

EL IMPACTO SOCIAL Y NORMATIVO DE LA INTELIGENCIA ARTIFICIAL

1ª SEMANA INTERNACIONAL

(13-16 DE MAYO 2025)

UNIVERSIDAD DE OVIEDO

por Paolo Inturri

PREFACIO

A continuación, se presenta el informe de la 1.ª semana internacional sobre "*El impacto social y normativo de la inteligencia artificial*" celebrada del 13 al 16 de mayo de 2025 en la Universidad de Oviedo.

En un esfuerzo por representar fielmente el desarrollo de las actividades, la exposición se ha dividido por días.

Para cada día, se ilustrará el contenido de las distintas mesas redondas, resumiendo los puntos clave de cada presentación.

ÍNDICE

MARTES 13 DE MAYO 2025

<u>INAUGURACIÓN INSTITUCIONAL DE LAS JORNADAS</u>	4
<u>EL CONTEXTO DE LA IA</u>	5
SUSANA MONTES RODRÍGUEZ	5
FERNANDO HIGINO LLANO ALONSO	6
IBAN GARCÍA DEL BLANCO	7

MIÉRCOLES 14 DE MAYO 2025

<u>LA REGULACIÓN DE LA IA</u>	9
MIGUEL ÁNGEL PRESNO LINERA	9
JOSÉ LUIS BOLZAN DE MORAIS	11
JAVIER FERNÁNDEZ RODRÍGUEZ	12
<u>EDUCACIÓN Y IA</u>	14
JAVIER FOMBONA	14
JULIO ALBALAD GIMENO	14
THOMAS CASADEI	16
<u>CIBERSEGURIDAD</u>	17
MARÍA JOSEFA RIDAURA MARTÍNEZ	17
LUIS MOLINA VALBUENA	18
STEFANO PIETROPAOLI	19

JUEVES 15 DE MAYO 2025

IA Y TRABAJO

IVAN ANTONIO RODRÍGUEZ CARDO	19
LUCIA MONTEJO	20
JAVIER CAMPA	21
ISABEL SANTOS	21

IA Y SALUD

SERGIO CALLEJA	23
ELENA DÍAZ RODRÍGUEZ	24
PEDRO JOSÉ ESPINOSA PRADOS	25

ECONOMÍA Y SOSTENIBILIDAD DE LA IA

LUIS MOLINA VALBUENA	27
CARLOS IBÁÑEZ SÁNCHEZ	28

VIERNES 16 DE MAYO 2025

IA Y JUSTICIA

JORDI NIEVA FENOLL	29
ALFREDO RODRÍGUEZ DEL BLANCO	29
AITOR CUBO	30

CONFERENCIA DE CLAUSURA

LORENZO COTINO HUESO	32
----------------------	----

MARTES 13 DE MAYO 2025

INAUGURACIÓN INSTITUCIONAL DE LAS JORNADAS

Comité Organizador: Agustina **Bouchet** Gutierrez, Roger **Campione**, Miguel Ángel **Presno** Linera

El comité organizador dio la bienvenida a los participantes en la sede del **Centro de Estudios sobre el Impacto Social de la Inteligencia Artificial (CEISIA)** en la ciudad de Lugones.

En sus intervenciones, Agustina Bouchet Gutierrez, Roger Campione y Miguel Ángel Presno Linera, los tres profesores de la **Universidad de Oviedo (UNIOVI)** y miembros del CEISIA, explicaron la idea fundamental que impulsó la iniciativa.

En este sentido, el comité organizador enfatizó la necesaria **interdisciplinariedad** de los estudios en materia de inteligencia artificial (IA) y la importancia del diálogo entre la academia, las instituciones y la sociedad civil, lo cual quedó patente en la cualificación de los ponentes presentes en las diferentes jornadas.

En la visión de los organizadores, el enfoque de la IA debe ser **transversal** y, por lo tanto, debe beneficiarse necesariamente de la colaboración entre humanistas y técnicos.

Los primeros se interrogan sobre el impacto social y normativo de la IA, careciendo, sin embargo, de los conocimientos técnicos necesarios; los segundos, aunque son expertos en los aspectos técnicos, no pueden prescindir de la aportación de los humanistas para comprender y orientar el impacto de la IA en la sociedad.

EL CONTEXTO DE LA IA

Ponentes: Susana **Montes** Rodríguez, Fernando Higinio **Llano** Alonso, Iban **García** del Blanco

Moderador: Roger **Campione**

Contenido: En la primera mesa redonda los ponentes debatieron sobre el contexto de la IA desde tres perspectivas distintas: (1) técnica (Montes); (2) ética (Llano); (3) política (García).

SUSANA MONTES RODRÍGUEZ, profesora de matemática y estadística en la UNIOVI y miembro del CEISIA, abrió la discusión presentando su perspectiva técnica sobre la IA.

En primer lugar, Montes comparó el desarrollo de la IA con la investigación médica, destacando cómo en ambos sectores la técnica primero crea las herramientas, pero luego corresponde a la ética y al derecho establecer qué es lícito o ético hacer.

Más en general, la ponente subrayó que "IA" es una expresión muy extendida, pero con demasiada frecuencia poco comprendida.

Los métodos predictivos en la base de la IA existen desde hace tiempo; sin embargo, actualmente se asiste a un verdadero "boom" debido a la creciente potencia de cálculo de los ordenadores y a la cada vez mayor disponibilidad de datos.

Aunque el desarrollo de la IA parte de la idea de que **las máquinas** pueden pensar, lanzada por Alan Turing en 1950, estas no piensan realmente.

Herramientas como ChatGPT, aunque son excelentes para generar texto, basándose en enormes volúmenes de datos procesados rápidamente, no ejecutan operaciones lógicas ni "razonan" en el sentido humano. Reproducen, basándose en la probabilidad, lo que han visto con mayor frecuencia en su entrenamiento. Precisamente por esta razón, la **actitud crítica** y la **supervisión humana** son indispensables. Sin supervisión, la IA puede ser peligrosa debido a su tendencia al error, pero en cualquier caso la responsabilidad por el uso de la IA siempre recae en el operador humano.

Basándose en los desarrollos actuales de la IA, que reportan resultados sorprendentes en tareas como la generación de texto y el reconocimiento de patrones, la ponente opinó que la IA no eliminará el trabajo humano, sino que lo transformará, haciéndolo más intelectual.

A pesar de ello, Montes destacó las principales dificultades técnicas que, a día de hoy, siguen siendo un desafío para el desarrollo de la IA:

(1) Sesgo: La IA refleja los prejuicios presentes en los datos de entrenamiento, lo que puede llevar a resultados discriminatorios para ciertos grupos, los cuales, sin embargo, no derivan de una intención discriminatoria, sino de la composición de los propios datos;

(2) Dependencia y Vulnerabilidad: Cuanto más delegamos tareas de nuestra vida a la IA, más dependemos de ella y corremos el riesgo de vernos vulnerables en caso de ciberataques o mal funcionamiento de la red;

(3) Privacidad y Seguridad de los Datos: La IA aprende de los datos introducidos por los usuarios, lo que requiere cautela al proporcionar información sensible o no privada.

(4) Desigualdades: La IA presenta diferencias, en términos de oportunidades de acceso a esta tecnología, entre quienes disponen de los medios y el conocimiento para explotarla y quienes, en cambio, carecen de ellos.

A la luz de esto, Montes concluyó su intervención subrayando la existencia de una responsabilidad colectiva de aprender a usar la IA correctamente.

FERNANDO HIGINO LLANO ALONSO, profesor de filosofía del derecho en la Universidad de Sevilla, abordó la discusión desde una **perspectiva ética y normativa**.

En particular, subrayó cómo la ética de la IA no se traduce en una aplicación particular de la ética tradicional, sino más bien en un conjunto de principios que las **empresas** deben respetar para ser **fiables y creíbles en el mercado europeo**.

Más en detalle, para regular su propio mercado, la Unión Europea ha tenido en consideración los **principios éticos fundamentales de la IA**: Justicia, Autonomía, Benevolencia, No-maleficencia y Explicabilidad.

Precisamente por el respeto de tales principios, el AI Act puede considerarse un modelo de regulación antropocéntrico, a diferencia del modelo estatista chino y del ultraliberal estadounidense.

Según Llano, para comprender plenamente la ética de la IA es de fundamental importancia considerar que **los algoritmos en sí no tienen ética**, sino que reflejan los sesgos, incluso inconscientes, de quienes los diseñan y desarrollan: seres humanos que, en cambio, son portadores de una ética específica.

La intervención concluyó remarcando la importancia de la **alineación de la IA con los valores humanos** y la necesidad de una amplia alfabetización digital, como un proceso que involucre con responsabilidad compartida a la administración pública, la ciencia y la sociedad civil.

IBAN GARCÍA DEL BLANCO, exeurodiputado y negociador del AI Act, aportó a la discusión una **perspectiva política y regulatoria** de la IA, con particular referencia al contexto europeo.

El punto de partida de su intervención fue la constatación de que, a pesar de la aprobación del **AI Act**, los rapidísimos desarrollos en el sector de la IA plantean diversas cuestiones fundamentales que tarde o temprano nuestra sociedad deberá abordar.

A pesar de ello, el ponente se mostró optimista respecto al AI Act, el cual, aunque no esté en vigor en su totalidad, ya ha logrado influir en los operadores del sector.

Estos últimos, en particular, parecen ver con desaprobación esta normativa y, a través de una intensa actividad de *lobbying*, buscan aligerar las obligaciones impuestas por el AI Act. En otras palabras, la "competición tecnológica" no puede justificar iniciativas de desregulación, ya que solo el derecho puede garantizar la centralidad del ser humano.

De hecho, aunque el AI Act considera la IA como un producto, su fin último es la **protección del ser humano** y de los derechos fundamentales.

Entre los aspectos del AI Act dignos de mención, el ponente destacó:

- (1) La **vocación geográfica global (extraterritorial)**: se aplica a todos los operadores que actúan en el mercado europeo, independientemente de su origen;
- (2) La exclusión del ámbito de aplicación de los sectores **militar** y de **seguridad nacional** (al ser de competencia exclusiva de los Estados miembros), para los cuales sería deseable un marco ético común;
- (3) La introducción de **nuevos derechos para los ciudadanos**, como el derecho a ser informados y a recibir una explicación razonable sobre el uso de la IA que les concierna;
- (4) La obligación de **consultar a los trabajadores** antes de la implementación de la IA en el lugar de trabajo.

Finalmente, el ponente subrayó la necesaria conexión entre IA y **democracia**. De hecho, dado el potencial transformador de la IA, los ciudadanos deben participar en las decisiones relativas a su implementación y sus límites. Al mismo tiempo, la adopción de **decisiones democráticas y colectivas** conscientes implica un nivel generalizado de **educación digital** entre la ciudadanía activa, so pena de crear una brecha de desigualdad digital.

MIERCOLES 14 DE MAYO 2025

LA REGULACIÓN DE LA IA

Ponentes: Miguel Ángel **Presno** Linera, José Luis **Bolzan** de Morais, Javier **Fernández** Rodríguez

Moderador: Miguel Ángel Presno Linera

Contenido: La mesa redonda inició la segunda jornada de la Semana Internacional de la IA. Los ponentes, reunidos en el aula "Severo Ochoa" del edificio histórico de la Universidad de Oviedo, abordaron el tema de la regulación de la IA desde tres perspectivas diferentes: (1) la europea (Presno); (2) la brasileña (Bolzan); la de la administración pública asturiana (Fernández).

MIGUEL ÁNGEL PRESNO LINERA, profesor de derecho constitucional en la UNIOVI y miembro del CEISIA, abordó el tema de la regulación de la IA en la UE, centrando su exposición en tres temas principales: (1) el procedimiento de aprobación del AI Act; (2) la definición jurídica de IA; (3) los puntos críticos del AI Act.

(1) EL PROCEDIMIENTO DE APROBACIÓN DEL AI ACT

Antes de reconstruir el procedimiento de aprobación del AI Act, Presno partió de los resultados de los estudios de la Profesora Anu Bradford.

En particular, esta última, en su obra "*Digital Empires: The Global Battle to Regulate Technology*", sistematiza los modelos de regulación de la IA de EE. UU., China y la UE, subrayando que ninguno de ellos se realiza en la práctica en su versión pura.

El modelo estadounidense se basa tendencialmente en el libre mercado y, por lo tanto, en una intervención normativa y estatal mínima.

El modelo chino prevé, en líneas generales, una regulación estatal total, con el objetivo de utilizar la IA como instrumento de control social y político.

El modelo europeo se centra en la garantía de los derechos, promoviendo el desarrollo tecnológico respetando los derechos fundamentales.

Coherentemente con este último modelo, la UE ha promulgado el AI Act con el fin adicional de desencadenar lo que la profesora Bradford denomina el "*Efecto Bruselas*", es decir, el efecto por el cual los Estados extracomunitarios promulgan normativas similares a las comunitarias para poder proporcionar a sus empresas un marco regulatorio homogéneo a nivel global.

Dicho esto, Presno reconstruyó el procedimiento normativo que condujo a la aprobación del AI Act.

Así, recordó: que el Libro Blanco sobre la IA (2020) representa el punto de partida del procedimiento de regulación de esta tecnología en la UE; que a este le siguió la propuesta de Reglamento de la Comisión (2021); que el Parlamento Europeo presentó enmiendas (junio de 2023) destinadas a reforzar la protección de los derechos fundamentales; que algunas de ellas fueron rechazadas en los trilogos (diciembre de 2023); que con la aprobación del AI Act (2024) se inició el complejo trabajo de traducción del texto a las lenguas oficiales de la Unión, de delicada importancia debido a su aplicabilidad directa y obligatoria en todos los Estados miembros.

(2) LA DEFINICIÓN JURÍDICA DE IA

Posteriormente, Presno abordó el tema del concepto jurídico de IA, un punto muy debatido durante la aprobación del AI Act, debido al hecho de que entre los diversos documentos oficiales publicados a nivel mundial figuran al menos 55 definiciones distintas.

El legislador europeo ha adoptado una definición que no se aleja demasiado de las utilizadas en otros contextos internacionales (véanse, por ejemplo, las adoptadas por la OCDE y el Consejo de Europa).

En esencia, el Reglamento define la IA como un sistema basado en máquinas capaz de inferir y, a partir de esta inferencia, hacer predicciones sobre eventos futuros, elaborar contenidos, dar recomendaciones o incluso tomar decisiones. La **inferencia, la intelectualidad y la autonomía** son los elementos clave. Por lo tanto, los sistemas que no poseen estas características, aunque sean algorítmicos, no estarán sujetos al Reglamento (por ejemplo, VioGén, utilizado en España para prever el riesgo de reincidencia en casos de violencia de género).

Esta concepción fundamental de la IA se refleja en una de las principales particularidades del Reglamento, que es que considera la IA tanto como objeto de regulación (por ejemplo, con prácticas prohibidas o permitidas) como sujeto de regulación, pudiendo constituir o influir en situaciones jurídicas subjetivas (por ejemplo, un sistema que decide la contratación de un trabajador).

(2) LOS PUNTOS CRÍTICOS DEL AI ACT

Finalmente, Presno destacó lo que, a su juicio, son los principales problemas del AI Act.

En primer lugar, el Reglamento fue concebido para ser **dinámico**, es decir, con la previsión de mecanismos de actualización periódica.

Sin embargo, por un lado, esto confiere a la **Comisión** un **poder** significativo en la aprobación de directrices (por ejemplo, sobre prácticas prohibidas) y actos delegados; por otro lado, no está garantizado que la promulgación de actos para interpretar un Reglamento ya de por sí muy extenso pueda efectivamente clarificar su alcance preceptivo.

Al mismo tiempo, **el AI Act entra en vigor por fases**: algunas partes ya están en vigor (por ejemplo, las prácticas prohibidas); otras no lo estarán antes de 2030.

Todo esto, unido a la rapidez del desarrollo del sector, plantea serias dudas sobre qué Reglamento estará realmente en vigor en los próximos años.

Se han planteado dudas adicionales con referencia a la previsión de la necesaria **supervisión humana** para los sistemas de alto riesgo, debido a la improbable capacidad humana de supervisar sistemas complejos sin ser influenciados por ellos.

Finalmente, el **modelo de gobernanza de la IA** basado en la garantía de los derechos choca con las exigencias del mercado único y la competitividad.

El Reglamento trata la IA como un **producto industrial**, centrándose solo en los riesgos significativos para los derechos fundamentales. Siempre desde esta perspectiva, parece dudoso el amplio margen de **autoevaluación** del riesgo de los productos concedido a los proveedores de los sistemas.

Asimismo, parece dudosa la exclusión del ámbito de aplicación de los sistemas militares y de investigación científica. De hecho, aunque motivadas respectivamente por razones de competencia comunitaria o de promoción del conocimiento, estas exclusiones podrían pasar por alto el impacto en los derechos en estos sectores.

Sobre todo, concluye Presno, parece paradójico que las exigencias del mercado permitan **exportar** (incluso a Estados que no reconocen el respeto de los derechos humanos) aquellos mismos sistemas de IA que están prohibidos en la UE.

JOSÉ LUIS BOLZAN DE MORAIS, profesor de derecho constitucional y digital en la Universidad de Vitória (Brasil), ilustró el contexto normativo en el que se sitúa el fenómeno de **la IA en Brasil**.

El punto de partida de la reflexión de Bolzan fue la constatación del **límite de los Estados** nacionales territoriales en la regulación de fenómenos tecnológicos como Internet y la IA, que tienen un alcance no territorial. Esto no significa, para Bolzan, abdicar de la regulación: esta es necesaria para equilibrar las exigencias del mercado con la tutela de los derechos, lo cual no es concebible con la autorregulación.

Dicho esto, el ponente destacó cómo el ordenamiento jurídico brasileño se caracteriza por la presencia de diversos actos y disposiciones destinadas a regular el fenómeno en cuestión.

En primer lugar, se mencionó el **habeas data** del art. 5, co. LXXII, de la Constitución brasileña de 1988, que, aunque concebido antes de la era de Internet, constituye un precursor de la normativa del sector, garantizando el derecho de acceso a la información personal en bases de datos públicas. El verdadero punto de inflexión para el ordenamiento jurídico brasileño debería haber sido el **Marco Civil da Internet** de 2014, considerado en el momento de su entrada en vigor como la "Constitución de Internet". Sin embargo, después de 10 años, la experiencia de esta normativa demostraría, según

Bolzan, lo afirmado al principio, es decir, la incapacidad de los Estados territoriales para regular fenómenos de esta magnitud.

En particular, en cuanto a la **responsabilidad de las plataformas digitales**, el modelo del Marco Civil impone la obligación al proveedor de eliminar un determinado contenido solo después de una solicitud del usuario a la que haya seguido una orden judicial de eliminación. Este mecanismo ha mostrado todas sus limitaciones con los problemas de desinformación relacionados con las elecciones presidenciales de 2018.

La **justicia electoral** en Brasil ha mostrado mayor agilidad: posee poderes normativos y administrativos que le han permitido regular las elecciones de 2022, promoviendo la colaboración con las plataformas para gestionar la desinformación, sin necesidad de esperar una orden judicial. También se ha implementado una herramienta ("Pardal") para la notificación administrativa de casos de desinformación.

Al mismo tiempo, existen algunos proyectos legislativos para modificar el régimen de responsabilidad de las plataformas, pero han sido bloqueados por campañas de *lobbying* por parte de las mismas plataformas.

En cuanto a la **protección de datos**, Brasil cuenta desde 2020 con la **Lei Geral de Proteção de Dados Pessoais (LGPD)**, que reproduce la estructura del GDPR europeo.

Aún se encuentra en fase de discusión un **proyecto de ley marco sobre Inteligencia Artificial (PL 2338/2023)**, modelado a semejanza del AI Act europeo.

Finalmente, el profesor Bolzan subrayó el amplio proceso de digitalización e informatización que ha afectado al sistema judicial brasileño, el cual, en esta fase, se está caracterizando por un uso cada vez mayor de la **IA en los tribunales**, tanto regulado por las resoluciones del Consejo Nacional de Justicia (n.º 332 y 615 de 2015), como informal.

JAVIER FERNÁNDEZ RODRÍGUEZ, director de la *Dirección General de Estrategia Digital e Inteligencia Artificial del Principado de Asturias*, concluyó la serie de intervenciones de la mesa redonda abordando el tema de la **regulación del uso de la IA en la administración del Principado de Asturias**.

Según Fernández, la IA representa un desafío para la administración pública debido a la expectativa de la ciudadanía de utilizarla para mejorar la eficiencia de los servicios y la lentitud general de los aparatos burocráticos para adaptarse a cambios tecnológicos tan rápidos e impactantes.

Para que la administración pública pueda explotar todo el potencial de la IA, es necesario adoptar una **estrategia** estructurada, que incluya en su visión los objetivos y prioridades de la administración pública, así como las experiencias del mercado y de las empresas privadas.

Solo así sería posible aprovechar y adaptar la IA a las **necesidades concretas de la administración**.

En este sentido, la introducción de la IA en la administración pública requiere la **adaptación de la organización y del personal**: es necesario un plan de recursos humanos que considere el impacto en los perfiles profesionales (algunas tareas podrían volverse superfluas, requiriendo recalificación o adaptación). Se necesitan directivos dedicados, un grupo de trabajo especializado, el apoyo de expertos externos y un sistema de gobernanza adecuado.

Para abordar estos desafíos, el Principado de Asturias está adoptando un **decreto interno para regular el uso de los sistemas de IA**, una herramienta para garantizar el respeto de las buenas prácticas y la reducción de riesgos (como, por ejemplo, el incumplimiento de la normativa de rango superior), proporcionando así seguridad jurídica a los proyectos de la administración pública sobre IA.

Más en detalle, el decreto, elaborado con la aportación de un grupo de trabajo multidisciplinar, tiene como objetivo regular el uso de los sistemas de IA, estableciendo criterios para la adquisición, desarrollo, implementación, seguimiento y evaluación continua. Además, tiene como objetivo definir el mecanismo de gobernanza de la administración autonómica, así como crear sandbox normativos para permitir a las empresas probar soluciones de IA en un entorno controlado.

En conclusión, Fernández ilustró los diversos **casos de uso** y proyectos piloto o en curso en la administración del Principado de Asturias, como Hoga, un coordinador virtual para los derechos sociales que, interactuando con los usuarios, permite proporcionar un nivel de atención al ciudadano personalizado que no sería posible obtener con recursos humanos.

EDUCACIÓN E IA

Ponentes: Javier **Fombona**, Julio **Albalad Gimeno**, Thomas **Casadei**

Moderador: Marián González Rúa

Contenido: La segunda mesa redonda del miércoles abordó el debate entre los ponentes sobre el potencial y el impacto de la IA en el sector de la educación.

JAVIER FOMBONA, profesor de ciencias de la educación en la UNIOVI y miembro del CEISIA, inició su intervención lanzando la idea de que la IA no es una simple herramienta, sino un compañero capaz de interactuar con los seres humanos.

Destacó cómo la IA está en constante cambio y rápido progreso en muchísimos sectores, respecto a los cuales, sin embargo, el de la educación sigue siendo marginal.

En concreto, en este contexto, la IA permite una **gestión inteligente de la web**, permitiendo a los estudiantes obtener directamente documentos y resúmenes de diversas fuentes sin acceder a los sitios tradicionales.

Las capacidades actuales de la IA para generar imágenes y procesar el lenguaje natural, interactuando así con el usuario, permiten a los estudiantes utilizarla para obtener traducciones, redactar textos académicos y crear materiales educativos.

De la integración de la IA con la **realidad aumentada y virtual** están surgiendo plataformas educativas capaces de generar unidades didácticas en las que es posible aprender de manera innovadora, por ejemplo, conversando con personajes históricos. Además, están emergiendo plataformas con **aprendizaje adaptativo**, que partiendo de un análisis individualizado de las capacidades del estudiante, personalizan la estrategia de aprendizaje según sus necesidades. Esto permite implementar **metodologías diferenciadas** que tienen en cuenta el ritmo y el estilo de aprendizaje individual, incluso en clases con muchos estudiantes.

Una eventual evolución de la educación en esta dirección corre el riesgo, al mismo tiempo, de reducir la autonomía de los docentes, quienes, si no desean delegar totalmente ciertas decisiones de aprendizaje al algoritmo, deben someterse a una actualización constante.

JULIO ALBALAD GIMENO, director del “*Instituto Nacional de Tecnologías Educativas y de Formación del Profesorado*” (INTEF) del *Ministerio de Educación, Formación Profesional y Deportes*, centró su intervención en tres aspectos principales: (1) el uso de la IA por parte de alumnos

y docentes; (2) los riesgos de la IA en la educación; (3) las iniciativas del Ministerio de Educación sobre la IA.

(1) EL USO DE LA IA POR PARTE DE ALUMNOS Y DOCENTES

En referencia al uso de la IA por parte de los **alumnos**, el ponente explicó que la nueva ley educativa (LOMLOE 2022) ha reformado los currículos de los estudiantes basándolos en las competencias, entre las cuales figura la **competencia digital**, que debe adquirirse al finalizar la educación secundaria obligatoria. La competencia digital implica un uso saludable y sostenible de las tecnologías, incluida la IA, que, sin embargo, solo se menciona en este contexto con referencia a materias específicas en los niveles superiores. El objetivo del Ministerio es lograr una **educación sobre la IA** (sobre su funcionamiento) y **para la IA** (para su uso crítico).

En cuanto a los **docentes**, el Marco de Referencia de la Competencia Digital Docente (2022) menciona específicamente la IA tanto para el desarrollo de las competencias digitales de los alumnos como para la promoción de su uso crítico por parte de los docentes.

(2) LOS RIESGOS DE LA IA EN LA EDUCACIÓN

El Ministerio de Educación ha publicado una **guía sobre el uso de la IA en la educación no universitaria** (2024) que analiza riesgos y desafíos para estudiantes, docentes y centros educativos, proponiendo estrategias.

Los mayores **riesgos para los estudiantes** son la **falta de competencias digitales** y de **posibilidades de acceso tecnológico**, especialmente para aquellos provenientes de familias desfavorecidas.

La mayor preocupación, sin embargo, la representan los **datos**: el uso de la IA por parte de los estudiantes, si es incentivado por los docentes, debe ir acompañado de la concienciación sobre los riesgos relacionados con la provisión de datos personales, los sesgos y las alucinaciones algorítmicas.

Para los **docentes**, el riesgo principal es la **falta de formación**, dada la edad media elevada.

En general, se debería animar al docente a ver la IA como una **herramienta de apoyo, no sustitutiva**.

En particular, el ponente subraya que el uso de la IA en la educación se considera de alto riesgo según el AI Act, por lo que, por un lado, toda evaluación debe permanecer bajo el control humano del docente; por otro lado, prácticas como el reconocimiento facial de las emociones de los estudiantes, utilizado en China, están prohibidas.

(3) LAS INICIATIVAS DEL MINISTERIO DE EDUCACIÓN SOBRE LA IA

Entre los proyectos piloto de IA del Ministerio de Educación, el ponente mencionó: el desarrollo de **chatbots para las familias** (para apoyarlos en los procesos administrativos más complejos); la

creación de **herramientas para la creación de recursos didácticos para los docentes** basados en contenidos validados para garantizar su calidad; la adaptación de los materiales para estudiantes con necesidades específicas.

THOMAS CASADEI, profesor de filosofía del derecho en la Universidad de Módena y Reggio Emilia, compartió reflexiones a partir de los resultados de un proyecto que él mismo dirigió, centrado en el **uso consciente de la red, las redes sociales y las nuevas tecnologías**, como la IA, por parte de las jóvenes generaciones.

El proyecto dirigido por Casadei parte de una encuesta estadística de Eurostat según la cual en Europa **el 10% de los jóvenes puede desarrollar comportamientos problemáticos o de riesgo cuando navega por la red**, un porcentaje contenido pero significativo que requiere prevención.

El proyecto se caracteriza por un enfoque interdisciplinar, reuniendo derecho, informática y ciencias sociales, y se centra en la relación entre ciberseguridad, nuevas tecnologías y protección de derechos.

En particular, las cuatro líneas de investigación son: (1) plataformas digitales, con foco en **TikTok**, la más usada por los menores; (2) **discriminaciones digitales** y *online hate speech*, con el objetivo de identificar los grupos más vulnerables; (3) educación de los menores sobre los riesgos relacionados con los **delitos informáticos**; (4) las buenas prácticas y los proyectos de escuelas, ciudades y universidades sobre la **protección de datos**.

Los principales resultados esperados son: (1) la redacción de una **guía para el uso consciente de las nuevas tecnologías** dirigida a todo tipo de educador, en la que se abordarán temas como *fake news*, *revenge porn*, ciberacoso, *hikikomori* y *dark web*; (2) la **creación de un observatorio**; (3) la apertura de una ventanilla de ayuda.

A la luz de los resultados de sus estudios, Casadei concluyó subrayando la necesidad de un enfoque sistémico ante el impacto de las tecnologías, considerando tanto el aspecto jurídico como los aspectos escolar, psicológico, relacional, sanitario y político.

Además, se mostró dubitativo respecto a la poca atención dedicada por el AI Act al sector de la educación, que, por el contrario, podría beneficiarse de un documento distinto sobre la **cultura y la conciencia digital a nivel europeo**.

CIBERSEGURIDAD

Ponentes: María Josefa **Ridaura** Martínez, Luis **Molina** Valbuena, Stefano **Pietropaoli**

Moderador: Josè Manuel Redondo

Contenido: La mesa redonda de la tarde sobre ciberseguridad, celebrada en la sede del CEISIA en Lugones, abordó las cuestiones problemáticas relacionadas con: los sistemas de reconocimiento facial (Ridaura), la identidad digital (Molina) y la relación entre el derecho y las nuevas tecnologías (Pietropaoli).

MARÍA JOSEFA RIDAURA MARTÍNEZ, profesora de derecho constitucional en la Universidad de Valencia, abrió la sesión abordando el tema de los **sistemas de reconocimiento facial**.

En particular, Ridaura destacó el distinto impacto en los derechos fundamentales entre los sistemas de verificación y los sistemas de identificación.

De hecho, los **sistemas de verificación** (uno a uno), como los utilizados para el desbloqueo del teléfono, presentan menos problemas de derechos fundamentales que los **sistemas de identificación** (uno a muchos), ya que mientras los primeros comparan la imagen capturada con otra previamente insertada, los segundos la comparan con una base de datos completa.

La cuestión fundamental que plantea el reconocimiento facial de identificación es la del **equilibrio entre los derechos fundamentales de los ciudadanos** afectados por el uso de la herramienta y la **seguridad pública**.

Según Ridaura, libertad y seguridad no están en tensión dialéctica, sino que ambas son necesarias para el ejercicio de los derechos fundamentales. Al mismo tiempo, las herramientas empleadas no deben limitarse a garantizar genéricamente la seguridad, sino una **seguridad democrática**, respetuosa con los derechos fundamentales.

Los sistemas de reconocimiento facial de identificación son actualmente muy potentes y están bastante extendidos (por ejemplo, en China, EE. UU., Colombia), pero plantean graves problemas para los derechos fundamentales, especialmente cuando se utilizan en la vía pública.

De hecho, en esta circunstancia pueden violarse no solo el **derecho a la privacidad**, sino también el **derecho a la libertad de circulación** y el **derecho de reunión**, ya que la presencia de estas herramientas puede crear un **efecto disuasorio** en la ciudadanía.

Además, estos sistemas pueden violar la **presunción de inocencia** en caso de **errores de identificación**, muy a menudo correlacionados con casos de **discriminación algorítmica** debidos a un entrenamiento en bases de datos no suficientemente representativas de categorías específicas (por

ejemplo, mujeres de color). Bases de datos creadas con imágenes de dudosa **procedencia y fiabilidad**, a menudo obtenidas de empresas privadas o de redes sociales.

En casos de error macroscópico, el error de identificación podría ir seguido de una detención, con la consiguiente e injustificada limitación de la libertad personal.

Esta incidencia en los derechos se agrava por el hecho de que el uso de estos sistemas en lugares públicos se realiza **sin el consentimiento** de los sujetos implicados. En los casos más extremos, esto puede conducir a escenarios de **vigilancia masiva**, difícilmente justificables en una sociedad democrática.

Una cuestión crucial es la relativa a la **proporcionalidad** de un sistema de control generalizado destinado a identificar a individuos concretos. En un estado de derecho, las intromisiones en la privacidad por seguridad deben tener justificación legal, objetivo legítimo, necesidad y proporcionalidad, como establecen la jurisprudencia del Tribunal Europeo de Derechos Humanos (TEDH) y del Tribunal de Justicia de la UE.

Con la entrada en vigor del AI Act, el uso de estas herramientas por parte de la policía solo se permite en **casos específicos y con muchas garantías**, pero resulta problemática la exclusión del ámbito de aplicación de los sectores de: defensa, seguridad nacional e inteligencia.

LUIS MOLINA VALBUENA, consultor legal de Castroalonso especializado en derecho digital, dedicó su intervención al tema de la **identidad digital** y los peligros derivados de los **contenidos sintéticos** hiperrealistas, los cuales pueden crearse actualmente con facilidad y sin conocimientos técnicos específicos.

Por identidad digital se entiende el conjunto de todos los datos atribuidos a una persona en el entorno digital, ya sean compartidos voluntariamente o recopilados sin consentimiento. Basándose en los datos de la identidad digital, incluso los aparentemente más inofensivos como un audio o un video de pocos segundos o una foto, es posible clonar la voz de una persona o **crear videos hiperrealistas**. Tales contenidos pueden generar daños enormes, que, aunque no constituyan tipos penales específicos, son idóneos para destruir reputaciones personales y carreras profesionales.

Nos encontraríamos, por tanto, ante un **vacío legal** en la protección de la **identidad digital auténtica**. A pesar de las prescripciones del GDPR, las garantías generales parecen escasas y la represión de las conductas criminales ineficaz.

Así, en el **derecho penal** falta una legislación específica contra la generación fraudulenta de identidades sintéticas. Al mismo tiempo, sigue abierta la cuestión de la **responsabilidad** de quien utiliza fraudulentamente las tecnologías en examen, que, en opinión de Molina, debería recaer en el

autor de la conducta, no en el creador. Además, se plantea el problema de la prueba de la generación de los contenidos sintéticos, ya que podría realizarse sin dejar rastros digitales.

Las soluciones propuestas por Molina para afrontar los problemas planteados son: (1) el **reconocimiento legal de la identidad digital auténtica**, a reconocer como un bien jurídico a proteger; (2) la **institución de un registro voluntario de identidades verificadas**; (3) la **actualización del código penal** para tipificar específicamente el uso malicioso de contenidos sintéticos; (4) la instrucción de la ciudadanía; (5) la adopción de buenas prácticas de seguridad.

STEFANO PIETROPAOLI, profesor de filosofía del derecho en la Universidad de Florencia, cerró la mesa redonda con una intervención centrada en los desafíos que plantean las nuevas tecnologías para el derecho.

El punto de partida de la intervención fue la constatación de que las nuevas tecnologías, aunque se manifiestan en una dimensión digital, generan un impacto real, y no virtual, en la vida y los derechos fundamentales de las personas.

Sin embargo, al abordar estos nuevos desafíos, el derecho se encuentra lidiando con una **crisis de los conceptos jurídicos tradicionales** (por ejemplo, propiedad, subjetividad, nexo causal) que ya no son adecuados para representar los fenómenos de la nueva dimensión digital. Según Pietropaoli, esta crisis de los conceptos se traduce en una **crisis del derecho mismo**. El derecho, como recordaba Hermogeniano, es una creación del ser humano para el ser humano ("*Omne ius hominum causa constitutum est*"). Los códigos normativos son tales porque es posible desobedecerlos, a diferencia de los códigos informáticos a los que las máquinas no pueden manifestar disenso.

El espacio cibernético es un entorno en rápida evolución donde los derechos de los ciudadanos se encuentran en una condición de vulnerabilidad, ya que carecen de una formación adecuada y están a merced de los grandes actores privados del mundo de la tecnología. En resumen, no puede haber una tutela efectiva de los derechos si no se comprenden los mecanismos, los riesgos y las lógicas de la informática. La **ciudadanía digital** requiere saber cómo y por qué nos conectamos y qué consecuencias pueden derivarse de interacciones aparentemente neutrales.

Finalmente, Pietropaoli abordó el tema de la **explicabilidad algorítmica**, sosteniendo la absurdidad de la búsqueda de los motivos de los resultados de los algoritmos. Estos carecen de una racionalidad en el sentido humano: se basan en la estadística, en la lógica inductiva, no deductiva. A partir de ello, es posible polemizar contra la idea misma de decisión algorítmica. De hecho, según el ponente, las decisiones algorítmicas no serían decisiones: las máquinas no deciden, porque no pueden expresar una voluntad. No hay que antropomorfizar. En todo caso, existe, en un nivel superior, la decisión humana de utilizar una herramienta para resolver un problema.

JUEVES 15 DE MAYO 2025

IA Y TRABAJO

Ponentes: Ivan Antonio **Rodríguez** Cardo, Lucia **Montejo**, Javier **Campa**, Isabel **Santos**

Moderador: Maria Antonia **Castro** Arguelles

Contenido: Las actividades del jueves 15 de mayo de 2025 se iniciaron en el aula "Severo Ochoa" del edificio histórico de la Universidad de Oviedo con la mesa redonda sobre "*IA y Trabajo*", en la que se confrontaron representantes del mundo académico (Rodríguez), sindical (Montejo, Campa) y empresarial (Santos).

IVAN ANTONIO RODRÍGUEZ CARDO, profesor de derecho del trabajo en la UNIOVI y miembro del CEISIA, ofreció una visión general desde la **perspectiva jurídica** del impacto de la IA en las relaciones laborales.

Para este propósito, la exposición del profesor Rodríguez partió de algunas consideraciones de la economía de la IA.

En particular, desde el punto de vista empresarial, la IA se percibe como una oportunidad para hacer más eficientes y económicos los procesos productivos. En esta línea, puede utilizarse para organizar las tareas de los trabajadores y para tomar decisiones autónomas (la llamada "*gestión algorítmica del trabajo*"). Además, existiría cierta expectativa de que pueda mejorar el empleo general, dando lugar a la creación de nuevos puestos de trabajo.

Sin embargo, desde el punto de vista jurídico, la IA constituye un fenómeno a regular para proteger los derechos de los trabajadores, para evitar la deshumanización de las decisiones laborales y, sobre todo, para lograr una transición justa que no conlleve la pérdida repentina de puestos de trabajo.

A esto se suman los riesgos relacionados con un control excesivo de los trabajadores, el tratamiento de los datos personales, y los sesgos que conducen a discriminaciones difícilmente verificables debido a la opacidad de los sistemas de IA.

Con respecto a estas cuestiones, el AI Act parece "tímido" desde una perspectiva laboral. De hecho, no presenta normas laborales *stricto sensu*, en el sentido de regular la relación entre trabajador y empleador, sino una norma sobre la seguridad del producto de IA. En otras palabras, se limita a disciplinar la producción de un sistema de IA, independientemente de la calificación del usuario final como trabajador.

No obstante, algunas de las prácticas prohibidas por el reglamento tienen un impacto directo en el mundo laboral. Así, no se permite el uso de sistemas de **reconocimiento de emociones** y de **categorización biométrica para deducir datos sensibles** (como, por ejemplo, la afiliación sindical).

Al mismo tiempo, las normas sobre transparencia y supervisión humana dictadas para los sistemas de alto riesgo no prevén (directamente) que los trabajadores o sus representantes puedan participar en las actividades de control de la IA.

Diferente es el caso de la disciplina de la **Directiva europea sobre el trabajo en las plataformas digitales**, que sí prevé la participación de los trabajadores en las obligaciones de transparencia y supervisión de la IA.

Finalmente, el **artículo 64 del Estatuto de los Trabajadores español** exige que el comité de empresa tenga derecho a ser informado por la empresa sobre los parámetros, las reglas y las instrucciones en las que se basan los algoritmos o los sistemas de IA que influyen en el proceso de toma de decisiones que pueda tener un impacto en las condiciones de trabajo, el acceso y el mantenimiento del empleo, incluida la elaboración de perfiles.

A pesar de esto, Rodríguez manifestó dudas respecto a un enfoque de protección centrado principalmente en la transparencia. De hecho, por un lado, persiste el problema de la opacidad algorítmica; por otro, a menudo los empresarios adquieren el sistema de terceros o externalizan su gestión. En lugar de centrarse en la apertura de la caja negra, el ponente abogó por controles centrados en los resultados que verifiquen la corrección de los resultados de la IA, recurriendo, por ejemplo, a la prueba estadística.

LUCIA MONTEJO, miembro de la “*Confederación Sindical de Comisiones Obreras*”, participó en la discusión aportando la perspectiva de los trabajadores que sufren los efectos del empleo de la IA.

Al abordar la cuestión, Montejo recordó que la IA no debe considerarse un bloque monolítico, sino que engloba diferentes tecnologías. Además, subrayó que las empresas utilizan algoritmos desde hace mucho tiempo, pero solo recientemente se está asistiendo a una verdadera "hiper-automatización".

Dado que las aplicaciones de la IA en el mundo laboral son diversas, involucran de diferentes maneras los derechos de los trabajadores y, por lo tanto, necesitan **estrategias sindicales dedicadas**.

Según Montejo, el marco normativo actual ofrece protecciones laborales inadecuadas: por un lado, el AI Act no se centraría adecuadamente en la protección de los derechos de los ciudadanos; por otro, el artículo 64 del Estatuto de los Trabajadores español ofrece una protección limitada. En resumen, es inútil recibir de la empresa 200 páginas de código fuente del algoritmo si las asociaciones sindicales no disponen de los medios y el personal para hacerlo **inteligible**. Por el contrario, es necesario que la

obligación de transparencia se cumpla proporcionando una explicación efectiva del funcionamiento algorítmico.

Además, no todos los problemas pueden resolverse conociendo los mecanismos de decisión algorítmicos.

A estos **problemas**, que derivan específicamente del uso de la IA en el contexto laboral, se suman los **estructurales** del sistema productivo y del mercado de trabajo que, sin embargo, se ven amplificados por la IA. Por ejemplo, incluso en presencia de un conjunto de datos de entrenamiento robusto, la IA podría replicar los sesgos (discriminatorios) preexistentes en la sociedad.

Actualmente, subrayó Montejo, los sindicatos parecen poco preparados para gestionar cambios tan rápidos, ya que la alfabetización digital de los representantes requiere tiempo y apoyo institucional.

En términos más generales, el auge de la IA está ligado a un proceso de aceleración de la concentración de la riqueza a nivel mundial, con repercusiones en las relaciones laborales actuales, en el futuro del trabajo y en la sociedad en general. Esto, concluyó Montejo, obligará a abordar la cuestión de la fiscalidad de las tecnologías.

JAVIER CAMPA, miembro de “*Unión General de Trabajadoras y Trabajadores*”, reconoce que la IA puede ser portadora tanto de oportunidades como de riesgos para los trabajadores.

Puede facilitar las condiciones laborales, pero también sacrificar numerosos puestos de trabajo en nombre del beneficio. En otras palabras, la IA tiene el potencial de realizar una **reconversión global** del mundo del trabajo, similar a las anteriores revoluciones industriales.

Precisamente por esto, Campa subrayó la necesidad de un marco normativo que regule la implementación de la IA en las empresas de forma transparente, que involucre a los trabajadores en la gobernanza y que tutele sus derechos. Sobre todo, concluyó Campa, es necesario que se realice una transición justa, que permita una recalificación de los trabajadores mediante su formación continua y que no deje atrás a quienes no estén en condiciones de poder recalificarse.

ISABEL SANTOS, miembro de “*Federación Asturiana de Empresarios*”, enriqueció el debate con la aportación del punto de vista empresarial.

Para las empresas, subrayó Santos, la IA no es el futuro, sino el presente. Están obligadas a invertir en esta tecnología para no quedarse atrás. Para gestionar mejor este rápido cambio, es necesaria una **coordinación entre empresa y trabajadores**.

La IA constituye una evolución natural de los procesos productivos, interviniendo en sectores específicos.

Algunos podrían ser completamente sustituidos, como los trabajos administrativos basados en papel o los manuales básicos.

Otros podrían utilizar la IA como un apoyo de mejora (trabajos creativos).

Finalmente, la IA creará nuevos puestos de trabajo en el sector tecnológico.

Santos concluyó planteando algunos de los riesgos de la transición digital, como la brecha entre países capaces de convertir y formar su fuerza laboral cualificada, eventualmente con el apoyo de subvenciones públicas específicas, así como el de depender de terceras empresas para el suministro y la gestión de la IA.

IA Y SALUD

Ponentes: Sergio **Calleja**, Elena **Díaz** Rodríguez, Pedro José **Espinosa** Prados

Moderare: Daniel Jove Villares

Contenido: La segunda mesa redonda de la mañana del 15 de mayo confrontó al sector sanitario (Calleja y Espinosa) y a la academia (Díaz) sobre las **aplicaciones más recientes de la IA en la medicina**.

SERGIO CALLEJA, neurólogo en el “Hospital Universitario Central de Asturias” (HUCA), contribuyó a la discusión relatando su **experiencia clínica** personal.

Calleja abrió la intervención describiendo la medicina como una relación entre quien presenta un problema (el paciente) y quien busca resolverlo (el médico). Una relación que se basa en la **incertidumbre**, obviamente, del paciente respecto a su propia patología, pero también del médico acerca del diagnóstico y el tratamiento. Un concepto que, recuerda Calleja, es sintetizado eficazmente por las palabras de William Osler: "La medicina es la ciencia de la incertidumbre y el arte de la probabilidad". Esta incertidumbre pesa mucho en nuestra sociedad: los errores médicos, recordó Calleja, son la tercera causa de muerte más común en un país como EE. UU. Sin embargo, esta circunstancia parece fisiológica. De hecho, multiplicando la incertidumbre por el número de pacientes que cada médico visita cada día, se hacen evidentes las dificultades de la profesión médica.

Por lo tanto, la medicina siempre ha acogido las innovaciones capaces de ayudarla a reducir este margen de incertidumbre. La IA podría potencialmente ayudar en esta dirección, contribuyendo a reducir aquellos errores que no son fruto de negligencia, sino de la complejidad de la profesión misma.

La IA pretende proporcionar soluciones directamente.

Sin embargo, es necesario reflexionar sobre lo que se esconde detrás de la demanda de atención del paciente y la respuesta del médico. Este último genera hipótesis, las confronta con la realidad, recopila datos, evalúa su origen y los posibles sesgos. Se enfrenta a una complejidad gracias a su propia experiencia, que no es ni predeterminada ni enseñable en la universidad. Un médico experto actúa por una intuición que se basa en la exposición repetida a casos reales.

Al recibir respuestas directamente de la IA, la adquisición de conocimiento por experiencia del profesional sanitario corre el riesgo de diluirse.

El método basado en la evidencia empírica ha transformado la medicina de un saber precientífico a una ciencia, lo que ha llevado tanto a indudables progresos como a subestimar aquellas capacidades no numéricas intraducibles en fórmulas matemáticas como la empatía, la habilidad de recopilar datos cualitativos y la intuición. Se estima que casi el 50% de los síntomas declarados por los pacientes son

médicamente inexplicables. Además, para muchos pacientes nunca se llega a un diagnóstico. En estos casos, el papel del médico no es encontrar la respuesta, sino apoyar al paciente.

En los últimos años, recordó Calleja, la medicina ha terminado siendo influenciada por las presiones de los intereses económicos: de las industrias farmacéuticas, de la sanidad privada, de los titulares de patentes. Todos interesados en obtener la mayor desregulación posible en el sector.

La calidad de la ciencia médica se resiente de la demanda de generar beneficios a corto plazo, por ejemplo, con prescripciones farmacéuticas inútiles o con cribados indiscriminados.

Esta digresión de contexto parece fundamental para interrogarnos sobre quién financia y a qué intereses están subordinados los sistemas de IA en el sector de la sanidad.

Finalmente, concluyó Calleja, hay que preguntarse: ¿quién controlará los sesgos? ¿Quién será responsable en caso de error? ¿Cómo se formarán los nuevos profesionales cada vez menos expuestos a la experiencia clínica empírica?

ELENA DÍAZ RODRÍGUEZ, profesora de fisiología en la UNIOVI y miembro del CEISIA, abrió su intervención destacando cómo en el campo de la sanidad ya se utilizan muchísimas aplicaciones de IA.

Uno de los primeros usos de la IA en el sector, recuerda Díaz, fue en el ámbito oncológico, donde los algoritmos son capaces de analizar las biopsias obteniendo resultados superiores a los de los especialistas. Por lo tanto, una herramienta similar podría tanto aumentar la precisión de los diagnósticos como reducir la carga de trabajo de los médicos.

Díaz presenta dos estudios de aplicación de la IA a la **cronobiología de la reproducción**.

Por cronobiología se entiende la disciplina que estudia los ritmos circadianos: aquellos procesos fisiológicos, bioquímicos, genéticos que se repiten cada 24 horas (por ejemplo, sueño-vigilia, comidas). El ritmo es dictado por un reloj biológico en el sistema nervioso central, influenciado por la luz a través de la melatonina. La alteración de estos ritmos se denomina **cronodisrupción**. Causas comunes: el uso de pantallas por la noche, los viajes y los horarios de trabajo irregulares. La investigación se centra en cómo la cronodisrupción influye en la salud reproductiva, en particular el parto prematuro.

El **primer estudio** tenía como objetivo identificar nuevos factores de riesgo para el parto prematuro relacionados con la cronodisrupción y desarrollar un modelo matemático para predecir el riesgo. Se recopilaron datos de 481 partos, de los cuales 257 fueron prematuros. Se consideraron variables clínicas habituales y variables relacionadas con la cronobiología obtenidas a través de cuestionarios telefónicos (hábitos y calidad del sueño, horarios de sueño en días laborables y festivos,

dormir con o sin luz en la habitación). La muestra final para el algoritmo fue de 223 partos. La fiabilidad del modelo resultó alta, con un porcentaje de éxito del 97,3%.

El resultado más importante del estudio es la clasificación de las variables. El factor más importante resultó ser el peso estimado, un elemento ya conocido, lo que confirmó la validez del algoritmo. En cambio, la segunda variable más importante identificada por la IA fue, sorprendentemente, la luz en la habitación durante la noche, superando en importancia variables clásicas como la edad de la madre. El árbol de decisión construido por el algoritmo ayuda a los clínicos a estimar el riesgo de parto prematuro a partir de las variables: la IA ordena los **factores de riesgo** por importancia, proporcionando una ayuda para un **diagnóstico rápido**.

El **segundo estudio** exploró marcadores de cronodisrupción en la salud reproductiva general de las mujeres, centrándose en una población muy expuesta a determinadas enfermedades: las enfermeras que trabajan por turnos. Los datos se recopilaban a través de cuestionarios en ocho áreas sanitarias y nueve hospitales de Asturias. Las preguntas se referían a datos sociolaborales, salud reproductiva, calidad del sueño y horarios de alimentación. Los problemas estudiados fueron: tiempo de concepción, duración de la gestación, problemas durante la gestación y enfermedades reproductivas generales. Se utilizaron tres algoritmos diferentes. No siempre estos últimos conducían a los mismos resultados, por lo que se procedió a su integración. Para el tiempo de concepción, la variable más importante resultó ser el punto medio de alimentación en los días de descanso, mientras que la segunda variable más importante fue la calidad del sueño. Para la duración de la gestación, la variable principal fue el "jet lag alimentario" (la diferencia en los horarios de las comidas entre los días laborables y los de descanso).

A la luz de sus propios estudios, Díaz concluyó afirmando que la IA constituye una herramienta de indudable utilidad para llevar a cabo investigaciones similares en el sector.

PEDRO JOSÉ ESPINOSA PRADOS, capo jefe de proyecto del espacio de datos de salud del “*Servicio de Salud del Principado de Asturias*” (SESPA), dedicó su intervención a la exposición de los resultados de un proyecto de creación de **dos repositorios centralizados de datos clínicos de los pacientes, actualizados en tiempo real**.

La procedencia de los datos de diferentes sistemas informativos sanitarios requirió un intenso trabajo de normalización de los mismos.

El repositorio **primario** sirve para el uso asistencial; el **secundario** está dedicado al uso de los datos para investigación, explotación, extracción de datos y análisis avanzado.

El repositorio primario permitirá alimentar un portal que proporcionará una **visión global y resumida del paciente**, adaptada al uso por parte de múltiples perfiles profesionales que podrán visualizar su historial y notas clínicas con una búsqueda similar a la de Google.

El repositorio secundario sigue un modelo de tres niveles: Bronze (datos brutos), Silver (datos transformados y orientados al modelo), Gold (datos validados y estandarizados).

Los datos se almacenan localmente en el SESPA, con la posibilidad de utilizar una nube para proyectos de investigación a gran escala o colaboraciones externas.

La identificación de los pacientes se reduce al mínimo mediante procesos de pseudoanonimización.

En esencia, se proporciona a los investigadores un **espacio de datos seguro** para procesar sus conjuntos de datos.

ECONOMÍA Y SOSTENIBILIDAD DE LA IA

Ponentes: Luis **Molina** Valbuena, Carlos **Ibáñez** Sánchez

Moderador: Augustina Bouchet

Contenido: La mesa redonda, celebrada en la sede del CEISIA en Lugones, se centró en el impacto de la IA en la sociedad, tanto en términos de sostenibilidad en un sentido amplio, que incluye la ética, los derechos fundamentales y la equidad social (Molina), como en términos de economía, con particular referencia a los costos y desafíos de las empresas emergentes en el sector de la legal tech (Ibáñez).

LUIS MOLINA VALBUENA, consultor legal de Castroalonso especializado en derecho digital, abrió la discusión sobre la cuestión de la sostenibilidad de la IA subrayando que este concepto no debe limitarse al medio ambiente, sino que también debe incluir los derechos humanos fundamentales.

En otras palabras, sostiene Molina, la **sostenibilidad** no puede medirse solo en términos de huella de carbono, sino también de ética, justicia algorítmica y responsabilidad.

En particular, la IA genera muchísimo valor, pero solo en beneficio de muy pocos individuos, que concentran en sí mismos un poder no democrático y sin contrapesos.

La fuente de este poder son los datos, "**el petróleo del siglo XXI**", extraídos directamente de los usuarios, que los proporcionan consciente o inconscientemente, y luego empaquetados, monetizados y revendidos.

Una economía digital, sostiene Molina, no puede ser sostenible si ignora la justicia social; no puede ser verde si es insaciable de energía, opaca en la gobernanza e ineficaz en los derechos. Así, en caso de errores algorítmicos, la ausencia de responsabilidad conduce a una economía injusta e inmoral. El entrenamiento de un LLM como GPT-4 produce un impacto ambiental comparable al consumo de energía de 100 familias europeas en un año y su uso produce "huellas de carbono invisibles".

Dicho esto, Molina excluye la posibilidad de perseguir la sostenibilidad algorítmica en ausencia de una **regulación jurídicamente vinculante**.

Los primeros pasos en esta dirección están representados por el AI Act y la **Directiva sobre Corporate Sustainability Due Diligence**, que impone la obligación de prevenir los daños ambientales y sociales a lo largo de la cadena de valor causados por sistemas automatizados y de adoptar políticas de debida diligencia efectivas y eficaces.

Molina reconoce el indudable costo de cumplimiento para las empresas, pero considera que para garantizarlo deben implementarse sanciones disuasorias.

A pesar de ello, Molina destaca la posibilidad de un **riesgo sistémico**: la propagación global de modelos opacos y entrenados con datos distorsionados. Así, por ejemplo, un modelo de IA puede bloquear el acceso al crédito o al trabajo, normalizando prácticas discriminatorias.

Basándose en esto, concluye Molina, gobernar la IA es la única manera de garantizar que la misma persiga fines beneficiosos para la sociedad.

CARLOS IBÁÑEZ SÁNCHEZ, CEO de ADEA, contribuyó a la discusión centrándose en la sostenibilidad de la IA desde el punto de vista económico-empresarial.

El objetivo de la sociedad ADEA, dirigida por Ibáñez, es proporcionar las mejores soluciones tecnológicas a los operadores legales. Para ello, se identificaron los cinco puntos clave que dichos sujetos buscaban en un servicio de este tipo: **eficiencia, precisión, costes contenidos, accesibilidad y transparencia.**

Para una empresa como ADEA, es fundamental mantener el ritmo de la evolución tecnológica, sin, sin embargo, limitarse a una actualización por sí misma, sino considerando cuál es la mejor tecnología posible en relación con los casos de uso reales.

Inicialmente, la empresa de Ibáñez utilizaba redes neuronales para procesar documentos legales. Esto requería el acceso a grandes volúmenes de datos, su clasificación y etiquetado, y el empleo de una alta capacidad computacional.

Con la llegada de los modelos fundacionales y la IA generativa, se ha pasado a un sistema híbrido: las redes neuronales se emplean para los entornos determinísticos (modelos, plantillas) porque son más rápidas y económicas; la IA generativa para los no determinísticos, que requieren interpretación y análisis del contexto.

Los **costos** del *machine learning* son mucho más contenidos que los de la IA generativa: el primero solo requiere una CPU, la segunda costosas GPU y un tiempo de entrenamiento mayor.

Sin embargo, la IA generativa permite generar datos sintéticos a partir de un pequeño conjunto de partida, lo que permite aumentar artificialmente el volumen de los datos de entrenamiento.

Después de describir los desafíos de desarrollo, las características y el funcionamiento del modelo de procesamiento de documentos legales desarrollado por ADEA, Sánchez concluyó abordando el tema de la **singularidad legal**.

En particular, a diferencia de la singularidad tecnológica, Sánchez considera que ya existen las herramientas para la automatización completa de muchísimos procedimientos legales. Visualiza la posibilidad de un derecho completo, sin lagunas, autoejecutable y autónomo. Sin embargo, esto requiere cierta gradualidad, para guiar de la mejor manera un eventual cambio en este sentido.

VIERNES 16 DE MAYO 2025

IA Y JUSTICIA

Ponentes: Jordi Nieva Fenoll, Alfredo Rodríguez del Blanco, Aitor Cubo.

Moderador: Marcos Loredó Colunga

Contenido: El aula "Severo Ochoa" del edificio histórico de la Universidad de Oviedo acogió la última mesa redonda sobre "IA y Justicia". Los ponentes exploraron la creciente integración de la IA en el sistema judicial, ilustrando: el potencial innovador para los procesos (Nieva); los riesgos para los derechos fundamentales en el proceso penal (Rodríguez); las recientes experimentaciones en la administración de justicia española (Cubo).

JORDI NIEVA FENOLL, profesor de derecho procesal en la Universidad de Barcelona, inició la discusión abordando el tema del **uso de la IA en los procesos** judiciales.

En este sentido, Nieva subraya la diferencia entre la automatización de los procedimientos y el uso de la IA en los procedimientos. Actualmente, recuerda Nieva, el sistema judicial se centra en la primera, con el objetivo de reducir sus costos. Sin embargo, este cambio debe ir necesariamente acompañado de una **modificación de las normas procesales**, de lo contrario obsoletas.

Los procedimientos no afectados por este cambio (la minoría) se caracterizan por el empleo necesario de todas las prerrogativas del jurista: reconstrucción de los hechos, valoración de la prueba, interpretación del derecho.

Al mismo tiempo, es precisamente en este ámbito donde la IA podría desplegar todo su potencial. Actualmente, la reconstrucción de los hechos y la valoración de las pruebas se basan en la experiencia del jurista, que se ve afectada por sesgos y prejuicios que esta tecnología puede ayudar a mitigar. Además, la IA podría **ayudar a evaluar** peritajes técnicos, a analizar pruebas documentales de forma rápida y a sugerir hipótesis reconstructivas.

Todo esto, concluyó Nieva, debe ocurrir siempre bajo el **control** de los juristas y del Estado, no de empresas tecnológicas de terceros.

ALFREDO RODRÍGUEZ DEL BLANCO, abogado penalista y miembro del CEISIA, contribuye a la discusión abordando los problemas del uso de la IA en el proceso penal en relación con las garantías procesales fundamentales.

Las características de la IA, subraya Rodríguez, la hacen particularmente útil en las actividades de **investigación**. Sin embargo, para evitar una peligrosa compresión de las garantías procesales, es necesario marcar sus límites, para evitar una violación inaceptable.

Para ello, es imprescindible, en opinión de Rodríguez, analizar la jurisprudencia nacional, constitucional e internacional sobre los principales derechos fundamentales implicados en la fase de investigación: igualdad, libertad personal y privacidad.

La **igualdad** se ve perjudicada por el posible uso de resultados de IA viciados por **sesgos**.

La **libertad personal** por aquellos sistemas opacos empleados para la aplicación de una medida cautelar.

La **privacidad** por un uso indiscriminado y generalizado de los datos personales, así como por el empleo de la IA en las actividades de interceptación.

A esto se suman los riesgos relacionados con la introducción de pruebas falsas a través de la IA generativa, de difícil evaluación con los medios actualmente a disposición de la autoridad judicial.

Esto sugeriría que, en ausencia de un adecuado control humano, la justicia podría deshumanizarse.

Para evitar este riesgo, es imprescindible para Rodríguez una **adecuada alfabetización digital de los abogados**, con vistas a la mejor defensa posible de las partes.

AITOR CUBO, director general de Transformación Digital del Ministerio de Justicia español, cerró la mesa redonda ilustrando el papel del Ministerio de Justicia español en el proceso de transición digital.

La digitalización de la justicia es vista por Cubo como una forma de hacer un servicio público más eficiente y preciso para el ciudadano.

En este sentido, el Ministerio de Justicia español está trabajando tanto en proyectos de automatización de la justicia como de empleo de la IA.

Con referencia específica a la IA, Cubo propone una distinción basada en el diferente impacto en las decisiones jurisdiccionales: sin impacto, con impacto indirecto, con impacto directo. Esta última hipótesis, naturalmente, requiere las mayores precauciones.

Para lograr la transición, la orientación de los datos al sector de la justicia es de crucial importancia. Los datos deben generarse ya en el momento inicial de los procesos.

Cubo concluyó mencionando algunas de las posibles aplicaciones en el sector.

En cuanto a los **procesos de automatización**, Cubo recordó: la **cancelación de oficio de los antecedentes penales al vencimiento de los plazos legales**; la automatización de los **procesos monitorios**; la **gestión de fianzas y pagos**.

En cuanto a las **aplicaciones de IA**, Cubo mencionó: el análisis y la clasificación documental; la simplificación de textos legales; la transcripción automática de grabaciones; la asistencia judicial en el análisis de expedientes de gran tamaño y en la identificación de relaciones entre las partes.

CONFERENZA DI CHIUSURA

Relatore: Lorenzo **Cotino** Hueso

Interventi: Miguel Ángel Presno Linera, Ignacio Villaverde (Rector de la UNIOVI).

LORENZO COTINO HUESO concluyó la semana de trabajos aportando su doble perspectiva como profesor de derecho constitucional y director de la “*Agencia Española de Protección de Datos*” (AEPD).

Para Cotino es fundamental la **colaboración entre la academia y las instituciones** públicas en el sector de la IA.

Esta última debe ser empleada para aumentar la productividad de la administración pública, así como un instrumento para lograr la efectividad de los derechos.

Además, Cotino abogó por el reconocimiento de un "**derecho constitucional a la innovación tecnológica**" al que correspondería un deber del Estado de eliminar los obstáculos y de iniciar iniciativas para un desarrollo similar.

Como institución pública, la AEPD está trabajando en primera línea para identificar los posibles "fines de uso" de la IA en todas las unidades, evaluando, al mismo tiempo, los riesgos. El objetivo de la AEPD es aumentar su productividad, pero sobre todo convertirse en un ejemplo virtuoso para el resto del sector público, en el que la IA se implemente respetando la privacidad y la seguridad.

A este respecto, Cotino propuso la creación de **registros públicos de algoritmos**: herramientas esenciales para la transparencia y la supervisión de la IA por parte de la autoridad pública.

Desde un punto de vista dogmático, Cotino criticó la escasa atención de los constitucionalistas hacia las **metodologías consolidadas de análisis de riesgos** y estudios de impacto sobre la protección de datos, que permitirían una mayor precisión en las evaluaciones de impacto sobre los derechos fundamentales.

Además, Cotino destaca la necesidad de coordinación interpretativa entre el GDPR y el AI Act, los cuales presentan algunos ámbitos de superposición, el primero de ellos el sancionador. Esto se traduce en una ampliación de las competencias de la AEPD extendida a aquellos aspectos de la IA estrictamente ligados al tratamiento de los datos personales.

En términos más generales, concluyó Cotino, parece necesario dotar el uso de la IA en sectores sensibles particulares (por ejemplo, justicia) de una **adecuada cobertura jurídica**, que establezca los estándares mínimos de aplicación.