

El Reglamento Europeo de Inteligencia Artificial: implicaciones prácticas para las organizaciones

58 MIN · YOUTUBE · [HTTPS://WWW.YOUTUBE.COM/WATCH?V=U8IPK0XMAZI](https://www.youtube.com/watch?v=U8IPK0XMAZI)
<https://www.youtube.com/watch?v=U8IPK0XMAZI>

SUMMARY

La sesión del Foro de Inteligencia Artificial se centró en las implicaciones del Reglamento de Inteligencia Artificial de la Unión Europea, destacando la transición de la IA de un tema de innovación a uno de gestión y responsabilidad en las organizaciones. Cristina Mesa, abogada experta, explicó los riesgos y obligaciones que las empresas deben considerar al implementar IA, así como la importancia de la gobernanza y la regulación.

- La IA ha pasado de ser un tema de innovación a uno de gestión y responsabilidad en las organizaciones.*
- El Reglamento de IA de la UE entra en vigor en agosto de 2024, con un período de adaptación para las empresas.*
- Un 93% de las empresas españolas están en fases incipientes de implementación de IA, lo que indica un uso personal más alto que empresarial.*
- La gobernanza de la IA es crucial para mitigar riesgos, especialmente en sistemas de alto riesgo que requieren certificación.*
- Las empresas deben ser conscientes de los usos de IA que pueden clasificarse como de alto riesgo, lo que implica mayores obligaciones regulatorias.*
- La cadena de valor de IA determina las responsabilidades y obligaciones de los diferentes actores (proveedores, importadores, usuarios).*
- Las sanciones por incumplimiento del reglamento pueden ser significativas, alcanzando hasta 35 millones de euros o el 7% de la facturación anual.*
- La colaboración interdepartamental es esencial para implementar la IA de manera segura y efectiva en las organizaciones.*

Speaker 1: Así que, bueno, pues bienvenidos a esta sesión dentro del Foro de IA, del Foro de Inteligencia Artificial del Club. Ya sabéis que es un foro, o bueno, muchos de los que estáis conectados, algunos no, porque algunos os incorporáis a esta sesión, pero algunos sabéis que es un foro que llevamos trabajando con él un par de años, ha tenido distintos momentos y arrancamos el primer año con muchas acciones de formación, sensibilización sobre Inteligencia Artificial.

Speaker 1: El año pasado entramos más en compartir algunas prácticas y al final lo que estamos viendo es que la Inteligencia Artificial ha pasado muy poco tiempo de ser un tema de innovación a convertirse en un tema de gestión, de estrategia, cada vez más de responsabilidad y que tenemos más integrado en en nuestro día a día, en nuestras organizaciones. Por tanto, hemos ido evolucionando el Foro, el Foro de Inteligencia Artificial o el Foro

de IA y estamos replanteando la estrategia de cara a este año.

Speaker 1: Hablaremos con vosotros, con los distintos miembros, para ver cómo planteamos la estrategia este año, en qué tiene sentido centrarnos, qué tenemos que priorizar, si tenemos solo que centrarnos en Inteligencia Artificial o hacer un foro más amplio. Bueno, estamos en ese punto de replantearlo, pero en cualquier caso, arrancamos este año con esta sesión, que era un tema identificado por los miembros del Foro de IA. El año pasado, en 2025, muchas organizaciones estáis ya, o estamos ya utilizando la IA y a veces lo hacemos de forma muy estructurada, otras veces de una forma más espontánea, pero la realidad es que no siempre somos conscientes de las implicaciones que esto tiene desde el punto de vista regulatorio y de gobernanza.

Speaker 1: Es algo que ha salido en todas las reuniones que hemos ido haciendo. Precisamente por eso hemos querido dedicar esta sección del Foro de IA a entender las implicaciones prácticas del Reglamento de Inteligencia Artificial, qué cambia realmente, dónde están los principales riesgos y qué decisiones conviene que tomemos desde las organizaciones desde hoy, desde el día de hoy. Y para ello, y yo he de decir que a Cristina no la conocía y fue uno de vosotros el que nos sugirió que contactásemos con ella, contamos con Cristina Mesa, abogada en Garrigues y una de las especialistas de referencia en ese ámbito en el Reglamento Europeo, que está asesorando a empresas, administraciones en cómo traducir el Reglamento en decisiones concretas de su día a día, de su gestión, de su contratación y de la gobernanza.

Speaker 1: Así que contamos hoy con Cristina. Muchas gracias, Cristina, por acompañarnos. Voy yo a dejar de compartir para que vayáis compartiendo.

Speaker 2: Pues muchas gracias, Susana. En primer lugar agradecer al club y también a Susana y a Denise, que me han ayudado un montón para preparar la presentación, la invitación. No sé quién ha sido el que os ha condenado a esta hora de presentación mía, pero bueno, muchas gracias. Y como comentaba Susana, pues bueno, el esfuerzo va a ser para un acercamiento generalista al reglamento de Inteligencia Artificial, entre otras cosas porque todavía falta un poquito para su aplicación global, por decirlo de alguna forma, pero estamos en un buen momento para tomar las decisiones sobre cómo se van a subir los riesgos sobre la implementación del reglamento y en general, incluso en ausencia de reglamento, cómo podemos mitigar los riesgos de la introducción del reglamento de inteligencia artificial en la empresa.

Speaker 2: Me he intentado enfocar en pequeña y mediana empresa, pero quería comenzar por un dato que creo que es especialmente llamativo. Es verdad que a los que nos dedicamos a estos temas siempre nos deprime un poco el hecho de que casi todas las empresas relacionadas con temas de inteligencia artificial son estadounidenses o chinas. Parece como que la Unión Europea y España no tienen mucho que decir aquí. Pero bueno, en el último estudio que lanzó Microsoft sobre la adopción de la inteligencia artificial en el mundo, pues salió a la luz un dato realmente sorprendente, que son las estadísticas de uso de la inteligencia artificial.

Speaker 2: Para nuestra sorpresa y creo que orgullo, no sé si estaréis de acuerdo, pues estamos en el top 6 entre los 6 países del mundo que más inteligencia artificial utilizan. Entonces, bueno, ser el creador de la tecnología está muy bien, pero ser los que la utilizan correctamente o intensivamente, pues también. Ahora bien, no sé si os

llama la atención como me ha llamado a mí ver cuál es el ranking de los países que más la están utilizando, porque me parece que ahora que vivimos momentos extraños de geopolítica, pues estas estadísticas también nos vienen a confirmar que la innovación no siempre está donde pensamos que está o no siempre está tan generalizada.

Speaker 2: Porque, por ejemplo, a mí sí me sorprende que en esta tabla no encontremos a Estados Unidos, siendo que las herramientas que utilizamos casi todos nosotros vienen de Estados Unidos. La de uso más extendido es GPT, luego después está Gemini, las dos son estadounidenses y sin embargo, el número de usuarios no los tenemos aquí. Entonces, pues bueno, el hecho de que se cree un tejido de uso también es importante. Tampoco al principio sí me sorprendió, pero luego ya no me sorprende tanto que los Emiratos Árabes estén arriba, que Singapur esté arriba y que en general los países europeos estén muy arriba y estén incluidos casi todos en los puestos superiores de adopción de la Inteligencia Artificial.

Speaker 2: Pero ahora bien, como he hecho una promesa, diré de propósito de año nuevo, que es intentar no buscar datos para lo que a mí me gustaría pensar, sino datos objetivos. Estos datos que parece que nos llevan a un sitio estupendo, se está haciendo todo muy bien, estamos implementando la IA, pues tienen una pequeña puerta de atrás. Estos son datos de uso, pero si pasamos al uso en empresa, los datos caen mucho. Es decir, los individuos estamos usando mucho la Inteligencia Artificial en Europa, en Oriente Medio, etc.

Speaker 2: Pero luego las empresas estamos teniendo un poquito más de dificultad a la hora de incorporar. Y este dato creo que es importante que lo tengamos todos en cuenta. En el informe que publicaba Deloitte, cuando se habla de uso empresarial, pues ya bajamos mucho en el ranking y además bajamos mucho en preparación. Perdona, bajamos mucho en preparación. Entonces, lo que vienen a decir de forma muy suave, es que el 93% de las empresas españolas todavía están en una fase muy incipiente de implementación de la IA.

Speaker 2: ¿A dónde me lleva esto a mí de forma práctica? A que una separación tan grande entre el uso de la IA por las personas, por nosotros, por nuestros empleados, por nuestros hijos, y el uso todavía escaso de la IA en la empresa, me lleva a pensar que, y he intentado hacerlo en una ilustración, que a lo mejor no lo sabemos, pero que nuestros empleados han decidido por su propia cuenta de riesgo incorporar la Inteligencia Artificial en nuestras empresas.

Speaker 2: Nosotros utilizamos, tengo un compañero que lo dice mucho, como dicen los cursis, un término que es el de Bring your own device. Es decir, si la empresa no te da una herramienta que tú consideras útil, pues a lo mejor el empleado utiliza su teléfono móvil, o utiliza su iPad, o utiliza su propio ordenador, o simplemente accede mediante las herramientas de la empresa a sistemas de Inteligencia Artificial. Probablemente esto les hace mucho más productivos. Y es bueno para la empresa.

Speaker 2: Sí, pero también introduce riesgos. Es lo que llamamos Inteligencia artificial en la sombra. Es decir, cuando la empresa no sabe que sus empleados están utilizando la IA, ni qué IA están utilizando sus empleados es materialmente imposible que pueda mitigar los riesgos, más que nada porque los desconoce. Entonces, este es uno de los puntos que quería traer a colación, porque me ha pasado mucho con determinadas empresas. No, no,

nosotros esto lo tenemos prohibido. Luego hablar con. Incluso hablo de asesorías jurídicas a las que nosotros estamos ayudando a implementar y que son nuestros clientes y que te dijese, bueno, prohibido sí, pero yo sí miro cosas.

Speaker 2: Entonces, claro, hay que ver qué cosas está mirando la gente y cómo lo están haciendo. ¿Y por qué es importante ahora? Porque como os decía, pese a que todavía, y seguro que si leéis las noticias habéis visto ahora con el tema de la ómnibus, que se está intentando simplificar el marco contractual para las empresas, y especialmente yo me dedico al digital y reconozco que es muy difícil gobernar toda la normativa que llega en materia digital. Entonces, pues claro, para una pequeña, mediana empresa, incluso para las grandes, se convierte en un trabajo titánico.

Speaker 2: En el caso del reglamento, pues este reglamento está aprobado, entró en vigor en agosto de 2024, pero como pasa últimamente, como son reglamentos tan complicados y requieren de la colaboración transversal de todos los departamentos de la empresa, se da un tiempo de adaptación relativamente generoso para poder conocer cuáles son las normas, qué pasos tenemos que seguir y para que no muramos todo un impacto en un día desde que se publica la norma hasta que tenemos que cumplir.

Speaker 2: De hecho, hay una cosa que a mí. Yo estoy dentro de la sedia como experta en el sandbox del Reglamento de Inteligencia Artificial, y hay una discusión que seguimos manteniendo. Como veis, en agosto de 2026, como podéis ver en pantalla, ya es de plena aplicación el reglamento, incluso los que se llaman, y ahora hablaremos un poquito de esos sistemas de alto riesgo, muchos de esos sistemas de alto riesgo, para poder funcionar, es decir, para que los podamos utilizar, necesitan una certificación.

Speaker 2: Y lo cierto es que a día de hoy no existe ningún organismo en el territorio de la Unión Europea que pueda emitir esas certificaciones. Entonces, somos muchos y muchas pymes las que se empiezan a angustiar porque, claro, no se ven capaces de cumplir con. Bueno, no es que no se vean capaces, es que es materialmente imposible tener un certificado que todavía nadie está en posición de dar. Entonces, pues bueno, ya se empieza a rumorear que va a haber extensiones porque la propia Unión Europea no está en posición de cumplir con su rol de certificador de sistemas de inteligencia artificial.

Speaker 2: ¿Por qué levanta tanto revuelo? Aparte de porque la inteligencia artificial es muy chula y siempre confronta al hombre con sus propios límites, el hecho de que tú digas, bueno, esta máquina está trabajando mejor que yo, está haciendo las cosas mejor que yo, en el imaginario colectivo, pues siempre la ciencia ficción también nos ha llevado a tener cierto miedo a la adopción de la inteligencia artificial. Pero esto se ve todavía reforzado si miramos las multas que nos puede imponer el regulador por una implementación incorrecta.

Speaker 2: Porque por los incumplimientos más graves, estamos hablando de multas de 35 millones de euros o el 7% del volumen de facturación anual mundial. Y además lo que sea más grande. Luego la mayoría de incumplimientos tienen un rango un poquito menor. Pero no penséis que estamos hablando de multas insignificantes, porque son multas de entre 15 millones de euros o el 3% del volumen de facturación anual. Siguen siendo multas muy elevadas. Además, para las pymes se hace una pequeña salvedad, pero me temo que

tampoco traigo grandes noticias, porque en este caso lo único es que, a diferencia de lo que pasa con las grandes empresas, en vez de mirar la que sea mayor, se va a mirar la que sea menor, es decir, la multa económica o el volumen de facturación.

Speaker 2: Pero quitando eso, las multas son considerables y además siguen la estela de los regímenes sancionadores en otras áreas, protección de datos, de lo que yo no voy a hablar hoy, o las plataformas de intermediación de Internet, son multas considerables. Dicho eso, voy a intentar aterrizarlo de una forma que no sea, y siempre lo intento, pero que no sea excesivamente legal, porque además es una de esas normas donde sorprendería, pero la gente que está más preocupada por su cumplimiento no son solo los abogados, son los ingenieros, la gente de compliance, la gente de recursos humanos, la gente que quiere.

Speaker 2: De hecho, yo he dado muchísimas charlas sobre el reglamento de Inteligencia Artificial en congresos específicos de recursos humanos, y ahora veremos por qué tiene sentido que ellos estén tan preocupados, que lo tiene. Pero dicho esto, lo relevante en esta materia es la gobernanza, la gobernanza de los sistemas. Como os decía, la IA puede ser muy buena para la empresa, un buen uso de la IA nos puede hacer eficientes. Y de hecho me acabo de dar cuenta que se me ha olvidado meter una slide que me gusta meter mucho, pero que no es original mía porque se la copia un compañero, que era básicamente como cuando se extinguieron los dinosaurios por los meteoritos.

Speaker 2: Hay gente que pensamos que determinadas formas de trabajar ya no van a ser permisibles. Nuestros clientes, los vuestros, piden, cada vez hay más presión con los precios, más presión con la eficiencia, entonces van a dar por hecho que somos capaces de determinar, de realizar determinado trabajo de forma más eficiente, y por eso hay que decir de forma más barata, con menos recursos, y por lo tanto van a asumir que se están implementando herramientas que ya están disponibles en el mercado.

Speaker 2: Entonces cada vez es menos opcional, especialmente para determinados trabajos más repetitivos o que ya se ha demostrado que pueden ganar eficiencias con la Inteligencia Artificial. Y por eso, más que la decisión de si adoptar o no IA, que creo que ya está superada, la decisión se mueve más a cómo introducir la Inteligencia Artificial en la empresa. Porque se me ocurre, y no tengo ningún problema en que me interrumpáis si pensáis distinto, se me ocurren pocos procesos dentro de la empresa donde la IA no pueda eficientar, o a lo mejor no eficientar, pero sí subir la calidad del trabajo.

Speaker 2: Y como digo, eso es lo que nos van a pedir nuestros clientes. Dicho esto, y aquí no me voy a detener mucho porque digo, no quiero enfocarlo mucho solo en la pata legal, sí que es importante que entendamos cómo se está organizando la regulación de la IA. Lo que dice la Comisión Europea, y eso es verdad, es que hay mucha preocupación, pero que lo cierto es que la mayoría de sistemas de Inteligencia Artificial que utilizamos son bastante inocuos, es decir, no entrañan un gran riesgo para los derechos fundamentales de los ciudadanos europeos, ni a la seguridad de productos.

Speaker 2: Y ese es el enfoque que ha dado el reglamento. ¿Como ves esta pirámide? Que seguro que habéis visto alguno. Lo que ha hecho el reglamento es imponer obligaciones, más o menos obligaciones, dependiendo del

nivel de riesgo que pueda suponer el uso de un sistema de Inteligencia Artificial. Y a mí me gusta además, por eso le he dado la vuelta muchas veces, se ve hacia arriba, pero a mí me gusta utilizar esa porque la mayoría, como se estaba diciendo de los usos son de bajo riesgo, y por eso me gusta que estén en la base de la pirámide, la más gruesa, Pues obviamente no es lo mismo un filtro de spam de correo electrónico para detectar si te entra un correo electrónico que es útil o no para tu trabajo, que no tiene básicamente ningún riesgo, que un sistema de inteligencia artificial que te permite hacer un social scoring de una persona, rastrear sus redes sociales, rastrear su vida, rastrear su declaración de impuestos para ver si es o no un buen ciudadano y conforme a eso darle determinados beneficios.

Speaker 2: Son sistemas que, por ejemplo, se están utilizando en China, pero que la Unión Europea ha considerado que no son acordes con los derechos fundamentales de los europeos y considera prohibidos sistemas de reconocimiento parcial, etc. Y luego tenemos otros que aunque son los de alto riesgo, que son en los que verdaderamente, porque sinceramente entre nosotros, los que están prohibidos son relativamente claros, es decir, la Unión Europea ya ha decidido el regulador que hay cuatro o cinco usos que no están permitidos y hay poco que podemos hacer al respecto.

Speaker 2: No están permitidos y punto. Luego entramos en la zona que es más difícil de cumplir, que es la de los sistemas de alto riesgo. Y aquí pueden ser sistemas que son muy valiosos, pero que por la forma en la que afectan a los ciudadanos se establecen muchas obligaciones para su uso. ¿Cuándo puede pasar esto? Como os decía, por ejemplo en el uso de recursos humanos, cuando estamos utilizando un sistema de inteligencia artificial para cribar currículums. ¿Puede tener esto un impacto en derechos fundamentales?

Speaker 2: Sí, lo puede tener si el sistema está. Y además ha habido ya varios casos, incluso antes de que se empezase a discutir sobre el fundamento de IA, en los que se han detectado sesgos y por ejemplo, se detectaba que había sistemas de inteligencia artificial que se utilizaban en recursos humanos, donde el sistema, porque aprendía, ha pasado, perpetuaba sesgos de género y por lo tanto siempre optaba por currículums de hombres frente a currículums puestos donde había candidatas mujeres, porque era lo que se había hecho en el pasado, o también había un sesgo de edad donde directamente se descartaban los currículums de mayores de 45 años.

Speaker 2: Entonces, ¿Tiene esto un impacto en derechos fundamentales? Sí, lo tiene. ¿Significa esto que no se pueden utilizar sistemas de cribado en recursos humanos? No, significa que esos sistemas tienen que someterse a determinados controles para que no pasen justo las cosas que os estaba ocultando. Luego, por ejemplo, un uso de alto riesgo que ha sido muy curioso y que ha sorprendido mucho a la comunidad educativa es el de la comunidad educativa, por ejemplo, corregir exámenes. Tú no puedes permitir que un sistema de inteligencia artificial corrija un examen solo sin supervisión humana, porque claro, esto tiene un impacto en los estudiantes.

Speaker 2: Y por ejemplo, uno que me llamó mucho la atención en cómo se conceden las becas. Si una inteligencia artificial tiene un sesgo en cómo se conceden las becas, pues obviamente esto tiene un impacto en los derechos fundamentales de los estudiantes. De nuevo, ¿Quiere decir que no se pueden utilizar? No, quiere decir que te van a pedir un estándar de cumplimiento muy elevado para saber que tu sistema no está afectando a los derechos fundamentales de los europeos.

Speaker 2: Luego pasamos a lo que estamos utilizando todos los chatbots GPT Gemini que se usan, pues bueno, en ese sexto uso vas a tener desde recetas de cocina, gente que el niño ha tomado hoy esto en el colegio, qué le hago de cena para que esto esté equilibrado, a empresas que los están utilizando para agilizar sus procesos, resumir, por ejemplo, en nuestro caso, leyes, sentencias, escritos, mejorar la calidad, diseñar estrategias, incluso agentes de IA que nos liberan o nos ayudan a liberarnos de procesos de facturación, procesos, el uso es infinito.

Speaker 2: Y luego los que directamente caen fuera del reglamento porque el uso es impacto mínimo, o yo incluso pienso que algunos no se deberían considerar ni inteligencia artificial porque no cumplen con los propios criterios. Para entender que estamos hablando de una inteligencia artificial que necesite regulación, un simple filtro de correo electrónico normalmente no va a entrar. Perdonad que me haya extendido aquí, pero para poder entender cuál es el enfoque de la Unión Europea y cuál es el que se tiene que hacer en la empresa.

Speaker 2: El tema es que este enfoque, pues creo que nos hace entender mejor por qué se tiene que regular la inteligencia artificial. Y como os decía, esos sistemas de alto riesgo en muchas ocasiones van a exigir un marcado cercano. Estáis muy acostumbrados a esto, los coches que llevamos, los juguetes de nuestros hijos tienen un marcado CE, pero yo es que no sé vosotros, yo no me quiero montar en un coche que no cumpla con un mínimo de seguridad y no le quiero dar un juguete a mi hijo que no cumpla con un mínimo de seguridad.

Speaker 2: Esta es la lógica. Y aunque parece muy sofisticado El Reglamento de Inteligencia Artificial responde a esta lógica, responde a la lógica de seguridad de los productos. El problema, como os decía, es que a día de hoy, que estamos, si no me equivoco, a 3 de febrero, no tenemos ningún organismo que esté habilitado y capacitado para emitir este certificado CE que el reglamento requiere a determinados sistemas de IA de alto riesgo.

Speaker 2: Dicho esto, una vez que ya tenemos enmarcado de qué se ocupa el Reglamento de Inteligencia Artificial, hay un segundo punto que creo que es importante que entendamos. No os preocupéis que no os voy a leer todas las obligaciones, este no es el objeto de la sesión de hoy, pero sí que quiero que entendáis un concepto que es el de la cadena de valor. La cadena de valor de la Inteligencia Artificial significa que tus obligaciones van a depender de cuál sea tu rol en la fabricación, o en la distribución o en el uso de un sistema de Inteligencia Artificial.

Speaker 2: Y solo quiero que veáis las diferencias que hay. Por ejemplo, si somos el proveedor, y aquí perdonadme, pero tengo que hacer la crítica, no podían haber elegido peor los nombres del reglamento, porque es que no son nada intuitivos. Para la clase de hoy, quedémonos en que el proveedor es como el fabricante, como un fabricante. Volvamos al símil con los coches. El proveedor es el que fabrica la Inteligencia Artificial, de la misma forma que el fabricante es el que fabrica el coche.

Speaker 2: El importador, no hace falta explicarlo. Si es un sistema estadounidense el que lo introduce en Europa, pues será el importador. Y luego, donde vais a estar la mayoría de vosotros, es en el deployer, el desplegador del sistema de Inteligencia Artificial. Vamos a traducir esto como el usuario, pero el usuario empresarial, que es donde estamos, este es el que decide introducir o contratar ChatGPT, o contratar un borde que tiene Inteligencia Artificial, o contratar herramientas de SAP que tienen Inteligencia Artificial.

Speaker 2: Sois vosotros los que sois el deployer, esos son los usuarios normales y corrientes. Como veis, las diferencias cambian muchísimo. Si nos vamos a un proveedor, y esta es la lista, y me he dejado algunas cuantas, todo esto es lo que tiene que hacer un proveedor de un sistema de alto riesgo. Fijaros, si bajamos a riesgo limitado y vamos a proveedor, pues la cosita ya se va reduciendo. Pero si volvemos a alto riesgo y me voy a deployer.

Speaker 2: Es decir, nosotros, empresas, vemos que tenemos también bastantes actividades, porque es como manejar uranio, por decirlo de alguna forma. No es lo mismo un sistema de inteligencia artificial, como digo, para que le preguntes cuál es la última sentencia o cuál es ayúdame a diseñar un curso para el manejo de las herramientas informáticas dentro de mi empresa, que un sistema que, como digo, va a utilizarse para cribar currículums o definir contrataciones o evaluar empleados.

Speaker 2: Pero en el deployer, y esto es a donde quería llegar, si es un deployer, es decir, una empresa que contrata un sistema de riesgo limitado, las obligaciones son muy poquitas, si os dais cuenta. Básicamente son obligaciones de transparencia. ¿Que quiere decir? Que si estás utilizando un sistema de inteligencia artificial, la gente tiene que saberlo, en el sentido de que tiene mucho que ver, que seguro que en las noticias lo habéis visto, con que la gente sea capaz de detectar, por ejemplo, cuando está interactuando con un charco.

Speaker 2: Si en nuestra empresa decidimos introducir un chat para que hable con nuestros clientes, pues no podemos intentar fingir que están hablando con un humano. Tenemos que ser muy claros a la hora de dejar claro que están interactuando con una inteligencia artificial. Luego esto nos afectará todo. Detección y etiquetado de deepfakes también tiene que ver con los chats, pero por ejemplo, imágenes generadas con inteligencia artificial que pueden dar la apariencia de ser reales cuando. Y en tercer lugar, que esto es importante, la alfabetización de los empleados.

Speaker 2: Es importante que nos aseguremos que nuestros empleados tienen una formación básica en inteligencia artificial. Y aquí, otra vez, perdonadme la informalidad, pero no nos volvamos locos, es una formación básica. No estamos hablando de que un empleado de cualquier área de la organización tendrá que saber qué es un peso o un sistema de vectorización de inteligencia artificial, sino cosas básicas sobre cuándo puede saltar un riesgo en materia de inteligencia artificial.

Speaker 2: Luego hablaremos un poquito más de eso. Pero bueno, como os decía, estas slides tenían más que ver con que seamos capaces de ver cómo cambian las obligaciones de ser elevadísimas en alto riesgo para el fabricante, que aquí es donde digo que el cumplimiento exige ingenieros, expertos en datos, probablemente abogados, departamento de recursos humanos, a cómo puede cambiar en deployer de riesgo limitado, donde básicamente son obligaciones perfectamente asumibles sin intervención experta.

Speaker 2: Dicho esto, como os decía, es una de esas normas donde la colaboración interdepartamental se vuelve esencial, pero donde sí podemos, vamos a intentar quedarnos aquí. Y en este slide intentaba transmitir justo lo que os estaba diciendo. Si conseguimos mantenernos en deployers, usuarios de sistemas de riesgo limitado, el reglamento de Inteligencia Artificial no va a tener un gran impacto en nuestra vida. Lo que pasa es que es

importante entender que lo que hace que la Inteligencia Artificial pueda ser de alto o bajo riesgo no es la herramienta que compramos.

Speaker 2: Es decir, no es que chatgpt sea de riesgo limitado y luego haya una herramienta como Word que pueda ser de alto riesgo, porque tiene que ver con Recursos humanos. No, no. El enfoque es el caso de uso. Es decir, la misma herramienta puede ser de riesgo limitado o de riesgo alto dependiendo del uso que nosotros le demos. Y esto es importantísimo. Porque claro, si yo utilizo, imaginaros, Si yo utilizo ChatGPT, que en principio es un chatbot que podría parecer que es de riesgo limitado, para resumirme sentencias, no estoy aceptando derechos fundamentales, es un uso de riesgo limitado.

Speaker 2: No tengo que preocuparme más que por esas pequeñas obligaciones de formar a mis empleados. Por ejemplo, ¿Cómo formaría a mis empleados yo aquí en el despacho? Pues diciéndoles que, oye, la Inteligencia Artificial alucina, se equivoca, no puedes dar por buenos los resultados. Eso estaría dentro de la formación que le tengo que dar. Pero claro, y aquí voy a adelantar un poquito una de las slides. ¿Qué pasa si yo implanto ChatGPT o Gemini de forma generalizada en la empresa y nuestro departamento de Recursos Humanos o las dos personas que llevan el departamento de Recursos Humanos deciden que es una grandísima idea, que a partir de ahora, antes de ponerse a mirar los currículums, los van a analizar con ChatGPT y si pueden descartar la mitad de los currículums que así con unos parámetros chatgpt considere que no son apropiados para el puesto, pues oye, se van a ahorrar?

Speaker 2: Imagínate que son súper eficientes, te presentan un plan de trabajo y te mira, si empezamos a filtrar los currículums con ChatGPT, vamos a ahorrarnos un 70% de horas que dedicábamos a leer esto. Y puede ser un ahorro de 100 horas, 200, 300 horas anuales para nuestro departamento. Suena maravilloso, pero sin quererlo, sin comerlo ni beberlo, nuestro empleado de Recursos Humanos nos ha metido en un lío grande, porque en ese momento estamos utilizando un sistema de Inteligencia Artificial para un uso de alto riesgo.

Speaker 2: Por lo tanto, insisto, no es que no se pueda hacer, pero sí que tendremos que ser conscientes de que es un uso de alto riesgo, de que nos estamos convirtiendo en usuarios de un sistema de alto riesgo y que esto eleva el nivel de nuestras obligaciones. Si recordáis, ya no estamos en este mundo tan bonito y en este paraíso nos vamos a este. Hay muchas formas, además, en las que sin quererlo, podemos caer en escenas de alto riesgo, que creo que es interesante que las veamos de forma rápida.

Speaker 2: Es un artículo del reglamento que suele pasar un poco inadvertido, pero que discuto más, como os decía, con gente de tecnología o con ingenieros, que con los profesionales abogados. Porque esto quiere decir que incluso cuando, como os decía, tú contratas ChatGPT para la empresa, pero si lo modificas o si cambias su finalidad en una herramienta que aparentemente es una herramienta de riesgo limitado, puedes convertirte sin quererlo en un proveedor de alto riesgo. Y puede suceder en tres entornos muy definidos por el propio reglamento.

Speaker 2: Por ejemplo, si tú decides personalizarlo, meter una capa arriba y que de cara a terceros, ese sistema ya no sea ChatGPT, sino un sistema construido sobre ChatGPT, que tiene la marca de tu empresa, en ese

momento ya no eres un mero usuario, te has convertido en un fabricante, porque tú ya has asumido que ese sistema es tu responsabilidad también si lo modificas, y esto es más habitual de lo que pensamos con determinados sistemas Gemini Cloud, el hecho de que tú puedas modificar los parámetros hace que, de la misma forma que si tuneamos un coche, pues oye, no es lo mismo pintarle de un color distinto a la fábrica, que ya hacer una modificación del motor o del funcionamiento del coche, hasta el punto de que ese coche pueda, por ejemplo, necesitar que se vuelva homologar, pues lo mismo pasa con la Inteligencia Artificial.

Speaker 2: Si yo modifico la forma en la que funciona de fábrica, ya no voy a ser un mero usuario, me convierto en un. Y luego el cambio de finalidad, que es justo el ejemplo que os he puesto con el tema de chatGPT, si yo utilizo una herramienta generalista para un uso que se considera de alto riesgo, me voy a convertir en un usuario de alto riesgo. Y esto además, tiene otros peligros, como veréis aquí, y os lo voy a adelantar un poquito, a la propia chatGPT esto no le hace nada de gracia.

Speaker 2: El hecho de que la gente lo utilice para usos para los que ChatGPT no están diseñado. Y si vamos a leer los términos y condiciones de uso de ChatGPT, y estoy hablando de las versiones profesionales, de las personales, ellos prohíben ese tipo de usos. Y curiosamente han hecho un listado y lo modificaron hace muy poquito, era unos tres, cuatro meses, como podéis ver ahí, tres, perdón, donde prohíben lo que la Unión Europea ha definido uno por uno, usos de alto riesgo.

Speaker 2: Como veis, automatización de decisiones de alto riesgo para infraestructuras críticas, para educación, para empleo, para CICL financieras, para seguros, para ámbito jurídico. Ellos ya te están diciendo, nos están diciendo que su sistema no se tiene que usar para eso. Por lo tanto, si nosotros ignoramos estos términos y condiciones, como empresarios, estamos corriendo dos riesgos. El primero, que estamos incumpliendo los términos y condiciones de Chat PT. ¿Y esto qué significa? Que nos pueden terminar el contrato en cualquier momento.

Speaker 2: Y el segundo, y mucho más importante, el hecho de que nos estamos convirtiendo en otra cosa. Ya no somos un mero usuario de un sistema limitado. Si, por ejemplo, estamos utilizando, como os decía, para finalidades de empleo, para prestar asesoramiento jurídico directo, insisto, etcétera, etcétera, podemos caer del alto riesgo. Y si no cumplimos con las obligaciones que el reglamento impone en los usuarios de alto riesgo, en los deployers, en los desplegados de alto riesgo, nos pueden sancionar.

Speaker 2: Y como hemos visto, las sanciones no son pequeñas. ¿Qué quiero decir con esto? Que tenemos que controlar qué se está utilizando en la empresa y tenemos que tener procesos por los que sepamos muy bien qué estamos contratando y para qué lo estamos contratando. Por eso, en esta pata contractual, que no me quiero extender mucho porque ya es que me enrolló mucho, me tenía que haber regañado Susana, en esta pata de contractual os quería dar unos consejitos pequeños sobre cómo intentar lidiar con estos riesgos que supone la implementación.

Speaker 2: Y lo primero es eso, identificar quiénes somos en la cadena de valor. Si yo estoy comprando una suscripción Enterprise de Gemini, pues probablemente voy a ser simplemente un desplegador del sistema. Pero

claro, tengo que ver para qué lo uso. Porque claro, no es lo mismo ser, como hemos visto, un desplegador del sistema de riesgo limitado, que es un depredador del sistema de riesgo. Luego, además, tengo que intentar que ese proveedor, igual que haría si estoy comprando una flota de coches, me dé garantías de que su sistema funciona bien.

Speaker 2: Por ejemplo, en el caso del alto riesgo, si estoy comprando un sistema, imaginaros, utilizo GPT en toda la empresa, pero mi equipo de recursos humanos me ha pedido, porque lo necesita, un sistema para gestionar los procesos de onboarding o los procesos de selección de candidatos, yo decido adquirir o contratar un sistema de inteligencia específico para recursos humanos que puede tener el riesgo de ser de alto riesgo. Y entonces yo tengo que entrar en una negociación relativamente dura con ese proveedor para oye, justifícame que tienes la declaración de conformidad de la Unión Europea del reglamento, justifícame que cumples con todos los requisitos del reglamento.

Speaker 2: Y muy importante, y voy a decirlo así por encima, justifícame también que sé lo que os digo en esta última slide, que en el caso de que esto no sea así, me vas a mantener indemne. Es decir, que si a mí me pasa algo por tus incumplimientos, tú vas a hacerte cargo de que yo no sufra ningún daño. Si tengo que acabar pagando una indemnización a un tercero porque ha pasado algo, tú me vas a acabar pagando a mí, es decir, una cláusula de goce pacífico o de indemnidad.

Speaker 2: Y aquí entra en juego una parte que no es jurídica, sino de sentido común. Es importante saber con quién estamos contratando el uso de este tipo de sistemas. Porque claro, si lo contrato, y aquí no quiero me malinterpretéis, pero si lo contrato a una persona que no tiene, que acaba de entrar en el mercado, que no tiene recursos para hacer frente a una posible indemnización, a una posible multa o una indemnización que me puedan imponer de tres, cuatro millones de euros, pues la cláusula de indemnidad me va a servir de muy poco, porque esa empresa va a quebrar, no me va a pagar y a mí no me va a mantener indemne absolutamente nadie.

Speaker 2: Pues entonces, ¿Qué quiero decir? ¿Que no se contrate con el pequeño? No, no quiero decir esto. Muchas veces los pequeños son muy buenos y pueden dar garantías de cumplimiento o se mueven de forma más ágil y tienen justo el producto que necesitamos. Pero en ese caso, pensad en los seguros, en el hecho de que esa empresa pueda contratar un seguro donde a ti se te asegure, nunca mejor dicho, que te van a poder mantener. Inexplicable. Son pequeñas, pequeñas recomendaciones contractuales estratégicas para que tengáis en cuenta a la hora de implementar esos sistemas.

Speaker 2: ¿Qué más tenemos? ¿Una pata que es extremadamente importante? Voy a pasar aquí el triaje de riesgos. Esto suena como muy sofisticado, pero al final es haceros cuatro preguntas. ¿Lo que queréis hacer con Inteligencia Artificial en vuestra empresa está prohibido? Aquí la respuesta es extremadamente sencilla. Si está prohibido, no lo hagáis. Es que no hay nada que podamos hacer para solventarlo. En el punto 2 ya empezamos en el mundo es alto riesgo.

Speaker 2: Perfecto, pues tenemos que tener mucha atención con el proveedor que estamos contratando, porque las multas son altas, porque el impacto es alto, entonces tenemos que hacer un screening de proveedor muy

elevado. Y si no nos da garantías, pues a lo mejor tenemos que replantearnos si es nuestro proveedor. Tenemos que tener un control muy grande, como decía, de lo que autorizamos o dejamos de autorizar en la empresa. Los empleados tienen que tener muy claro lo que pueden hacer y lo que no pueden hacer con las herramientas.

Speaker 2: Y para esto, como veremos, son importantes las políticas internas de Inteligencia Artificial. Igual que tenemos uno de protección de datos, igual que tenemos uno de confidencialidad, pues a lo mejor llega el momento de modificarlas o incluso si son usos muy específicos, de hacer una específica para el uso de Inteligencia Artificial en la empresa. Pero el paso cero para todo lo que os estoy contando es, como digo, inventariar. Ahora que ya sabemos qué es lo que espera el reglamento de nosotros, qué diferentes riesgos establece el reglamento, cómo cambian esos riesgos dependiendo de si yo diseño o fabrico mi propia IA o si compro la IA de un tercero.

Speaker 2: Pues es que el primer paso, y perdonad, pero esto es muy cíclico, es definir el alcance, hacer un inventario, conocer, saber lo que está pasando en la empresa. Y a veces nos vamos a sorprender. Lo que os decía, ¿No? Mi empresa no se está utilizando. Habéis preguntado, habéis visto cómo está trabajando la gente, habéis visto qué hace el departamento de marketing con sus creatividades, o cómo están desarrollando código vuestros departamentos de tecnología, porque esto tiene implicaciones, por ejemplo, en materia de marketing, y aquí hago un inciso, o de desarrollo de código.

Speaker 2: Estamos teniendo nuevas preguntas en los procesos de inversión y de due diligence, porque claro, esto no era el objetivo de hoy, pero simplemente apuntároslo. Los resultados que se generan mediante el uso de Inteligencia Artificial Generativa no son protegibles por propiedad intelectual. Por lo tanto, el código fuente que se genera con una herramienta de IA entra directamente en el dominio público, no somos sus propietarios. Los materiales de marketing que se generan directamente con Inteligencia Artificial Generativa, lo mismo.

Speaker 2: Por lo tanto, si yo estoy asesorando a una empresa para que compre o invierta en una startup dedicada al software, y cuando le pregunto a esa empresa, ¿Estás utilizando Inteligencia Artificial Generativa para desarrollar el código? Me dice que sí. Yo tengo que levantar la mano, hagamos un red flag y decirle a mi cliente que ojo con la inversión porque no tienen la propiedad del software, que ese software lo puede utilizar cualquier, que eso a lo mejor se puede gestionar, sí, pero lo tienen que saber, o con los materiales de marketing, sí, precioso, pero si ves tu fotografía o tu ilustración en un anuncio de la competencia de un tercero, no vas a poder hacer absolutamente nada para impedirle ese uso al tercero.

Speaker 2: Son pequeños apuntes que tenemos que tener en cuenta. Una vez que ya tenemos detectado, pues hacemos una ficha. Tan sencillo como eso. Es aburrido, pero es práctico. Cómo se llama el sistema, para qué se va a utilizar, para qué no se puede utilizar, cuál es el rol de la organización, etc. Como os decía, detectar IA en la sombra y luego tener en cuenta que este ejercicio no hay que hacerlo solo ahora, porque está a punto de entrar en el reglamento de Inteligencia Artificial en vigor.

Speaker 2: Es un inventario dinámico. Yo, para mi desgracia, tengo que revisar todos los que se valoran en Garrigues y somos 2.500 para su implementación. Porque la gente de marketing quiere hacer esto así, porque la

gente de Recursos Humanos quiere hacer esto así, porque nosotros en IP queremos hacer cosas con alguna herramienta de IA. Pues hay que revisarlos todos, ver si son adecuados y ver si entrañan algún riesgo para la empresa. Y es rara la semana que no recibo una petición para un sistema nuevo y no creo que vaya a ir a menos, creo que va a ir a más.

Speaker 2: Y aquí nada, os había puesto un ejemplo de unas fichas que estamos haciendo para algunos clientes de cómo pueden gestionar esto internamente para que no sea ingobernable, cómo tener inventarios dinámicos de los sistemas que están utilizando podía acceder, pero como vamos un poco más de tiempo, pues no. Política interna de uso. Y aquí ya Susana terminó, que no me quiero comer el tiempo. Política interna de uso es sencillo, es intentar explicar con los términos más sencillos posibles lo que la organización permite y no permite.

Speaker 2: Establecer si se puede una formación para la gente que vaya a utilizar herramientas de inteligencia artificial. Tener muy claro que la gente que vaya a realizar usos de alto riesgo necesita una formación adicional. Por ejemplo, si nuestros empleados de recursos humanos empiezan a detectar ciertos sesgos o fallos en un sistema, tienen que ser capaces de levantar la mano y tenemos que formarles para esto. Y sobre todo tener muy claro que nuestra empresa, si puede mantenerse en el paraíso del usuario de sistemas de riesgo limitado, tiene que hacer todo lo posible para que de forma inadvertida se pase a un riesgo no controlado, a un uso de alto riesgo no controlado.

Speaker 2: Entonces tenemos que tener muy claro que los empleados no pueden empezar a utilizar la inteligencia artificial por su propia cuenta y riesgo, porque nos pueden meter en un lío. Y aquí ya sabéis que aplica el principio de culpa y no eligiendo culpa y vigilando. No vale decir, es que yo no lo sabía. No, no hay que saberlo, hay que saber lo que se está haciendo con las herramientas corporativas de la empresa. Y luego, simplemente para que tengáis en cuenta dos conceptos que creo que es interesante, y con esto ya sí que termino, que son los conceptos de cómo reducir los riesgos dentro de la empresa.

Speaker 2: Bueno, hay herramientas de mercado, hay tecnologías que incluso se pueden diseñar internamente que son importantes, sobre todo, por ejemplo, cuando tenemos la intención de utilizar chatbots, chatbots que van a utilizar con nuestros clientes. Y esto ya no es solo que también porque puedan meternos en un lío de no cumplir con el reglamento de inteligencia artificial, sino también porque pueden arruinar la reputación de nuestra empresa si no funcionan bien. Entonces, ¿Qué dos herramientas os quería comentar? Lo que se llama el red teaming, que parece así como un palabro muy tal, pero que no deja de ser hacer como un hacker.

Speaker 2: Es decir, tú mismo le sometes a pruebas a tu inteligencia, a tu sistema de inteligencia artificial para ver si falla. Por ejemplo, imaginaros un banco que no puede dar asesoramiento financiero a través de su chat, ¿Qué sería una buena prueba de recursos para que te dé el asesoramiento financiero? Preguntar, preguntar, preguntar, darle la vuelta, empezar a hacer trucos para ver si el sistema falla, para someterle a lo que llamamos una prueba de estrés. Y luego vendría ese segundo punto, los guardarraíles, que son una vez detectados los fallos, como pasa con las carreteras, estoy con ejemplos de coches, pero una vez detectados los fallos, se rediseña el sistema o se crean guardarraíles para que esos fallos no vuelvan a pasar.

Speaker 2: Y os cuento un caso real. Nosotros estuvimos prestando, valorando, que me tocó además valorarlo a mí, con un colaborador con el que nosotros colaboramos, un tercero, que es una empresa de asociada al Barcelona, es un spin off del Barcelona Supercomputing Center. Hicieron una prueba de red teaming con un sistema que nosotros queríamos contratar. Bueno, falló en el 100 de las ocasiones, obviamente lo descartamos. No voy a decir que el sistema es, pero era un sistema de una especie de avatares que te ayudaban a resolver dudas de uno de nuestros productos.

Speaker 2: Pues bueno, acababan diciéndote cualquier cosa a la que le preguntabas, pues eso, que ayudaban a hackear, a construir bombas y hacía falta. Pues obviamente ese sistema no es algo que pueda permitirse, ni quiera utilizar Garrigues, porque incumple todas nuestras políticas y el reglamento de Inteligencia Artificial. Estaría dispuesto a dar asesoramiento jurídico, como estáis viendo en pantalla, ayudarte o darte tips de cómo hacer un scam financiero. Entonces, obviamente no son sistemas que queramos implementar en nuestras empresas.

Speaker 2: Y bueno, con esta sí que termino la última. ¿Esto qué significa? Pues que al final es un trabajo que nadie puede hacer solo el reglamento, y la dinámica es tan complicada que necesita una buena visión de conjunto, que es lo que yo he intentado ofreceros hoy por parte de los directivos, por parte de los CEOs, sea cual sea, por parte de la C Suite que dicen los cursis otra vez, para poder implementar y mitigar los riesgos de forma segura, pero que necesita la colaboración de las distintas unidades de la empresa, principalmente de las que quieren utilizar la IA en primer lugar, que son los que en principio van a permitirnos ser más eficientes.

Speaker 2: Luego, obviamente, del departamento jurídico, que es el que tiene que decirte lo que puedes y no puedes hacer, pero también de tecnología, que por ejemplo tiene mucho que decir, pues en las pruebas de red teaming, en las pruebas de ver si podemos mitigar riesgos, etc. Nosotros para las empresas más grandes, donde hay muchos empleados, pues sí que recomendamos la creación de comité, que son los que puedan valorar con toda la seriedad que requiere, cuáles son los riesgos y que se va a recomendar.

Speaker 2: Y ahora sí que sí, que con esto ya he terminado.

Speaker 1: Pues muchísimas gracias, Cristina, la verdad es que súper interesante, mucha información. Y eso te iba a decir, digo vamos a. Porque hay unas cuantas preguntas, sobre todo para que haya tiempo ahora de lanzarlas. Los que estáis conectados, si podéis conectar la cámara mejor, porque creo que así es un poco más cercano, sobre todo para lanzar las preguntas. Hay varios comentarios y algunas preguntas, entonces voy leyendo. Pero por favor, si continúa la persona que ha lanzado la pregunta, pues que abra el micro y lo diga.

Speaker 1: Sara García nos ¿Y si existe un sistema que incluye opciones de personalización y un distribuidor te lo personaliza para poner en tu web, ¿Qué papel juega la empresa que compró la licencia inicial y contrató la personalización a un distribuidor?

Speaker 2: Pues es muy buenísima pregunta. Entrás de lleno en el riesgo del artículo 25. Y el tema es que pasa muchísimo más de lo que pensáis, porque sobre todo con los sistemas grandes, por cómo funciona, y esto va a ser un tema geopolítico, los grandes proveedores de sistemas de inteligencia artificial, insisto, GPT, Gemini, Claude,

son cada vez más infraestructura. ¿Qué está pasando? Que se está construyendo muchas capas. Por ejemplo, os doy un ejemplo. En el mundo jurídico, que es el mío, hay una herramienta, o hay dos que se están utilizando mucho, que son herramientas que se dicen que están personalizadas para el mundo jurídico, son Harvey Levora, algunos las conocéis.

Speaker 2: Entonces esa herramienta lo que han hecho es construir una capa sobre el modelo de GPT. Entonces luego encima suelen tener un integrador que te vende la herramienta. Entonces ya empiezas a tener un montón de gente y al final del día no sabes quién te da garantía, quién te está simplemente revendiendo, quién es un mero distribuidor, quién está asumiendo la responsabilidad y ese es el primer paso. ¿El primer paso es cuando tú hables con la persona a la que le has contratado la solución oye, pero tú me das la garantía sobre el cumplimiento completo?

Speaker 2: Y para mi sorpresa, muchas veces te han dicho, yo la pata que hace GPT no te la puedo garantizar. Entonces claro, ¿Dónde me dejas a mí? Yo contrato solo una franja contigo, pero yo no tengo ningún contrato ni le puedo exigir ninguna responsabilidad GPT. ¿Esto cómo funciona? ¿Yo le tengo que pedir responsabilidad aguas arriba o aguas abajo? ¿Como la queréis mirar? Pero esa cadena de distribución es esencial. Por eso sí a nosotros se nos caen muchos proyectos y a veces tenemos que desmontar algunos porque no tenemos garantías ni nos las pueden dar.

Speaker 2: En el caso de Harvey si tiene un acuerdo grande, entonces no les quiero que no parezca que como les he puesto de ejemplo, pero claro, no todos los proveedores son así y cuando te personalizan y personalizan bajo tu cuenta de riesgo, es muy probable que tú te conviertas en fabricante.

Speaker 3: Y si la pregunta era mía, es que estoy pensando en un caso concreto que estamos moviendo ahora para contratar en breve. ¿Y qué forma habría de mantenerte? Como diciendo, bueno, esto es un servicio que me está prestando un tercero de alguna manera evitando esa personalización, informando en una nota diciendo, aparte de lo que ya es para un chatbox, en este caso ya teníamos pensado decir esto es una IA, las respuestas, o sea, en caso de duda prevalece lo que ponga la normativa, ese tipo de cosas.

Speaker 3: Pero en ese caso nosotros podríamos decir, este es un servicio prestado con tal herramienta y bajo su responsabilidad, algo así,

Speaker 2: la responsabilidad va a ser siempre tuya. Otra cosa es como tú contratas aguas arriba, vale, esto no vas a poder hacer y yo lo que tendría que

Speaker 3: contratar es dejar explícitamente que mi sistema debe cumplir, seguramente es riesgo bajo, pero bueno, que mi sistema debería cumplir eso y contratarle al proveedor que me garantice que eso se va a cumplir.

Speaker 2: Te lo garantice ya, acuérdate, no sé el tamaño del proveedor, pero si es muy pequeño, acuérdate del seguro. Sí, también, pero a la hora de vender, claro, es que empieza a haber

Speaker 3: un lío de quién crea el producto. Este además utiliza una licencia de un tercero para el modelo de edad, o sea, pierdes totalmente el control de lo

Speaker 2: que hay en mi vida diaria. Por eso siempre es verdad que uno lo ve sobre el papel y dice, qué rollo esto, pero es que al final es donde están atascadas todas las empresas porque hay un lío de cadena de distribución tremendo y no se nos puede olvidar, es que normalmente ellos dependen de un tercero, tiene una servidumbre muy fuerte porque su producto se cae, es decir, si GPT se ha colgado ese día, el producto que te han vendido aquí.

Speaker 2: Sí, sí.

Speaker 3: La siguiente pregunta que hay también es mía, si quieres, era un poco cuando se estaba contando sobre si al final una misma herramienta puede ser de un riesgo más alto o más bajo, dependiendo de mi caso de uso concreto. ¿Quién decide el riesgo de ese caso de uso? Yo misma. Cuando me evalúo digo, mira, voy a hacer este fin con un propósito sencillo, no voy a vulnerar nada, yo soy de riesgo moderado, por ejemplo, vamos, el

Speaker 2: uso de mi aplicación, ahí aplicas la ley, es decir, lo importante es que te expliquen muy bien desde como has dicho, el uso, para qué lo vais a utilizar. Entonces en esa ficha tú justificas con sentido común que ese uso que tú estás haciendo cae dentro, o mejor dicho, no cae dentro de lo que son usos prohibidos ni altos riesgos, justifica por qué. Y en tu política de uso o de autorización de ese sistema deja bien claro, oye, chicos, esto se puede usar para esto, que es para lo que me has pedido, no vengas luego con ideas felices para utilizarlo para esto, porque yo te lo he utilizado para, no para mí.

Speaker 2: Por eso digo que las fichas hay que mantenerlas y son importantes. Creo que estás en Newt, Sara,

Speaker 3: perdón, que entonces la idea es que nosotros cuando yo estoy en la parte de desarrollo IP, bueno, estos temas de. Entonces nuestra idea sería hacer una ficha diciendo, este es el uso de lo que se nos ha pedido aquí, este es el riesgo que se determina, lo veremos con nuestro departamento jurídico, este es el riesgo que tenemos y entonces estas son las medidas a implementar y en caso de que en algún momento la UE o alguien decida hacer una auditoría, deberíamos sacar eso.

Speaker 3: ¿Pero ahí no tendremos que tener ningún otro certificador de un tercero, ni con su día de hoy? No, en principio espero que sea de riesgo limitado. En algún momento podría llegar a haber un uso de un riesgo más alto y habría que decir, vale, pues en este momento qué pido a mi proveedor, o sea, este riesgo hay que hacerlo muy en la definición de los requisitos

Speaker 2: del sistema y cuando la gente, la gente tiene que tener accesible una recomendación que seguimos nosotros mismos y siempre damos a nuestros clientes, la gente tiene que saber para qué puede utilizar cada.

Speaker 3: Muchas gracias, muy claro. Y otro otra cosita más, yo que

Speaker 2: vivía muy feliz de deployer saben aquí los comentarios. Estoy totalmente de acuerdo con todos, o sea, Inés, estoy 100% de acuerdo contigo. Esto no es sólo del FISO, de hecho este es el mensaje que he querido trasladar, si no ha quedado claro, espero que quede claro, es decir, todas las cabecitas que tengamos en la empresa están aquí, no solo el CISO, ni muchísimo menos, y creo que es un error enfocarlo de otra forma. Otra cosa es que para entender muchas cosas tengamos que hablar con que son cosas distintas y sí el 60% de las flechas factor humano, te diría que incluso más.

Speaker 1: Que han estado ahí comentando. ¿Yo creo que la pregunta de André Andrés te ha respondido Cristina o te queda alguna duda?

Speaker 2: Bueno, mi pregunta tenía que ver para ver si en el propio contrato del proveedor podemos poner las limitaciones a través de lo que la IA te puede permitir para que la regulación limite el servicio en aquellos aspectos que fueran de riesgo alto o riesgos que no se pueden. Es decir, que el propio con que puedas establecer tú las condiciones en tu contrato con el proveedor para que el proveedor se autolimite sin necesidad, pero eso teóricamente me parece posible, pero tiene un riesgo que le trasladas a él tu cumplimiento y si él no lo hace, el culpable vas a ser tú, pero contractualmente él es el que tiene los límites, no siempre tú. Puedes también intentar poner guardarrailes en una capa superior, es una solución, pero si vas a.

Speaker 2: Depende del sistema también depende mucho de la herramienta que vayas. Con sistemas generalistas ni de broma, no lo van a hacer con ningún concepto, como mucho te van a poner algo y además son para todo el mundo porque esta gente va a escalar, o sea, no venden soluciones ad hoc. Y luego con los integradores, pues claro, lo puedes hacer. Si lo dejaría solo a la pata contractual, me daría un poco de miedo y en ese caso vigilaría mucho cómo lo estáis haciendo, qué test de estrés que estáis utilizando, o sea, estaría muy vigilante.

Speaker 2: Se puede hacer, depende del caso de uso, lo recomendaría en algunos y en otros no. El problema es que a lo mejor se crea un sistema demasiado paralizado a evolucionar. Me da miedo depender tanto de un tercero, aunque colaborar con el tercero no me parece mala idea si se puede. Sistemas dedicados como muy especializados por industria y tal, puede ser interesante. Gracias.

Speaker 1: Si queréis cerramos con esta última pregunta que dice Mireia. Si tenemos contratadas licencias de Microsoft 365, que ya incluye Copilot, ¿Habría que hacer algún otro contrato documento si se va a hacer uso de riesgo limitado?

Speaker 2: Es una buenísima pregunta. El tema es saber si consigue que Microsoft se lo haga. La pregunta es muy buena. A ver, Microsoft quiere que sus herramientas se adapten de forma masiva. Microsoft además todavía sigue utilizando la base de chat GPT, con lo cual también tiene dependencia de un tercero. Lo cierto es que Microsoft es de los poquitos que te dan indemnidades por incumplimiento porque ellos mismos querían vender. Por poneros un ejemplo, dentro de las condiciones de uso de Copilot hay una garantía, que yo además que me dedico al IP, pues me gusta mucho, de los poquitos que te dan en una situación muy caótica como la que vimos ahora, una indemnidad por infracciones de propiedad intelectual, como puede pasar y se está viendo, al final puede haber problemas de IP por los resultados, o sea, de derechos de autor, perdona, por los resultados

generados porque infrinjan derechos de terceros, porque no se ha entrenado con licencia, etcétera, etcétera.

Speaker 2: Y al final, como en IP la responsabilidad es objetiva, te acaban pidiendo una indemnización al p. Microsoft se hace responsable.

Speaker 1: Muy bien, yo creo que se han contestado todas las preguntas. No sé si alguien quiere hacer alguna otra.

Speaker 2: A Inés es que no la contesto, pero porque estoy de acuerdo, o sea que no tengo nada que decir.

Speaker 1: Inés ha puesto comentarios más que preguntas, que efectivamente ya has comentado algunos y si no hay ninguna más, si te parece, Cristina, lo dejamos aquí. Muchísimas gracias por la claridad de tu presentación, porque además creo que has aterrizado los conceptos para los que no somos expertos y muy interesante, creo que has aclarado muchísimas dudas. En cualquier caso, si te parece bien, si alguien le queda alguna duda o alguna cosa, pues podemos pasar tu contacto, la grabación y la presentación la pasaremos a todos.

Speaker 1: Así que nada, pues que muchísimas gracias a todos por estar conectados y por acompañarnos. Hoy.

Speaker 2: Gracias Cristina, a vuestra otra estás. Un placer.