarios** de sistemas de IA que están **situados en un tercer país cuando el resultado producido por el sistema se utiliza en la UE**.

Iniciativas reguladoras de la Inteligencia Artificial

Propuesta Reglamento sobre el uso de la Inteligencia Artificial (III)

- ❑ La Comisión Europea propone los siguientes **instrumentos de control** de la IA:
 - Las autoridades nacionales de vigilancia del mercado controlen la **aplicación de las nuevas normas**.
 - Sugiere la creación de **un Comité Europeo de IA** que facilitará su aplicación e impulsará la creación de normas en materia de IA.
 - Se proponen códigos de conducta voluntarios para la IA que no entrañe un alto riesgo, así como espacios controlados de pruebas (*sandboxes*) para facilitar la innovación responsable.
 - **Obligaciones para los operadores:**
 - Establecer e implementar un sistema de gestión de calidad en su organización, de documentación técnica, de registro para permitir a los usuarios supervisar el funcionamiento del sistema de IA de alto riesgo.
 - Someterse a una evaluación de la conformidad y potencialmente re-evaluación del sistema, en caso de modificaciones significativas.
 - Registrar el sistema de IA en la base de datos de la UE.
 - Colaborar con las autoridades de vigilancia del mercado.
 - **Régimen sancionador:** La Comisión Europea abre la posibilidad de importantes multas por incumplimiento de esta nueva reglamentación, como el caso del RGPD, de hasta 30 millones de euros o el 6% de los ingresos globales de la empresa afectada.

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

5.1. Principio de proporcionalidad

5.2. Principio de Trato Justo y No Discriminación

5.3. Principio de Transparencia y Explicabilidad

5.4. Principio de Supervisión Humana

5.5. Principio de Gobierno del Dato

5.6. Principio de Solidez y Rendimiento

Anexos

Principios de gobierno para una IA ética y responsable en el sector asegurador

Introducción

- ❑ En los últimos años han proliferado **varias iniciativas** a nivel internacional, europeo y nacional con el objetivo de promover una IA ética y confiable.
- ❑ Aprovechando estas iniciativas intersectoriales, en particular en el área de Ética Directrices para una IA fiable desarrolladas por el Grupo de Expertos de Alto Nivel (HLEG) en IA de la Comisión Europea, el Grupo consultivo de expertos sobre Ética Digital (GDE) de la EIOPA ha desarrollado unos **principios de gobernanza de IA** para promover una Inteligencia Artificial ética y confiable en el sector de seguros europeo teniendo en cuenta las especificidades de este sector:
 - Principio de proporcionalidad.
 - Principio de trato justo y no discriminación.
 - Principio de transparencia y explicabilidad.
 - Principio de supervisión humana.
 - Principio de gobierno del dato.
 - Principio de solidez y rendimiento.
- ❑ No existe específicamente un **principio de responsabilidad**, pero se menciona explícitamente en el principio de transparencia y explicabilidad, y también está directamente relacionado con todos los demás principios, en particular, con el principio de gobierno del dato.
- ❑ UNESPA ha elaborado 6 principios por un uso ético de la IA muy relacionados con los anteriores: trato justo, responsabilidad proactiva, seguridad, transparencia, formación y evaluación y revisión.

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

5.1. Principio de proporcionalidad

5.2. Principio de Trato Justo y No Discriminación

5.3. Principio de Transparencia y Explicabilidad

5.4. Principio de Supervisión Humana

5.5. Principio de Gobierno del Dato

5.6. Principio de Solidez y Rendimiento

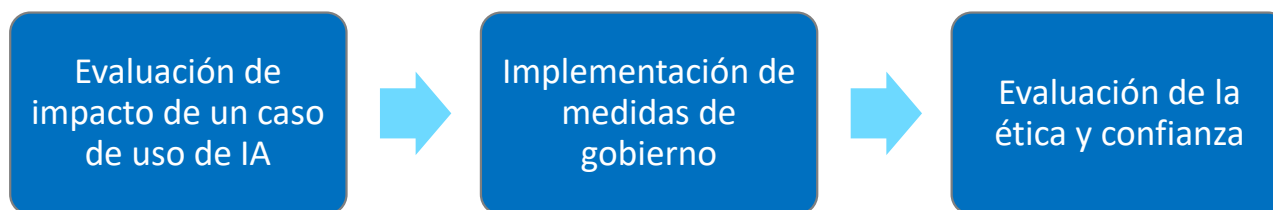
Anexos

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Proporcionalidad (I)

❑ El Principio:

- El principio de proporcionalidad es un **principio común de la legislación europea** en materia de seguros.
- Las empresas de seguros deben establecer las **medidas de gobernanza necesarias que sean proporcionadas** a la naturaleza, escala y complejidad de sus operaciones (Artículo 44.1 del Real Decreto 1060/2015 de Ordenación, Supervisión y Solvencia de Entidades Aseguradoras y Reaseguradoras, RDOSEAR). Esto también se aplica a el uso de IA dentro de la organización.
- **Diferentes impactos:** Los casos de uso de IA en la cadena de valor no tienen el mismo impacto en los consumidores y las entidades. Así, las **medidas de gobernanza** implementadas deberán ser **proporcionadas** a las **características** (impactos) específicas de cada **caso de uso de IA**.
- Proceso de implementación y **evaluación de medidas:**



Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Proporcionalidad (II)

☐ Evaluación del impacto de un caso de uso de IA (Consumidores):

	Alto	Medio	Bajo
Nº de consumidores afectados	Impacto directo en un elevado número de consumidores	Impacto directo en un reducido número de consumidores	No existe impacto directo o se limita a muy pocos consumidores
Intereses de los consumidores	El caso de uso se implementa sobre un proceso con interacción directa con el cliente y/o intereses esenciales del cliente (salario, salud, etc.)	El caso de uso se implementa sobre un proceso con impacto moderado en los derechos y obligaciones de los clientes	Caso de uso implementado en el back office de la entidad sin impacto material en los clientes
Tipo de consumidores afectados	Potencial impacto negativo en consumidores que podrían considerarse vulnerables (ancianos, bajo nivel de estudios, bajos ingresos, etc.)	Baja probabilidad de impacto significativo en los consumidores que podrían considerarse vulnerables	Potencial impacto positivo en consumidores que podrían considerarse vulnerables
Autonomía humana	Los sistemas de IA pueden moldear e influir en el comportamiento de los consumidores a través de mecanismos que pueden ser difíciles de detectar	Los sistemas de IA pueden moldear e influir en el comportamiento de los consumidores, pero a través de mecanismos transparentes y fácilmente entendibles por los consumidores	Los sistemas IA no tienen impacto significativo en el comportamiento de los consumidores
No discriminación	Uso de grandes conjuntos de datos que están sesgados (los datos de entrenamiento tienen una mayor proporción de conjuntos de datos a favor de una etnia o género en particular) o utiliza datos que reflejan discriminación pasada (antecedentes penales) o algoritmos complejos de IA que son capaces de reconstruir información oculta (protegida).	Uso de conjuntos de datos no excesivamente grandes que en gran medida ya se han utilizado en el pasado y se han realizado algunos esfuerzos para eliminar los sesgos. El procesamiento de los conjuntos de datos se realiza mediante complejos algoritmos de inteligencia artificial respaldados por medidas de gobernanza relevantes.	Uso de algoritmos simples y fácilmente entendibles que ampliamente utilizados por la entidad (experiencia) donde se ha demostrado la no existencia de sesgos
Línea de negocio	El caso de uso de IA es aplicado en líneas de negocio consideradas esenciales para los clientes (auto, salud, hogar)	El caso de uso de IA es aplicado en líneas de negocio consideradas no esenciales para los clientes (asistencia)	El caso de uso de IA es aplicado en líneas de negocio pequeñas consideradas no esenciales para los clientes (bicicletas)

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Proporcionalidad (III)

☐ Evaluación del impacto de un caso de uso de IA (Entidades re/aseguradoras):

	Alto	Medio	Bajo
Continuidad de negocio	El caso de uso de IA está implementado en actividades críticas (suscripción, gestión de siniestros...)	El caso de uso de IA está implementado en actividades sensibles	El caso de uso de IA está implementado en actividades no críticas
Riesgo financiero	El fallo del caso de uso tiene consecuencias significativas en los compromisos financieros de la entidad (ratio de solvencia, FFPP)	El fallo del caso de uso tiene impacto en el corto plazo, pero un plan puede mitigar dicho impacto	El fallo del caso de uso no tiene impacto material
Riesgo legal	El fallo del caso de uso supone incumplimiento legal con importantes consecuencias (reversión de autorización, sanciones graves)	El fallo del caso de uso supone incumplimiento legal con ciertas consecuencias (sanciones medias)	El fallo del caso de uso no supone incumplimiento legal alguno o es inmaterial
Riesgo reputacional	El fallo del caso de uso tiene consecuencias importantes sobre la imagen de la entidad	El fallo del caso de uso tiene consecuencias menores en la reputación de la entidad que pueden ser enmendadas	El fallo no afecta a la imagen de la entidad

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Proporcionalidad (IV)

☐ Probabilidad de fallo de un caso de uso de IA:

	Alta	Media	Baja
Evaluación (incluido predicción y perfilamiento)	El sistema de IA califica riesgos, gestiona el fraude, usa BBDD externas o analiza el comportamiento del consumidor	El caso de uso proporciona recomendaciones basadas en las preferencias del consumidor y con su consentimiento	Sin evaluación
Toma de decisiones automatizada	El proceso es autónomo en la toma de decisiones en nombre de la entidad	El caso de uso proporciona recomendaciones para la toma de decisiones	Sin autonomía/involucración en la toma de decisiones
Observación	Los datos se recopilan sin conocimiento del consumidor	Los datos se recopilan con consentimiento del consumidor	observación
Nº variables del modelo	Modelos complejos que involucran un alto número de variables	Uso de un número “maneable” de variables en los modelos o algoritmos	Modelos de baja complejidad como los que implican una descripción determinista simple de la realidad utilizando pocas variables
Procesos	Los sistemas IA implementan nuevos procesos en la organización (nuevos productos, recomendaciones, etc.)	Los sistemas IA implementan procesos existentes con nuevas reglas de negocio (nueva forma de evaluar el riesgo)	No hay nuevos procesos/actividades involucradas (automatización de un proceso existente)
Tratamiento de categorías especiales de datos	Categorías especiales de datos según se definen en el artículo 9.1 de RGPD, así como datos personales relacionados con condenas o delitos penales como se define en el artículo 10 del RGPD, o datos que tienen un alto riesgo de ser un proxy de este tipo de categorías especiales de datos o uso de datos que nunca antes se gestionaron en la organización y que van a ser procesados masivamente o datos sensibles internos	Datos personales distintos de las categorías especiales definidas en los artículos 9.1 ó 10 del RGPD, o uso de datos que, aunque existentes, no fueron procesados o no procesados masivamente	No se usan datos personales o no están en las categorías anteriores
Uso de fuentes de datos	Fuentes de datos externas donde no es posible determinar de manera efectiva cómo se han recopilado / manipulado / procesado los datos	Mix de datos internos y externos	Sólo uso de fuentes de datos internas donde se asegura que los datos han sido procesados responsablemente a lo largo de la cadena de valor

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

5.1. Principio de proporcionalidad

5.2. Principio de Trato Justo y No Discriminación

5.3. Principio de Transparencia y Explicabilidad

5.4. Principio de Supervisión Humana

5.5. Principio de Gobierno del Dato

5.6. Principio de Solidez y Rendimiento

Anexos

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (I)

❑ **EL Principio:** Los consumidores no son discriminados indebidamente como consecuencia de tener ciertos valores de las características protegidas (ver Anexos). La legislación de la UE distingue entre discriminación directa e indirecta:

- **Discriminación directa:** Trato diferente en situaciones comparables debido a una característica protegida.
- **Discriminación indirecta:** Cuando una disposición, criterio o práctica aparentemente neutral pondría a personas con un valor de características protegidas en desventaja con personas con un valor diferente de la misma característica, a menos que esa disposición, criterio o práctica esté objetivamente justificada por un fin legítimo y los medios para lograr ese objetivo son apropiados y necesarios.

❑ **Regulación:**

- **Artículo 17.1 de la Directiva 2016/97 de Distribución de Seguros (DDS):** Los distribuidores de seguros actuarán siempre de forma honesta, justa y profesional de acuerdo con los mejores intereses de sus clientes (en línea con el principio de procesamiento justo de datos del artículo 5 del RGPD).
- **Artículo 20.1 DDS:** Cualquier producto propuesto al cliente debe ser coherente con sus demandas y necesidades.
- **Artículo 25 del DDS y el Reglamento Delegado 2015/35:** Establecen requisitos de supervisión y gobernanza de productos para los distribuidores de seguros, incluida la necesidad de identificar un mercado objetivo y evaluar periódicamente que el producto sigue siendo coherente con las necesidades del mercado objetivo.
- El **principio de trato justo** también se reconoce en el artículo 5 del RGPD y en el artículo 5 de la Directiva 2005/29 relativa a las prácticas comerciales desleales de las empresas en sus relaciones con los consumidores.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (II)

❑ Procesos de gobernanza de los sistemas de IA:

- Los **procesos de gobernanza deben ser sólidos y transparentes** para asegurar un trato justo y no discriminatorio:

Trato justo procedimental (Foco en la gobernanza)

- Requisito de una **conducta comercial justa de la entidad aseguradora frente al consumidor**.
- **La DDS requiere que:**
 - Los distribuidores actúen de forma honesta, justa y profesional de acuerdo con los mejores intereses de sus clientes (artículo 17.1 de la DDS).
 - Los productos de seguros propuestos a los consumidores sean coherentes con sus demandas y necesidades.
 - Los consumidores reciban información objetiva de forma comprensible para permitirles tomar decisiones informadas (artículo 20.1 de la DDS).
- Varias de las medidas de gobernanza presentadas aquí tienen como objetivo implementar el trato justo procesal cuando se utiliza la IA, por ejemplo, **promoviendo una transparencia y explicabilidad adecuados**.

Trato justo distributivo (Foco en los resultados)

- **Accesibilidad de la cobertura a precios asequibles y sin prejuicios ni discriminación.**
- Distribución justa tanto de los beneficios como de los costes.
- **Particularmente relevante cuando se utilizan sistemas de IA de “caja negra”** para procesar grandes conjuntos de datos donde los controles tradicionales no se pueden implementar fácilmente (por ejemplo, peso de las variables, varianza, asimetría, etc.). **Un área de investigación en crecimiento en la comunidad de la ciencia de datos busca desarrollar métricas de trato justo y no discriminación para evaluar los resultados de los sistemas de IA.**
- La evaluación de impacto de casos de uso de IA **insta a que las entidades evalúen el impacto de los sistemas de IA en grupos de consumidores vulnerables, la discriminación ilegal y la accesibilidad a ciertos productos que son importantes para garantizar la inclusión financiera.**

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (III)

❑ Desarrollo justo de sistemas de AI buscando el equilibrio de los intereses de todos los *stakeholders*:

- Las entidades deben **encontrar un equilibrio entre los intereses** (variados y cambiantes) **de las diferentes partes** al considerar los desafíos éticos de la IA (como parte de su compromiso con la Responsabilidad Social Corporativa):
 - **Intereses de los consumidores:** Buscan obtener la cobertura de un riesgo a un precio asequible.
 - **Intereses de la entidad y sus accionistas:** Sostenimiento de un negocio rentable en un mercado competitivo.
 - **Intereses de los mutualistas:** Los intereses del conjunto de asegurados pueden diferir de los intereses particulares (el alto riesgo de determinados mutualistas puede producir un incremento de la prima).

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (IV)

❑ Productos particularmente relevantes para la inclusión financiera:

- El GDE (*Group on Digital Ethics in Insurance*) de EIOPA indica que los siguientes **productos** son **especialmente relevantes** en relación a la **inclusión financiera de los consumidores**:

Producto	Relevancia
Auto	La falta de seguro de automóvil puede afectar negativamente el nivel de movilidad requerido para la empleabilidad, así como el nivel de vida mínimo social (por ejemplo, donde el transporte público es inadecuado).
Salud	Un seguro de salud inadecuado puede impedir el acceso a atención médica adecuada con el consiguiente impacto en la sociedad.
Hogar	Protección frente a la pérdida de bienes. Particularmente importante en familias con préstamos hipotecarios. En algunos estados miembros de la UE, es un requisito para alquilar o comprar vivienda.
Responsabilidad civil	Similar al caso anterior del seguro de hogar al proteger la pérdida patrimonial.
Vida / Pensiones	Aporta tranquilidad tras la jubilación.
Accidentes	Proporciona ingresos y beneficios médicos a las personas que no pueden trabajar debido a las lesiones sufridas.
Desempleo	Proporciona ingresos a las personas que han perdido su trabajo.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (V)

❑ Potencial de los sistemas de IA para mejorar la inclusión financiera de los consumidores de alto riesgo:

- **Definición de exclusión financiera:** Proceso por que las personas encuentran dificultades para acceder a productos que se adapten a sus necesidades y les permitan llevar una vida normal.
- Podemos encontrar **situaciones positivas** (uso de sistemas telemáticos) y **situaciones negativas** (inclusión de pruebas de ADN) para los consumidores en relación a la granularidad en la evaluación de riesgos:
 - **Situación negativa:** Los avances en IA permiten a las entidades **evaluar los riesgos con una mayor nivel de granularidad**:
 - El uso del ADN revelaría condiciones pre-existentes desconocidas que supondrían la imposibilidad de contratar el seguro de vida para algunas personas.
 - Personas que vivan en zonas afectadas por el cambio climático podrían tener problemas a la hora de asegurar sus viviendas en base a la evaluación de riesgos con una mayor nivel de granularidad.
 - **Situación positiva:** La mejor comprensión de los riesgos en combinación con la mitigación de riesgos pueden mejorar la inclusión financiera de algunos consumidores de alto riesgo accediendo a una cobertura asequible.
- Un informe de EIOPA de 2019, respecto a los seguros de salud y auto, **no encontró evidencia de que una granularidad cada vez mayor de las evaluaciones de riesgos estuviera causando importantes problemas de exclusión para personas de alto riesgo**. Sin embargo, las entidades que participaron en el estudio, **esperaban que el impacto aumentara en los años sucesivos** como resultado del uso de algoritmos más precisos en conjunción con nuevos datos.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (VI)

❑ Consumidores de especial atención (I):

- **Otros consumidores a los que hay que prestar especial atención:** Además de los clientes de alto riesgo que pueden verse perjudicados por la evaluación de riesgo con mayor nivel de granularidad, existen **otros consumidores** a los que hay que prestar una especial atención por su **alta vulnerabilidad**:
 - Personas que, por ejemplo, por **pobreza, discapacidades físicas o mentales, edad**, etc., han reducido las oportunidades de participación social.
 - Un **informe** de la **OCDE** de 2019 concluyó que los **consumidores** con **menos recursos y experiencia** con el entorno **digital** son cada vez **más vulnerables**.
 - Un **estudio** de **FCA** de Reino Unido en 2019 encontró que **1 de cada 3 consumidores** que **pagaron primas elevadas** mostraron **al menos un característica de vulnerabilidad**, como un bajo nivel de capacidad financiera. Este mismo estudio concluyó que los consumidores con menor renta pagaban mayores primas en los seguros de hogar frente a aquellos con mayores rentas.
- **Adopción de políticas:** Las entidades aseguradoras podrían adoptar políticas para evitar situaciones como las anteriores:
 - Por ejemplo, en los seguros de coche, podrían **usar los dispositivos telemáticos para sustituir la evaluación de características del cliente como la edad por el análisis de la conducción** del cliente, por dónde conduce, su estilo de vida, etc.
 - Otra opción, con ayuda de los gobiernos, sería la de **diseñar productos asequibles con coberturas básicas para aquellos colectivos vulnerables**.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (VII)

❑ Consumidores de especial atención (II):

- En la siguiente tabla se muestran **tipos de vulnerabilidades de los consumidores**:

Tipo de vulnerabilidad	Ejemplos de consumidores
Características personales	- Ancianos. - Bajos ingresos. - Bajo nivel educativo. - Miembros de minorías. - Migrantes. - Jóvenes/estudiantes. - Personas que vivan en zonas afectadas por el cambio climático.
Situación personal	- Desempleados. - Vagabundos. - Padres divorciados. - Personas sobre-endeudadas. - Personas con impagos. - Presos. - Lesionados en accidentes. - Víctimas de violencia de género.
Estado de salud	- Discapacitados. - Personas con enfermedades hereditarias. - Embarazadas. - Personas con problemas de salud mental o en terapia.
Habilidades digitales	- Personas con bajas habilidades en el entorno digital. - Personas con problemas o sin acceso al entorno digital (internet).

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (VIII)

❑ El papel de las autoridades públicas:

- **Las autoridades también tienen un papel relevante en el fomento de la ética y de la inclusión financiera** en el mercado de seguros **mediante el desarrollo legislación adecuada** (por ejemplo, protección de datos, conducta de legislación comercial, etc.).
- En Francia hay leyes restringiendo el tipo de datos que se pueden usar en seguros de salud.
- En Dinamarca no se permite el uso de datos de ADN en la vida y seguro de salud.
- En Francia, Bélgica, Luxemburgo y los Países Bajos han adoptado o están en proceso de adoptar medidas legales para limitar el uso de datos de salud (y, más particularmente, en lo que respecta a los sobrevivientes de cáncer).

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (IX)

❑ Uso justo de datos:

- **Finalidad adecuada:** Los **datos utilizados** por las entidades deben ser **adecuados al fin que se persigue** (Principio de limitación del propósito del Artículo 5 del RGPD).
- **Transparencia:** Las **entidades** de seguros deben ser transparentes **en cuanto a cómo utilizan los datos y deben ser capaces de explicar fácilmente el uso que dan a los datos**.
- **Privacidad:** Ciertos conjuntos de datos pueden ser especialmente sensibles y mostrar comportamientos y hábitos muy privados de los consumidores. Ejemplos:
 - **Datos sanitarios** utilizados por entidades de seguros de salud.
 - **Datos recopilados a través de *wearables*** o dispositivos telemáticos en automóviles.
 - **Información de los hábitos de compra** de los consumidores a través del análisis del uso de sus tarjetas de crédito.
- **Reducción del riesgo:** Los datos se pueden utilizar para proporcionar valiosos servicios a los consumidores, como **sugerencias para mejorar habilidades de conducción o estilos de vida saludables**, o los datos del paciente pueden ser esenciales para investigar enfermedades y nuevos tratamientos personalizados.



Información privada como "¿dónde vas de compras?", "si vas a comer a restaurantes de comida rápida" o información como "¿quiénes son tus amigos?", "¿sufres de enfermedad mental?" pueden derivarse de la telemática o de los *smartphones* (por ejemplo, geolocalización), pero **dicha información debería excluirse de los modelos de análisis de datos**, por ejemplo, para la fijación de primas.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (X)

☐ Opciones en el caso de consumidores que no estén dispuestos a compartir sus datos:

- **Entidades:** Tendrían 2 posibilidades, en este caso:
 - Explicar claramente a los clientes que podrían ofrecerles pólizas más personalizadas en base a esa información requerida.
 - Ofrecer a los consumidores alternativas de cobertura.
- **Consumidores:** No deben aprovecharse del análisis genético para adquirir seguros sin informar de forma transparente a la entidad sobre el hecho de que tienen esta información.

☐ Consumidores con escasas habilidades digitales, o con dificultades o falta de voluntad para acceder a los servicios digitales:

- Las entidades de seguros deben **ser conscientes de estas limitaciones** (falta de datos de esta tipología de clientes) **al entrenar sus sistemas de IA y definir ofertas de productos** para este mercado objetivo.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (XI)

❑ Protección de la autonomía humana:

- **Los sistemas de IA** utilizados por las entidades de seguros **deberían ayudar a los consumidores en la toma de decisiones informadas** al comprar productos de seguros.
- **No deben subordinar, engañar o manipular injustificadamente a los consumidores:** Esto es particularmente relevante en el área de la economía del comportamiento donde los sistemas de IA se pueden implementar para moldear e influir en el comportamiento humano a través de mecanismos que pueden ser difíciles de detectar (aprovechando los procesos subconscientes del ser humano).
- **Ejemplos de buenas prácticas:**
 - Las entidades de seguros pueden ayudar a los consumidores a tomar decisiones más informadas al **"empujarlos" de manera transparente hacia una situación más saludable, más seguro o hacia opciones económicamente más beneficiosas para ellos**. Esto puede ser importante para lograr una mayor satisfacción del cliente (p. ej. elegir determinados centros de salud o talleres de reparación).
 - Los llamados sistemas de recomendación basados en IA pueden ayudar a los consumidores **recomendándoles productos o coberturas que de otra forma no se habrían tenido en cuenta por parte de estos últimos**.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (XII)

❑ Sesgos en la información:

- Existen diversas **fuentes de sesgo**:
 - **Uso de características protegidas**: Cuando un conjunto de datos contiene características protegidas o variables *proxy* aparentemente neutrales que se correlacionan mucho con las características protegidas.
 - **Algoritmos supervisados de aprendizaje automático** (el 90% de todos los sistemas de IA usan variables / conjuntos de datos etiquetados). Esto implica juicios humanos que potencialmente podrían reflejar los prejuicios de la persona que etiquetó los datos (por ejemplo, los peritos de siniestros que etiquetan un siniestro como fraudulento).
 - **Proceso de recopilación de los datos**: Si los datos no son representativos y no reflejan la distribución real, un sistema de IA que utiliza estos conjuntos de datos para el entrenamiento funcionará mal.
- **Cualquier sesgo, error o incorrección en los datos** utilizados para entrenar el modelo de IA, accidental o intencional, **se reflejará en la salida del sistema de IA**.
- **No es posible eliminar completamente los sesgos en los datos**. Los datos pueden estar más ponderados que otros (por ejemplo, siempre habrá más hombres o más mujeres conduciendo determinados vehículos).

❑ **Cómo evitar los sesgos en la información (I)**: Las entidades de seguros deben hacer un esfuerzo razonable para eliminar sesgos en los datos utilizados por los sistemas de IA:

- **Eliminando** de los conjuntos de datos usados **las características protegidas**.
- **Eliminando variables *proxy* que podrían correlacionarse con características protegidas**, a menos que su uso esté objetivamente justificado, sea apropiado y necesario.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (XIII)

❑ Cómo evitar los sesgos en la información (II):

- **Análisis de la idoneidad:** Las entidades de seguros deben **evaluar la idoneidad de sus modelos** (algoritmo y datos necesarios para el algoritmo), **especialmente, para aplicaciones de IA de alto impacto.**
- **Revisar los resultados de los sistemas de IA para identificar posibles sesgos.** El cuadro siguiente muestra métricas sobre cómo se podría hacer esto en la práctica:

Métrica	Descripción
Paridad demográfica	El objetivo es asignar un resultado positivo a tasas proporcionalmente iguales a cada subgrupo de una clase protegida donde el resultado positivo se refiere a la decisión favorable. Por ejemplo, en el contexto de un escenario de reclutamiento, la "paridad demográfica" podría significar que los candidatos masculinos y femeninos son invitados a entrevistas de trabajo en proporciones iguales, proporcionalmente al número de solicitudes.
Calibración	Garantiza la igualdad en las predicciones entre los diferentes subgrupos. Por ejemplo, en un escenario de diagnóstico médico, un modelo calibrado podría asegurar niveles iguales de confianza en las predicciones para pacientes de diferente género o antecedentes genéticos.
Igualdad de probabilidades	Significa las mismas tasas de verdaderos positivos y verdaderos negativos para todos los subgrupos. Por ejemplo, cuando una entidad utiliza sistemas de IA para escanear los CV y las solicitudes de empleo en los procesos de contratación, la "igualdad de probabilidades" garantizaría que las probabilidades para hombres y mujeres de ser invitados a la entrevista de trabajo son iguales.
Igualdad de oportunidades	Requiere que las tasas de error para el resultado favorable sean las mismas, pero permite desviaciones para el resultado desfavorable. Por ejemplo, en el marketing online, cuando el objetivo es informar a hombres y mujeres por igual sobre una oferta de seguro, la "igualdad de oportunidades" podría garantizar que los segmentos relevantes de ambos grupos reciban la información a la misma velocidad. Sin embargo, la tasa de exposición a personas para las que la oferta es realmente irrelevante puede diferir.
Equidad individual	"Equidad individual" abandona la idea de pertenencia a un grupo y sugiere que cualquier individuo similar debería ser tratado de manera similar. Por ejemplo, todas las personas con el mismo perfil de riesgo deben pagar la misma prima por el mismo producto de seguro.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (XIV)

❑ Caso de uso: Tarificación y suscripción (I):

- **Tarificación basada en factores de riesgo adecuados** según la legislación antidiscriminación:
 - **Las entidades de seguros tienen que cumplir con la legislación contra la discriminación** tanto a nivel de la UE como a nivel nacional.
 - Las leyes nacionales de antidiscriminación pueden diferir de un país a otro en Europa, pero suelen contener algunas disposiciones específicas en relación a la tarificación que indican qué atributos son inadmisibles en la evaluación de riesgos y qué atributos son admisibles (típicamente, discapacidad y edad).
 - La Comisión Europea emitió una guía aclarando que los factores de calificación que se correlacionan con el género, y, por lo tanto, pueden causar una discriminación indirecta, pueden estar "objetivamente justificados por un fin legítimo " cuando se utilizan para la estimación de riesgos/costes.
 - Un enfoque similar debe utilizarse con otras características protegidas. El requisito de que "los medios sean adecuados y necesarios" es traducido por la Comisión Europea en la noción de "**verdaderos factores de riesgo**".

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (XV)

❑ Caso de uso: Tarificación y suscripción (II):

- Las entidades de seguros **deben evaluar y desarrollar medidas para mitigar el impacto de ciertos factores de valoración del riesgo** como el crédito, la localización, los ingresos, la ocupación o el nivel de educación sobre poblaciones vulnerables y protegidas:
 - **Factores de valoración del riesgo correlacionados con el riesgo de siniestro** (por ejemplo, las personas con perfiles de crédito bajos comunican más siniestros que las personas con perfiles de crédito altos). Su uso puede tener un **impacto negativo desproporcionado sobre los consumidores vulnerables**, como las poblaciones de bajos ingresos o ciertos grupos minoritarios, y contribuir así a reforzar la desigualdad:
 - Cuando las poblaciones de migrantes llegan a un nuevo país tienen evaluaciones de riesgo bajas (malas) porque hay poca información crediticia sobre ellos.
 - En tiempos de crisis económica como en la actual pandemia de Covid-19, muchas personas que han perdido su trabajo verán un empeoramiento de sus calificaciones de riesgo.
 - **La información de comportamiento puede ser inexacta y estar estrechamente correlacionada con características protegidas** y, además, es posible que los consumidores no sepan que dicha información sobre ellos se utiliza para valorar su riesgo. Por lo tanto, dicha información no debe utilizarse con fines de fijación de primas.
 - **Micro-segmentación por variables no relacionadas con el riesgo a cubrir**: Las entidades de seguros deben **evitar el uso de factores de calificación en el seguro de auto que consideren no relacionados directamente con la conducción** (código postal, ocupación, estado civil o nivel de estudios).

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (XVI)

❑ Caso de uso: Tarificación y suscripción (III):

- **Las prácticas de optimización de precios deberían evitarse cuando perjudiquen injustamente a los consumidores**, sobre todo, en consumidores vulnerables o clases protegidas:
 - **Algunas entidades de seguros ajustan la prima al precio de mercado y optimizan la prima final utilizando una serie de diferentes técnicas que son, en gran medida, independientes de la del perfil de riesgo del consumidor.** Normalmente se basan en correlaciones de los atributos relevantes del riesgo, así como otros factores que no se basan en el riesgo como los ingresos, el nivel de estudios, tipo de dispositivo utilizado (marca de teléfono inteligente, tableta, portátil, etc.), canal de distribución, localización, aplicaciones descargadas, etc.
 - Por ejemplo, **clientes leales con una baja sensibilidad al precio o baja propensión a comparar precios pueden no ser conscientes de que pueden estar pagando primas mucho más altas** que los consumidores con un perfil de riesgo similar, pero con mayor sensibilidad al precio o mayor propensión a comparar precios.
 - En el Reino Unido, la **FCA realizó una investigación de mercado en 2020** sobre las prácticas de fijación de precios de las entidades de seguros donde **encontró evidencias de que las entidades a menudo aumentan deliberadamente las primas a los consumidores menos propensos al cambio.**
 - La **FCA propuso nuevas reglas para evitar que las entidades aumentaran el precio de renovación para los consumidores a lo largo del tiempo**, exigiendo a las entidades que ofrezcan un precio de renovación a los consumidores igual que el que ofrecen a los nuevos clientes.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Trato Justo y No Discriminación (XVII)

❑ Caso de uso: Gestión de siniestros:

- Prácticas de **predicción** de la disposición de un consumidor a **aceptar** la **indemnización propuesta**:
 - Ofrecer a alguien con las mismas características, la misma lesión en el mismo tiempo y con el mismo producto una compensación materialmente diferente no es apropiado.
 - Los sistemas de IA que utilizan datos de comportamiento para estimar la elasticidad precio del consumidor, la probabilidad de caída o, más en general, su disposición a aceptar una oferta/indemnización concreta, no deberían admitirse.
- **Detección de fraude**:
 - **Red personal**: Ciertas técnicas de prevención del fraude, como las técnicas de análisis de redes, podrían perpetuar los patrones existentes de vulnerabilidad/discriminación a través de la "**discriminación de la red**", mediante la cual las personas son penalizadas (o recompensadas) en función de las características de su familia, amigos u otros miembros de su red personal.
 - **Datos de entrenamiento de algoritmos**: Al igual que en otros casos de uso de IA, las entidades de seguros deben realizar esfuerzos para **eliminar posibles sesgos en los datos de entrenamiento de los algoritmos**.
 - **Supervisión humana**: Las entidades de seguros deben establecer medidas de gobierno para mitigar dichos riesgos, como establecer niveles adecuados de supervisión humana. **Ningún modelo debería clasificar a un consumidor como fraudulento sin la involucración de expertos**.
 - **Explicabilidad**: Los modelos de fraude deberán ser explicables, aportando suficiente información como para entender por qué se ha identificado un asegurado o siniestro en particular.

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

5.1. Principio de proporcionalidad

5.2. Principio de Trato Justo y No Discriminación

5.3. Principio de Transparencia y Explicabilidad

5.4. Principio de Supervisión Humana

5.5. Principio de Gobierno del Dato

5.6. Principio de Solidez y Rendimiento

Anexos

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Transparencia y Explicabilidad (I)

❑ El Principio:

- Las entidades de seguros deben **adaptar la explicabilidad de sus sistemas a los casos de uso de IA y a los *stakeholders*.**
- Las entidades tienen **2 opciones**:
 - Deben esforzarse por **utilizar modelos de IA explicables**, sobre todo, **en casos de uso de IA de alto impacto**, o
 - **Pueden combinar la explicabilidad del modelo con otras medidas de gobernanza** en la medida en que se garantice la responsabilidad de las entidades.

❑ Regulación:

- El artículo 20.1 de la DDS indica que los proveedores **proporcionarán al cliente información objetiva** sobre el producto de seguro de forma comprensible **para permitir que el cliente tome una decisión informada.**
- Se reconoce el **principio de transparencia en el artículo 5 del RGPD.**
- Los artículos 13 y 14 del RGPD requieren que las entidades:
 - **Informen oportunamente a los consumidores de manera prudente y transparente sobre cómo se procesan sus datos personales.**
 - **Informen a los consumidores sobre la existencia de procesos automatizados de toma de decisiones** (a menudo respaldados por la IA) y, lo que es más importante, **aportarles información "significativa" sobre la lógica involucrada**, así como **el significado y las consecuencias de dicho procesamiento.**

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Transparencia y Explicabilidad (II)

- ❑ El uso de sistemas de IA presenta **desafíos y oportunidades** cuando estamos hablando de **transparencia y explicabilidad**:
 - **Desafíos:**
 - Los **sistemas de IA “caja negra”** usados para la tarificación y la suscripción, suelen ser **difíciles de explicar** y, por tanto, **difíciles de entender por los consumidores a la hora de saber por qué se ha denegado una cobertura**.
 - La **falta de explicabilidad** también **compromete la auditoría de estos sistemas**.
 - **Oportunidades:**
 - Se puede argumentar que **la automatización de ciertas decisiones y los procesos impulsados por sistemas de IA aumentan su transparencia y explicabilidad** en la medida en que la automatización de decisiones previamente tomadas por humanos es ahora más transparente porque, **siempre que exista un adecuado registro de datos y metodologías utilizadas**, dejan un **rastro digital** y, por lo tanto, son potencialmente replicables y más fácil de monitorear.
 - En el área de **gestión de siniestros**, la **decisión** de aceptar o rechazar un siniestro **ha sido tradicionalmente realizada por gestores de la entidad con sus posibles prejuicios** de estos individuos podrían haber pasado desapercibidos. **Esto sería más difícil si la decisión se automatiza por medio de un sistema de IA explicable** con parámetros e insumos predefinidos debidamente trazables/documentados.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Transparencia y Explicabilidad (III)

- ☐ Información a proporcionar a los diferentes *stakeholders* cuando se emplean sistemas IA en la tarificación y suscripción:

Información	Stakeholders		
	Consumidores	Auditor/Supervisor	Dirección
El proceso está automatizado con sistemas de IA	✓	✓	
Conjuntos de datos utilizados	✓	✓	
Criterios utilizados por el sistema de IA	✓	✓	
Explicación de los factores de calificación más relevantes	✓	✓	
Descripción de la integración del sistema de IA en la infraestructura IT		✓	
Metodologías de recopilación, preparación y post-procesamiento de datos		✓	
Limitaciones del sistema de IA		✓	
Código y datos utilizados para entrenar y testar el modelo		✓	
Medidas de seguridad del sistema de IA		✓	
Razones del uso de la IA y coherencia con las estrategias de la entidad		✓	✓
Personal involucrado en el diseño e implementación del sistema de IA		✓	✓
Rendimiento del modelo (incluidos KPIs)		✓	✓
Evaluación de la ética y confiabilidad del sistema de IA		✓	✓
Documentación sobre el cumplimiento de la normativa vigente		✓	✓
Existencia de auditoría externa del sistema de IA		✓	✓
Explicación de la lógica del sistema a un no experto		✓	✓
Tecnologías implementadas por terceros y riesgos asociados		✓	✓

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Transparencia y Explicabilidad (IV)

❑ Mejorar la explicabilidad de los casos de uso de IA:

- En todo caso, las entidades de seguros deben **usar la mayor cantidad posible de modelos de IA explicables**.
- **Casos de uso de IA de alto impacto para consumidores y/o entidades de seguros:** Las entidades deberían poder explicar suficientemente el proceso de toma de decisiones del modelo de IA, incluso a expensas del rendimiento del modelo.
- **Explicabilidad vs Rendimiento del modelo de IA:** La explicabilidad se ve afectada por la elección del modelo y, a menudo, existe una compensación entre la explicabilidad del modelo y el rendimiento (por ejemplo, los árboles de decisión son más explicables, pero quizás no tan precisos como las redes neuronales que son mucho más opacas).
- Existen **3 enfoques principales** que las entidades pueden utilizar para la creación de modelos de IA transparentes/explicables:
 - **Uso de sistemas de IA explicables.**
 - **Uso de sistemas de “caja negra” combinados con herramientas que faciliten la explicabilidad de este tipo de sistemas de IA:** Se deben documentar adecuadamente las limitaciones de estas herramientas y las pruebas a las que se ha sometido el sistema de IA.
 - **Uso de sistemas híbridos:** Los sistemas de IA explicables se pueden combinar con sistemas de “caja negra” que se entrenan con los mismos datos con el propósito de compararlos. Así, por comparación, se puede mejorar la explicabilidad del sistema.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Transparencia y Explicabilidad (V)

- ❑ **Combinar explicabilidad con medidas de gobernanza:** Existen otras medidas de gobernanza que las entidades de seguros pueden implementar para garantizar sistemas de IA éticos y confiables:
 - En caso de que los **beneficios económicos** del uso de sistemas de IA de “caja negra” **superen** indudablemente los **riesgos**, y/o en ausencia de alternativas, se requiere:
 - **Mayor nivel de supervisión humana** (*human-in-the-loop*), y
 - **Gestión adecuada de datos a lo largo del ciclo de vida del modelo de IA.**
 - **Falta de explicabilidad + supervisión humana:**
 - **Compensar la falta de explicabilidad incluyendo un humano en el proceso puede no ser suficiente si el sistema de IA proporciona resultados sesgados injustificados** porque se está entrenado con conjuntos de datos sesgados.
 - Además de tener niveles adecuados de supervisión humana, **las entidades deben establecer procesos sólidos de gestión de datos en tales casos de uso de IA. Las entidades de seguros necesitan pruebas estadísticas adicionales para revelar dónde podría surgir una posible discriminación.**
 - Para **modelos automatizados con supervisión humana limitada**, las entidades deben adoptar una **mayor explicabilidad** del modelo.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Transparencia y Explicabilidad (VI)

❑ Información al consumidor sobre los datos usados en los sistemas de IA:

- De conformidad con los artículos 13 y 14 del RGPD, **las entidades de seguros deben informar a los consumidores sobre los tipos, fuentes y propósitos del uso y procesamiento de datos personales en los modelos de IA.**
- El artículo 6 del RGPD establece la **necesidad de obtener el consentimiento del cliente para ciertos tipos de tratamiento de datos desde una perspectiva de transparencia:**
 - Los consumidores deben recibir la **información necesaria** y, a través de un **proceso apropiado**, tener la **capacidad de comprender la información clave** que el consumidor necesita **para elegir la opción correcta.**
 - Esta información **no debe proporcionarse en clausulados muy extensos, difíciles de entender o declaraciones llenas de jerga jurídica.**
- **Tarificación y suscripción – Información a proporcionar al consumidor:**
 - **Principales factores de valoración del riesgo que afectan a la prima** para reforzar la confianza y permitirles adoptar decisiones informadas.
 - **Tipos de datos para poder verificar si los datos considerados por la entidad son relevantes y precisos.** Así, se evitaría que las entidades de seguros se basen en datos que los consumidores podrían considerar potencialmente sensibles, intrusivos o discriminatorios.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Transparencia y Explicabilidad (VII)

❑ Información al consumidor sobre la interacción con sistemas de IA y sus limitaciones:

- **Uso de *chatbots*** para interactuar con los consumidores en procesos no sensibles (por ejemplo, ayudar al consumidor a navegar a través de su sitio web):
 - Es importante que los **consumidores** sean **conscientes** de que están **interactuando** con un **sistema de IA** y no con un ser humano.
 - Los consumidores también deben recibir **información significativa y oportuna sobre las capacidades y limitaciones del *chatbot***.
 - Los consumidores deberían poder **solicitar la intervención de un empleado** en algún momento del proceso.
- **Uso de robots asesores** para brindar asesoramiento a los consumidores (por ejemplo, sobre opciones de inversión en seguros de vida):
 - Los consumidores deben recibir una **comprensión básica de los algoritmos detrás de las recomendaciones** (por ejemplo, nivel de apetito al riesgo, preferencia por productos "verdes" o ciertos sectores en los que invertir, etc.)
 - Los consumidores deben recibir **información** sobre si la entidad también **proporciona asesoramiento humano y cómo se puede acceder a él**.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Transparencia y Explicabilidad (VIII)

❑ Caso de uso: Tarificación y suscripción:

- Las entidades de seguros deben poder **explicar a supervisores y auditores que los principios detrás del modelo de tarificación son sólidos** y los **consumidores son informados de los principales factores de valoración del riesgo** que influyen en la prima para reforzar su confianza:
 - La necesidad de transparencia/explicabilidad es una de las razones por las que los modelos GLM se utilizan normalmente para llegar a las tarifas, son simples de explicar y permiten el análisis de "qué pasaría si" de una manera simple.
 - La necesidad de explicabilidad también tiene una importancia primordial desde la perspectiva de la equidad del consumidor. El consumidor no sabe qué sistemas están detrás del cálculo de su prima, por lo que deben ser informados sobre los principales datos/factores que afectan a la prima. Esto también es válido en el caso de factores "no tradicionales" como la valoración del comportamiento de conducción a través de aplicaciones telemáticas.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Transparencia y Explicabilidad (IX)

❑ Caso de uso - Detección del fraude:

- En este caso, nos referimos a **aplicaciones que usan sistemas de IA “caja negra”** y que proporcionan indicaciones/recomendaciones a los gestores de siniestros sobre cuáles deben **priorizar para investigarlos con más detalle** (inclusión humana en el proceso), mejorando el proceso de gestión de siniestros.
- Adicionalmente, las entidades aseguradoras deben asegurarse de que **los sistemas de AI están debidamente entrenados con conjuntos de datos no sesgados** para evitar que los gestores de siniestros reciban recomendaciones inadecuadas.
- También se deben **establecer procedimientos de reparación efectivos ante eventuales problemas** ocasionados por los sistemas de IA.

❑ Caso de uso - Detección del fraude con sistemas OCR (*Optical Character Recognition*):

- Los sistemas OCR permiten a las entidades el procesamiento automático de siniestros escaneando la información proporcionada por el asegurado.
- En este caso, los sistemas de IA toman los datos proporcionados por los asegurados en los formularios de siniestro, detectando errores e inconsistencias en la declaración del siniestro.
- Los sistemas de IA basados en redes neuronales pueden procesar imágenes de cara a identificar patrones, mejorando la eficiencia, los tiempos y la precisión de la gestión de siniestros.
- Tanto **los sistemas basados en OCR como los de reconocimiento de imágenes, son sistemas de IA “caja negra” que requieren de medidas de gobernanza y supervisión humana cuando se superen ciertos umbrales** (sobre todo, en las primeros momentos de puesta en producción de estos sistemas) de cara a verificar que los sistemas de IA están funcionando como se espera.

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

5.1. Principio de proporcionalidad

5.2. Principio de Trato Justo y No Discriminación

5.3. Principio de Transparencia y Explicabilidad

5.4. Principio de Supervisión Humana

5.5. Principio de Gobierno del Dato

5.6. Principio de Solidez y Rendimiento

Anexos

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Supervisión Humana (I)

❑ El Principio:

- Las entidades aseguradoras deberían **establecer** unos **niveles adecuados de supervisión humana a lo largo del ciclo de vida de los sistemas de IA**.
- Las entidades aseguradoras deberán **asignar roles y responsabilidades a los empleados** vinculados con los sistemas de IA.
- Las entidades deberán **proporcionar a los empleados** vinculados con sistemas de IA una **adecuada formación** por 2 razones principales:
 - Las tareas de los empleados pueden cambiar radicalmente.
 - Algunos trabajadores pueden ser desplazados por procesos automatizados.

❑ Regulación:

- El artículo 41 de la Directiva de Solvencia II (artículo 44 del RDOSEAR) establece que todas las entidades aseguradoras y reaseguradoras dispondrán de un **sistema eficaz de gobierno que garantice una gestión sana y prudente de la actividad** y que comprenderá, como mínimo, una estructura organizativa transparente y apropiada, con una clara distribución y una adecuada separación de funciones, y un sistema eficaz para garantizar la transmisión de la información.
- La Directiva de Solvencia II hace referencia a la creación de las **funciones clave** (riesgos, actuarial, cumplimiento y auditoría) que, junto con la figura del *Data Protection Officer* (DPO), en cumplimiento de lo indicado en el artículo 37 del RGPD (Designación del delegado de protección de datos), proporcionan **líneas de defensa para gestionar los posibles problemas relacionados con la IA**.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Supervisión Humana (II)

❑ Establecimiento de adecuados niveles de supervisión humana:

- La selección de un nivel superior o inferior de supervisión humana **debe ser proporcional a la naturaleza, escala y complejidad de los riesgos inherentes al uso específico de la IA**, teniendo en cuenta la combinación de medidas de gobernanza establecidas sobre los casos de uso de IA.
- **Sistemas con baja explicabilidad/transparencia:** Las entidades pueden **compensar** la falta de explicabilidad con niveles mayores de supervisión humana y gestión de datos a lo largo del ciclo de vida del modelo de IA.
- **Sistemas con limitada supervisión humana:** Las entidades deberían **adoptar medidas para garantizar una mayor explicabilidad de estos modelos y otras medidas de gobernanza** (adecuada gestión de datos, garantizar la solidez y rendimiento de los sistemas de IA, etc.), **especialmente para los casos de uso de IA de alto impacto para los clientes.**

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Supervisión Humana (III)

❑ Responsabilidades del personal involucrado con sistemas de IA (I):

- Las entidades de seguros deben **documentar sus actividades de IA en políticas**, ya sea en una política de IA independiente o mediante la actualización de otras políticas.
- Estas políticas **deben establecer los procesos, roles y responsabilidades del personal involucrado** con sistemas de IA.
- Las políticas **deberán revisarse y actualizarse periódicamente**, así como si hay cambios materiales en el uso de la IA en la entidad.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Supervisión Humana (IV)

❑ Responsabilidades del personal involucrado con sistemas de IA (II):

Rol	Responsabilidad
Consejo de Administración	- Deberían estar adecuadamente informados sobre el uso de la IA en su entidad y los potenciales riesgos asociados. - Responsables en última instancia del uso de la IA. - El marco de gestión del riesgo aprobada por el Consejo debe tener en cuenta las actividades de IA. - El Consejo debe establecer procedimientos de comunicación interna sobre la estrategia de IA.
Cumplimiento	- Debe monitorear el catálogo de herramientas implementadas y asegurarse de que cumplan con el ordenamiento jurídico vigente. - Las políticas deben actualizarse y comunicarse.
Gestión de riesgos	- Las entidades deben revisar sus sistemas de gestión de riesgos en relación a la IA. - Los controles sobre los datos internos y externos deberán mejorarse para garantizar que son aptos para el propósito buscado y están libres de sesgos.
Auditoría interna	Una parte clave del marco de gobernanza es la gestión de riesgos y auditoría interna. La auditabilidad permite evaluar tanto la calidad como la eficacia de los algoritmos, así como la eficacia del sistema de gobierno.
Actuarial	- Debe considerar el impacto que tienen las actividades de IA en la función actuarial o en áreas en las que la función actuarial se apoya. - Debe demostrar suficiente comprensión de las técnicas de IA empleadas en su función o que afectan a las actividades actuariales para evaluar los riesgos potenciales, incluidas las cuestiones éticas, y poder tomar decisiones informadas.
DPO	- Las entidades deben conocer sus obligaciones en materia de protección de datos que se aplican a los sistemas de IA. - Las consideraciones del RGPD deben integrarse en el diseño, construcción, prueba e implementación de una aplicación de IA.
Comités	Buena práctica: Establecer comités multidisciplinarios compuestos por expertos multifuncionales que abarcan la función actuarial, riesgo, cumplimiento, protección de datos, IT para supervisar el uso de la IA.
Responsable de IA	- Buena práctica: Nombrar un responsable de IA para garantizar la consistencia y coherencia del trabajo realizado por la entidad y los equipos de IT (aplicando el principio de proporcionalidad). El responsable deberá tener experiencia suficiente para proporcionar supervisión y asesoramiento a todas las funciones de la entidad. - Su principal función será garantizar que cada equipo conoce las políticas y marco de gobierno de IA. También puede coordinar todas las actividades de IA en la entidad, actuando como punto central de experiencia en este ámbito.
Desarrollador sistemas IA	- Los científicos de datos, actuarios y los responsables de la implementación de sistemas de IA deben tener conocimiento específicos y recibir formación periódica sobre cómo implementar sistemas de IA éticos y responsables. - La formación debe promover la noción de que el uso de cualquier solución de IA debe adherirse a los principios éticos fundamentales de equidad, transparencia, explicabilidad, responsabilidad, solidez y precisión, y respeto por la autonomía humana. - Estos equipos deben mantener un catálogo de herramientas de IA para permitir a otras funciones la verificación del cumplimiento de las políticas y el marco de gobierno. Este catálogo debe estar disponible para la alta dirección y funciones clave.
Responsable de seguridad (IT)	- Se deben establecer unas medidas para proteger la seguridad de los datos y de los modelos de IA. - Los responsables de la seguridad deben garantizar que la seguridad del modelo de datos está integrada en el marco de gestión de la seguridad IT.
Usuario final	Juegan un papel muy importante durante todas las fases de desarrollo de sistemas de IA éticos y responsables.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Supervisión Humana (V)

❑ Caso de uso: Tarificación y suscripción:

- **La supervisión humana debe considerarse en el diseño, desarrollo, calibración y pruebas de los sistemas de IA** relacionados con la tarificación. **Muy importante en el caso de sistemas de IA “caja negra”** que afectan a la transparencia y explicabilidad.
- El uso de la IA en la tarificación supone afrontar **desafíos**: garantizar la no discriminación, legalidad y aspectos éticos como la equidad o un adecuado nivel de explicabilidad.
- Los principios son los mismos, pero la complejidad y la interacción con el consumidor cambian significativamente (por ejemplo, dispositivos telemáticos en los automóviles) → **Necesidad de implementar nuevos procesos para gestionar adecuadamente los desafíos**:
 - **Análisis** más frecuente de los **comentarios** y **sentimientos** de los **consumidores** (por ejemplo, en las RRSS), incluso mediante el uso de encuestas sobre la satisfacción de los clientes en su interacción con los sistemas de IA.
 - **Análisis de las salidas del sistema de IA por expertos** que contrasten que el sistema está arrojando las salidas esperadas.
 - **Gestión del funcionamiento del modelo**: Reducir el potencial impacto en el caso de que no funcione según lo previsto o incurra en discriminación.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Supervisión Humana (VI)

❑ Caso de uso: Gestión de siniestros:

- Los sistemas automatizados de gestión de siniestros tienen la ventaja de generar diversos **beneficios a las entidades**:
 - Reducción de costes.
 - Reducción de tiempo de gestión y resolución de siniestros.
 - Estimación de costes más precisa.
 - Mayores ratios de retención de clientes y mejora de la experiencia del cliente.
- Pero, también hay **riesgos**:
 - Rechazo de un siniestro legítimo.
 - Estimaciones de costes de reparación no precisos, liquidaciones injustas de siniestros, etc.
- **Medidas para mitigar posibles daños a los clientes**:
 - **Diseñar procesos justos e involucrar a expertos** en el desarrollo, implementación y test de sistemas de IA.
 - **Definir controles que determinen la necesidad de supervisión humana** en base a un potencial impacto sobre los clientes (por ejemplo, intervención humana cuando la liquidación supere un cierto importe).
 - **Supervisión humana para hacer frente al fraude** teniendo en cuenta el impacto significativo en los consumidores si se sospecha erróneamente de fraude.

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

5.1. Principio de proporcionalidad

5.2. Principio de Trato Justo y No Discriminación

5.3. Principio de Transparencia y Explicabilidad

5.4. Principio de Supervisión Humana

5.5. Principio de Gobierno del Dato

5.6. Principio de Solidez y Rendimiento

Anexos

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Gobierno del Dato (I)

❑ El Principio:

- Las disposiciones nacionales y las incluidas en el RGPD deben ser la base de la implementación de un sistema de gobierno del dato a lo largo del ciclo de vida de los sistemas de IA.
- Las entidades deben **garantizar** que los **datos empleados** en los **sistemas de IA** son **precisos, completos y apropiados** y que aplicarán los sistemas de gobierno con independencia de la fuente de los datos.
- Los **datos** deben ser **guardados en entornos seguros**, sobre todo, en los casos de alto impacto, **garantizando la trazabilidad y auditabilidad de los mismos**.

❑ Regulación:

- **Principios de precisión de los datos y limitación del propósito** (artículo 5 del RGPD): Cualquier uso posterior de los datos se ajuste a su propósito original.
- **Principio de minimización de datos**: Las entidades deben recopilar y tratar los datos personales que sean estrictamente necesarios.
- **Registro de actividades**: El artículo 30 del RGPD indica que las entidades deben mantener registro de todas las actividades de procesamiento de datos personales.
- **El artículo 82 de la Directiva de Solvencia II** (art. 52 del RDOSSEAR), sobre el cálculo de las PPTT, establece una serie de **requisitos de calidad de los datos**.
- **El Reglamento Delegado 2015/35 incluye requisitos de calidad de los datos** en relación al cálculo del CSO mediante modelos internos (artículos 121.3 y 231) y para el cálculo de provisiones técnicas (artículo 19).

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Gobierno del Dato (II)

- ❑ En el mundo digital, existen más fuentes de datos (RRSS, IoT, imágenes, videos, etc.) que pueden ser procesados por sistemas de IA complejos, **brindando oportunidades y retos**.
- ❑ Sistema de gobierno a lo largo del ciclo de vida del dato en los sistemas de IA:

Recopilación de datos	Preparación y transformación de datos	Entrenamiento y prueba del sistema de IA	Post-Procesamiento
▪ Selección adecuada de fuentes y tipos de datos para el sistema de IA. ▪ Considerar la diversidad de los datos posibles (consumidores no digitales). ▪ Los intermediarios pueden jugar un papel importante en la prevención del uso de datos de mala calidad ya que, en sus actividades, pueden detectar malos usos de datos en todas las áreas de la cadena de valor en las que están involucrados.	▪ Entendimiento de la naturaleza, características y formato de los datos. ▪ Limpieza de los datos (duplicados, incompletos). ▪ Asegurar que los datos son precisos, completos y apropiados antes de su uso por sistemas de IA. ▪ Trazabilidad de la transformación de los datos.	▪ Evitar los sesgos de la información de los datos de entrenamiento de algoritmos. ▪ Garantizar que los datos no se contaminen como resultado de errores. ▪ Garantizar la privacidad de los datos de acuerdo con el RGPD, así como los derechos ARCO.	▪ Evaluación del resultado del sistema de IA desde una perspectiva de calidad de los datos (incluida la búsqueda de posibles sesgos discriminatorios).

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Gobierno del Dato (III)

- ❑ **Registro:** Teniendo en cuenta la dificultad de entender el funcionamiento de los modelos de IA, las entidades deben guardar **registro de las siguientes informaciones:**

Registro	Descripción
Razones del uso de la IA	Explicación del objetivo de negocio perseguido con el uso de la IA y sus consistencia con los objetivos y estrategia de la entidad . Explicación de cómo fue implementado en los sistemas de IA.
Integración con la infraestructura IT	Descripción de cómo los modelos de IA son integrados en los sistemas de IT de la entidad , así como de los cambios.
Personal involucrado en el diseño y desarrollo de sistemas de IA	Identificación de todos los roles y responsabilidades del personal involucrado en el diseño e implementación de modelos de IA , así como de las necesidades requeridas para su entrenamiento.
Recopilación de datos	Documentación de la metodología implementada para la identificación y eliminación de posibles sesgos en los datos , explicando cómo se seleccionaron, recopilaron y etiquetaron los datos.
Preparación de datos	Registro de los datos usados en los modelos de IA.
Post-Procesamiento de datos	Descripción de los procesos seguidos para la mejora continua de los datos empleados en los modelos de IA.
Alternativas técnicas	Documentación de los algoritmos empleados en los modelos de IA, librerías y referencias utilizadas. Las limitaciones de los algoritmos deben ser documentadas.
Código	Registro del código fuente de los modelos de IA. En las aplicaciones de alto impacto se deben registrar los datos de entrenamiento y los parámetros establecidos (en aplicación del Pº de proporcionalidad, las entidades deben implementar alternativas que garanticen la auditabilidad del modelo de IA).
Rendimiento	Documentación de los KPIs , incluyendo frecuencia de revisión.
Seguridad	Descripción de los mecanismos establecidos para garantizar la seguridad frente a ataques e intentos de manipular datos y algoritmos.
Ética	Evaluación del impacto del caso de uso de IA. Documentación de las medidas de gobernanza implementadas para asegurar la ética y no discriminación de los sistemas de IA.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Gobierno del Dato (IV)

❑ Caso de uso: Tarificación y suscripción:

- **Mejorar las evaluaciones de riesgos mediante la combinación de fuentes de datos nuevas (IoT) y tradicionales** conduce a cambios en la fase de preparación de datos. Por ello, **es necesario**:
 - Garantizar la calidad y la imparcialidad de los datos dentro del nuevo conjunto de datos combinado.
 - Volver a entrenar y ajustar el algoritmo de IA.
 - Actualizar los registros con la nueva recopilación de datos.
 - Actualizar los registros de la fase de preparación y pos-procesamiento de los datos.
 - Actualizar el impacto de la nueva fuente de datos y el rendimiento del modelo.
- Las prácticas de **optimización de precios basadas en datos de comportamiento individuales** crean la necesidad de un ciclo de vida del modelo de IA más dinámico y flexible (complejo):
 - La calidad de los datos y la imparcialidad son cruciales.
 - Los datos de comportamiento individuales exigen **re-entrenamiento y re-evaluación con mayor frecuencia**, así como la **reconciliación de datos** (complejidad adicional en el sistema de IA).
 - **Una nueva iteración con datos nuevos de comportamiento requiere**:
 - Actualizar los registros de recopilación, preparación y post-procesamiento de datos.
 - Actualizar los los registros de decisiones técnicas, rendimiento del modelo y evaluación de impacto.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Gobierno del Dato (V)

❑ Caso de uso: Gestión de siniestros - Reconocimiento de imágenes:

○ Fase de recopilación de datos:

- **Necesario especificar los requisitos con respecto a la imagen.** A diferencia de los datos, una imagen puede ser de mala calidad, estar enmarcada o desplazada y, posteriormente, ser interpretada incorrectamente por el sistema de IA. Por lo tanto, el asegurado debe cumplir con estos requisitos para garantizar una alta calidad de los datos y la imparcialidad.

○ Fase de preparación de datos: Se pueden realizar comprobaciones adicionales para que los datos, y especialmente las imágenes, se ajusten a los requisitos exigidos.

○ En caso de que el proceso esté completamente automatizado, se deben mantener registros de:

- Las razones para usar el sistema de IA.
- Integración en la infraestructura de IT de la entidad.
- Personal involucrado en el diseño e implementación del modelo de IA.
- Recopilación de datos, preparación y procesamiento posterior de datos.
- Decisiones técnicas.
- Código fuente, rendimiento del modelo, seguridad del modelo y evaluación de impacto.

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

5.1. Principio de proporcionalidad

5.2. Principio de Trato Justo y No Discriminación

5.3. Principio de Transparencia y Explicabilidad

5.4. Principio de Supervisión Humana

5.5. Principio de Gobierno del Dato

5.6. Principio de Solidez y Rendimiento

Anexos

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Solidez y Rendimiento (I)

❑ Principio:

- Las entidades de seguros deben **utilizar sistemas de IA robustos (interna y externamente)**, teniendo en cuenta su uso previsto y el potencial daño que puedan causar.
- **Los sistemas de IA deben ser adecuados para su propósito y su desempeño debe ser evaluado y supervisado de manera continua**, incluido el desarrollo de métricas de desempeño relevantes. Es importante que la calibración, validación y reproducibilidad de los sistemas de IA se realice de una manera sólida que garantice que los resultados de los sistemas de IA sean estables en el tiempo.
- **Los sistemas de IA deben implementarse en infraestructuras de IT seguras y resistentes** a ataques cibernéticos.

❑ Regulación:

- El artículo 25 del RGPD incluye el requisito de que las entidades garanticen la **protección de datos** por diseño y por defecto, lo que significa que las entidades de seguros deben integrar la protección de datos en todas las actividades de procesamiento de datos y prácticas comerciales en todo el ciclo de vida de los sistemas IA.
- Solvencia II también incluye las disposiciones relevantes relacionadas con la solidez, desempeño y buen funcionamiento en el contexto de los modelos internos y el cálculo de las provisiones técnicas en seguros.
- La Asociación Internacional de Actuarios también ha desarrollado Estándares Internacionales para la Práctica Actuarial para asegurar la robustez y precisión de los modelos.
- El Reglamento sobre Resiliencia Operativa Digital (DORA) introducirá nuevos requisitos de seguridad de IT y resiliencia operativa.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Solidez y Rendimiento (II)

❑ Requisitos de solidez y rendimiento en seguros (I):

Requisito	Descripción
Sistemas de IA apropiados para su uso previsto	- Definir claramente cuál es el objetivo que pretenden lograr con un sistema de IA específico. - Definir las métricas de rendimiento adecuadas para ese sistema de IA específico Garantizar que el modelo no se utilice para un propósito diferente al que fue creado. Si el alcance o la naturaleza del propósito cambia, también debe hacerlo el sistema de IA. - Para ciertas aplicaciones de IA de alto impacto, las entidades deberían explicar suficientemente el proceso de toma de decisiones del modelo de IA.
Evaluación y seguimiento de los sistemas de IA	Comprender la naturaleza, escala y probabilidad de error de los sistemas de IA. Evaluar si el riesgo de incorrecto funcionamiento de los sistemas de IA se encuentra dentro de un apetito de riesgo definido por la entidad. Conocer las limitaciones del sistema de IA y en qué circunstancias se pueden producir problemas. Esto puede requerir restricciones en el uso del sistema para garantizar que se implemente sólo en circunstancias que conduzcan a los niveles deseados de rendimiento o que minimicen el bajo rendimiento.
Métricas de rendimiento (exactitud, recuperación, precisión, etc.)	- Por ejemplo, en los sistemas de clasificación de IA utilizados en la detección de fraudes, las entidades de seguros deben decidir si el objetivo es maximizar la precisión de la predicción (número de siniestros fraudulentos detectados), reducir el número de falsos positivos (siniestros legítimos etiquetados erróneamente como fraudulentos) o falsos negativos (siniestros etiquetados como legítimos que al final son fraudulentos). Monitorear el desempeño de un sistema de IA con respecto a los consumidores vulnerables y desarrollar métricas apropiadas para este propósito. Documentar la selección de métricas de desempeño junto con su justificación.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Solidez y Rendimiento (III)

❑ Requisitos de solidez y rendimiento en seguros (II):

Requisito	Descripción
Gestión de datos	• Evaluar la integridad de los datos y adecuación para el propósito, considerando si son precisos, completos, imparciales y representativos. • Evaluar la integridad y adecuación de los datos cada vez que se desarrolle un nuevo sistema de IA, justificando por qué los datos utilizados se consideran adecuados. • Definir, ejecutar y actualizar periódicamente controles de calidad e integridad de los datos (dado que contribuyen en gran medida al rendimiento del modelo). • Evaluar la consistencia entre el uso de IA y el uso que se consideró cuando se desarrolló el sistema teniendo en cuenta las limitaciones del modelo y el contexto en el que se pretendía utilizar el modelo de IA.
Calibración y validación del sistema de IA	• Re-entrenar, re-calibrar y re-evaluar periódicamente los sistemas de IA (particularmente en el caso de cambios significativos en los datos de entrada, factores externos relevantes y/o en el entorno legal o económico). • Documentar adecuadamente los criterios para un cambio de modelo significativo para cada aplicación de IA. Estos criterios deben ser más estrictos a medida que aumente la materialidad del sistema de IA. • Establecer procedimientos de evaluación suficientes: • El proceso de evaluación de los sistemas de IA de alto impacto debe incluir el uso de análisis de escenarios y pruebas de estrés. • Las entidades deben determinar la frecuencia mínima para las evaluaciones de cada sistema de IA. • Deben existir controles adecuados para evitar que el sistema de IA ponga un énfasis excesivo en los patrones a C/P que afecte al rendimiento del modelo a L/P.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Solidez y Rendimiento (IV)

❑ Requisitos de solidez y rendimiento en seguros (III):

Requisito	Descripción
Estabilidad de resultados a lo largo del tiempo	• Prestar atención a los datos utilizados para el entrenamiento del modelo de IA, para garantizar que las diferentes iteraciones o subconjuntos de dichos datos de entrenamiento no den resultados significativamente diferentes. • Documentar y justificar cualquier acción tomada a este respecto durante el entrenamiento y prueba del sistema. • Evaluar la solidez del sistema mediante métricas estadísticas que miden las diferencias entre los resultados del modelo. Estas métricas deben seleccionarse teniendo en cuenta el uso previsto del sistema y los resultados deseados. • Evaluar la solidez mediante la redundancia de sistemas. Un sistema puede considerarse robusto cuando su resultado es similar al de otros sistemas independientes.
Sistemas e infraestructuras de IT resilientes	• Los sistemas de IA son algorítmicos y, por lo tanto, requieren más potencia de cálculo (por ejemplo, utilizando una infraestructura en la nube). Las entidades de seguros deberían considerar la posibilidad de actualizar su infraestructura de IT para abordar las posibles limitaciones de software y hardware al implementar soluciones de IA. • Implementar medidas de seguridad resistentes para responder a los ciberataques. Los sistemas de IA implementados en producción para uso automático o con supervisión humana limitada pueden causar un daño significativo antes de que esto se note y aborde. Los pueden afectar a la confianza depositada en los sistemas de IA ciberataques (por ejemplo, entradas de datos diseñadas intencionadamente para hacer que el modelo cometa un error) .

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Solidez y Rendimiento (V)

❑ Requisitos de solidez y rendimiento en seguros (V):

Requisito	Descripción
Planes de “marcha atrás”	• Establecer planes alternativos en caso de que los sistemas de IA no funcionen según lo previsto o sean víctimas de un ciberataque.
Sistemas de IA de terceros	• Los proveedores de servicios se asegurarán de que las aplicaciones de IA que comercialicen cuenten con los más altos estándares de calidad y proporcionen a las entidades de seguros información suficiente para permitirles tener un conocimiento adecuado de su funcionamiento y las limitaciones de los sistemas de IA. • Las entidades de seguros también deben tener en cuenta la concentración potencial de riesgos derivados de una alta dependencia de un proveedor de servicios específico, siendo en estos casos importante mantener y probar estrategias de salida para soluciones subcontratadas, en particular cuando son materiales.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Solidez y Rendimiento (VI)

❑ Caso de uso: Suscripción y tarificación (I):

- Los riesgos sujetos a una **gran incertidumbre** generalmente darán lugar a **modelos menos precisos** y hay casos en los que las entidades de seguros aceptarán un rendimiento del modelo que no será suficiente al modelar los riesgos sujetos a menos incertidumbre:
 - **Se debe justificar por qué el desempeño alcanzado se considera aceptable** para informar a los *stakeholders* sobre los desafíos del modelo y la incertidumbre subyacente del riesgo.
 - **Se recomienda identificar las principales fuentes de incertidumbre** a fin de evaluar qué acciones (si las hay) pueden tomarse para reducirla.
- **La falta de datos o no tener acceso a datos relevantes también puede contribuir a la falta de rendimiento del modelo:**
 - Ejemplo: Fijación de primas de una nueva línea de negocio para la que no hay experiencia previa.
 - **Se requiere el juicio de expertos para proporcionar una indicación de cómo las deficiencias de los datos afectan al rendimiento del modelo** (aumento de la volatilidad, probable subestimación o sobreestimación del coste del riesgo para una cohorte específica, etc.) de cara a mitigar o buscar alternativas viables ante estas situaciones.
 - **Cualquier acción tomada para reducir el impacto de las deficiencias de datos debe identificarse y divulgarse** adecuadamente.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Solidez y Rendimiento (VII)

❑ Caso de uso: Suscripción y tarificación (II):

- Se debe **prestar especial atención a las consecuencias del desempeño inadecuado del modelo** y se debe proporcionar información a los *stakeholders* que les permita comprender las consecuencias para la entidad y sus consumidores.
- **Establecimiento de umbrales de rendimiento:**
 - El **umbral para lo que se considera un rendimiento adecuado del modelo** debe basarse en las posibles consecuencias de que el modelo no funcione como se esperaba.
 - Por ejemplo, una nueva línea de negocio con crecimiento controlado y que representa una pequeña parte de la cartera de una entidad es más probable que tolere un rendimiento del modelo más bajo que una que represente la mayor parte del negocio de una entidad donde se producirían consecuencias importantes en la solvencia y la estabilidad financiera de la entidad.
- **Medidas para garantizar la robustez de los modelos de tarificación:**
 - **Deben entenderse y explicarse claramente la existencia de diferentes primas para riesgos aparentemente similares.**
 - **Las diferencias significativas en las primas de un período respecto al siguiente deben tener una justificación** bien documentada en lugar de atribuirse genéricamente a cambios en los datos.
 - Las entidades de seguros deben considerar, dentro de la gestión de la continuidad del negocio, **qué alternativas deben aplicarse en el caso de que se sospeche que sus sistemas de tarificación presentan bajo rendimiento, discriminación, etc.** y que requieran que su sistema actual sea retirado de la producción en un plazo breve.

Principios de gobierno para una IA ética y responsable en el sector asegurador

Principio de Solidez y Rendimiento (VIII)

❑ Caso de uso: Gestión de siniestros - Técnicas de OCR y de procesamiento de imágenes:

- Generalmente, **brindan tasas de precisión extremadamente altas en muchos casos, aunque pueden presentar una cantidad significativa de errores** cuando se aplican a documentos que incluyen muchas palabras.
- Esto **podría ser aceptable si una entidad de seguros utiliza esta información para complementar los datos existentes de los clientes o cuando existe una supervisión humana adecuada.**
- **Medidas adicionales para garantizar un rendimiento adecuado:**
 - **Sistemas de apoyo:** Uso de múltiples sistemas que ayuden a superar las debilidades de los demás.
 - **Proceso automatizado para pedir a los asegurados confirmación o aclaración** cuando se ha demostrado que el sistema arroja mayores tasas de error, etc.
 - **Podrían usarse para la clasificación en lugar de para cuantificar la compensación real del siniestro:** Dado que este proceso ya incorporaría la supervisión humana, los requisitos de solidez podrían ser menores.
 - **La calidad de la imagen o del documento pueden afectar significativamente el desempeño del modelo:** Las entidades deben ser conscientes de esta deficiencia y tenerla en cuenta en sus procesos. Esto podría resultar en una **combinación de a) solicitar datos de mejor calidad a peritos y asegurados, b) pasar la gestión del siniestro a expertos para un análisis en detalle.**

FIN

¡¡¡MUCHAS GRACIAS!!!

Índice

1. Introducción

2. La Inteligencia Artificial en la cadena de valor del seguro

3. Aplicación de la Inteligencia Artificial en el seguro

4. Iniciativas reguladoras de la Inteligencia Artificial

5. Principios de gobierno para una Inteligencia Artificial ética y responsable en el sector asegurador

Anexos

Anexos

Características protegidas

- ❑ El artículo 21 de la Carta de Derechos Fundamentales de la UE establece las siguientes características protegidas en materia de no discriminación:

Características protegidas

- Sexo.
- Raza.
- Color.
- Origen étnico o social.
- Lengua.
- Religión o convicciones.
- opiniones políticas o de otra índole.
- Pertenencia a una minoría.
- Patrimonio.
- Nacimiento.
- Discapacidad.
- Edad.
- Orientación sexual.