

Horizontes éticos en la era algorítmica

Santiago Ullauri
Betancourt (editor)

Horizontes éticos en la era algorítmica

**Memoria de la I Jornada Internacional sobre
el Fenómeno de la Inteligencia Artificial: La
revolución de la IA y sus fronteras éticas**

Santiago Ullauri - Betancourt (editor)

Título:

Horizontes éticos en la era algorítmica. Memoria de la I Jornada Internacional sobre el Fenómeno de la Inteligencia Artificial: “La revolución de la IA y sus fronteras éticas”.

Editado por: Santiago Ullauri Betancourt

Primera edición

ISBN: 978-9942-752-24-6

DOI: <https://doi.org/10.31207/uhediciones8>

Publicado digitalmente en septiembre de 2025

© UHEdiciones • Universidad de Los Hemisferios

Paseo de la Universidad N° 300 y Juan Díaz, Ñaquito Alto

Quito - Ecuador.

Telf. (02) 401-4100

www.uhemisferios.edu.ec

Diagramación: Diego Ortiz Jaramillo.
Revisión editorial: Daniel López Garzón.
Revisión ortográfica y Corrección de estilo: Isaac Carbajal.

Esta obra fue publicada con el aval de lectores especializados y arbitrada según las normas de publicación del Centro de Publicaciones - UHEdiciones de la Universidad Hemisferios.

Este documento ha sido desarrollado y publicado por el Instituto para el Desarrollo de la Cultura y Sociedad (IDECS) de la Universidad Hemisferios con la colaboración de la Fundación Centro Interdisciplinario en Ética, Política y la Fundación Info Kratos. Las opiniones expresadas en este documento son criterios de los autores y no necesariamente reflejan las posiciones institucionales de la Universidad Hemisferios. Se permite copiar, distribuir y comunicar públicamente solamente copias inalteradas. Esta obra no puede ser utilizada con finalidades comerciales, a menos que se obtenga permiso por escrito de la editorial.

Cita sugerida (APA 7):

Ullauri - Betancourt, S. A. (Ed.). (2025). *Horizontes éticos en la era algorítmica. Memoria de la I Jornada Internacional sobre el Fenómeno de la Inteligencia Artificial: “La revolución de la IA y sus fronteras éticas”*. Universidad Hemisferios. DOI:



COMITÉ ORGANIZADOR

Universidad Hemisferios, Ecuador

Dra. (c) María Gabriela Rodríguez.

Dr.(c). Diego Ortiz Jaramillo.

Universidad Católica del Norte, Chile

Dr. Boris Briones

Fundación Centro Interdisciplinario en Ética, Política y Economía, Chile

Dr. (c) Isaac Carbajal

Mtro. Nicolás Fuentes Valdebenito

Mtra. Adriana del Valle

Lic. Salvador Quirilao

Lic. Bethlehem Rodríguez

Fundación Info Kratos

Lic. Juan Esteban Echeverry

INSTITUCIONES COLABORADORAS

International Association for Political Science Students (IAPSS), Canadá.

Instituto de Ciencias Religiosas y Filosofía de la Universidad Católica del Norte, Chile.

Filosófica. Fundación de Estudios Filosóficos, Políticos y Culturales, Ecuador.

Instituto de Estudios Políticos Andinos, Perú.

Escuela de Humanidades, Universidad Gabriela Mistral, Chile.

Cátedra Unesco de Ética y Sociedad en la Educación Superior, Universidad Técnica Particular de Loja, Ecuador.

Centro Latinoamericano de Estudios Superiores, México.

COMITÉ CIENTÍFICO

Dra. Yadyra Piñera Concepción - Universidad Bolivariana del Ecuador

Dr. Pablo Ruiz Osuna - Universitat Rovira i Virgili

Dr. Luis Quintero - Universidad Privada Dr. Rafael Belloso Chacín

Dr (c). Santiago Pérez Samaniego - Universidad Técnica Particular de Loja

Dr. (c) Javier Fattah Jeldres - Universidad Católica de la Santísima Concepción

Dra. (c) Pía Bustamante Barahona - Universidad San Sebastián

Dr. (c) Félix Andueza - Universidad Católica Andrés Bello

Dr. (c) Felipe Álvarez Osorio - Universidad de Chile

Dr. (c) Guillermo Colella - Universidad Nacional de Quilmes

Dr. (c). María Luisa Azanza - Universidad Hemisferios

Mg. María Fernanda Idrobo - Universidad Técnica Particular de Loja

Mg. Christopher Acuña - Universidad Andrés Bello

Mg. Camilo Sepúlveda - Universidad de Chile

Prólogo

Esta memoria reúne los resúmenes presentados en la I Jornada Internacional sobre el Fenómeno de la Inteligencia Artificial: “La revolución de la IA y sus fronteras éticas”, un evento académico organizado de manera conjunta por el Centro Interdisciplinario en Ética, Política y Economía (CIEPE), la Fundación Info Kratos y el Instituto para el Desarrollo de la Cultura y Sociedad (IDECS) de la Universidad Hemisferios. La jornada se celebró los días 26, 27 y 28 de mayo de 2025 en modalidad virtual y congregó a más de cuarenta ponentes de distintos países y disciplinas.

Concebida como un espacio de encuentro y deliberación, la Jornada abrió un diálogo plural y riguroso sobre los desafíos que plantea la inteligencia artificial a nuestras instituciones, a nuestras formas de vida y a nuestra propia comprensión de lo humano. Investigadores, estudiantes y profesionales de ámbitos tan diversos como el derecho, la filosofía, la medicina, la psicología, la educación, la ingeniería y las humanidades compartieron reflexiones que dan cuenta de la riqueza, complejidad y urgencia del tema.

El programa estuvo compuesto por once mesas temáticas, tres conferencias magistrales y un panel de expertos. Las conferencias fueron impartidas por el Dr. Antonio Luis Terrones, la Dra. Gabriela Arriagada y el Dr. Miguel Ángel Fernández, cuyas exposiciones ofrecieron marcos de análisis de gran profundidad sobre los fundamentos éticos, epistemológicos y sociales de la inteligencia artificial. El panel, integrado por el Dr. Hugo Tapia Silva, el Dr. Marcelo Correa y el Dr. Cristhian Almonacid, generó un debate interdisciplinario que permitió articular visiones complementarias sobre los retos normativos, antropológicos y políticos de estas tecnologías.

Las contribuciones recogidas en esta memoria no constituyen artículos académicos acabados, sino resúmenes de investigación que ofrecen un panorama amplio de los debates conceptuales, metodológicos y éticos que atraviesan hoy la discusión sobre la inteligencia artificial. Gracias a esta cartografía, el lector podrá advertir cómo la IA transforma las prácticas jurídicas, reconfigura la educación, modifica el cuidado médico, introduce

sesgos en la política y la economía, e interpela nuestras categorías filosóficas más elementales, desde la conciencia y la emoción hasta la libertad y la responsabilidad.

El IDECS desempeñó un rol central en la concepción de este encuentro, reafirmando su misión de vincular la investigación en cultura y sociedad con los desafíos emergentes de la tecnología. Para el Instituto, esta Jornada marca un hito fundacional en una línea de trabajo que continuará creciendo con investigaciones, publicaciones y nuevos espacios de encuentro internacional. Más que un evento aislado, estas memorias son la apertura de un proceso de reflexión interdisciplinaria sostenida, orientada a consolidar una gobernanza responsable y humanista de la inteligencia artificial.

Agradecemos profundamente a los ponentes, moderadores y colaboradores que hicieron posible esta primera edición, así como al conjunto de instituciones académicas y culturales que brindaron su apoyo. Gracias a ellas, fue posible dar vida a un espacio de discusión que trasciende lo técnico y se adentra en lo ético, lo político y lo cultural.

Estas memorias no buscan clausurar un debate, invitan a continuarlo. Porque más allá de la fascinación tecnológica, lo que está en juego con la inteligencia artificial es la orientación moral y política de nuestras sociedades. Pensar la IA significa preguntarnos no solo qué puede hacer la tecnología, sino qué debemos hacer nosotros como comunidad humana.

Santiago Ullauri-Betancourt

Instituto para el Desarrollo de la Cultura y Sociedad (IDECS)

ÍNDICE

Mesa 1. IA y Educación

Gestión ética y seguridad en el uso de IA por estudiantes universitarios

Miguel Ángel Cordero Monzónp. 11

Mesa 2. IA, Ética y Economía

Inteligencia Artificial y Nuevas Formas de Comunidad: Un Análisis Ético a través del Pensamiento de Byung-Chul Han

Boris Osvaldo Saavedra Pérezp. 12

Inteligencia artificial y el riesgo de la manipulación económica

Gabriel Donoso Umañap. 13

La IA - Inteligencia Artificial y la privacidad: Análisis de la Recolección de Datos Personales

Freddy Lenin Villarreal Satamap. 14

Más allá de los principios abstractos: fundamentar la IA responsable en la ética del discurso

Dr. Heber Leal & Dr. Joaquim Giannottip. 15

Mesa 3. IA, Política y Educación

Gobernanza de la IA: Cuando la ética no puede reducirse a algoritmos, casos ecuatorianos

Valeria Argüello Castrop. 17

Mesa 4. IA y Derecho

Nuevos Paradigmas Penales en la Era Digital: La Instrumentalización de la IA en el Caso Saint George (2024)

Felipe Alonso Jara Sanhuezap. 19

Diez años de la Reforma de nulidad matrimonial canónica y el potencial de la Inteligencia Artificial

Pbro. Dr. Paolo Rossano Apontep. 20

Hacia la transparencia algorítmica como un derecho laboral. Algunas posibilidades e inconvenientes

Francisco A. Ruay Sáezp. 21

IA y Derecho: Hacia la Recuperación del Ars Iuris en la Era Tecnológica

María Luisa Azanzap. 22

Transparencia algorítmica

Saulo Ricardo Jaramillo Romerop. 23

Mesa 5. IA y Ciencias Sociales

Interpretabilidad algorítmica en el proceso de identificación humana: Un enfoque de machine learning explicable para la estimación de la edad en adultos en antropología forense

Noemí Aedo-Noa & Stefano De Lucap. 25

Producción de sentido y la crítica sociotécnica en la era de la IA: Una aproximación sociológica contemporánea

Andrea Vanessa Cáceres Silvap. 26

Los riesgos de la IA en Tecnologías Emergentes y en Tecnologías No Reguladas

Horacio Antonio Correa Arechavalap. 27

Neuroética y autonomía en la toma de decisiones políticas en la era de la inteligencia artificial

Martín Bórquezp. 28

Mesa 6. IA y Sociedad

Estatuto ético y antropológico de las inteligencias artificiales

Eduardo Gutiérrez Gutiérrezp. 30

La doble brecha en América Latina: Desigualdad económica y brecha digital urbano-rural en la era de la transformación digital

Christopher Acuña Martínezp. 31

Análisis de las dimensiones del aprendizaje y su influencia en el uso ético de la inteligencia artificial generativa en estudiantes universitarios ecuatorianos

Jorge Luis Buele Leónp. 32

Mesa 7. IA y Filosofía

El replicante como genio maligno moderno. Lecciones para el desarrollo de la inteligencia artificial

Ariel Abelardo Sánchez Jarap. 34

El problema de la manifestación del fenómeno humano en la experiencia de la IA: un análisis desde la fenomenología de la donación

Ignacio Alarcónp. 35

¿Puedes sentir sin existir? Análisis ontológico de las emociones sintéticas

Isaac Carbajalp. 36

Mesa 8. IA, Religión y Bioética

Robotofobia vs. robofilia: un análisis comparativo sobre la consideración del autómatas inteligente en las cosmovisiones occidental y asiática

Pablo Ruiz Osunap. 37

Religión y Nuevas Tecnologías: Tensiones y complementariedades

Alejandro Cerda Sanhuezap. 38

Diagnósticos sin rostro: inteligencia artificial y medicina despersonalizada

Pía Bustamante-Barahona & Fernando Chuecas Saldíasp. 38

Mesa 9. IA, Humanidades y Derecho

Colaboración Hombre-Máquina: Nuevos paradigmas en la Inteligencia Artificial y su impacto humanístico en el trabajo cognitivo

María Arantxa Serantesp. 40

El uso ético de la Inteligencia Artificial en la formación jurídica chilena: riesgos de pérdida de habilidades y la importancia insustituible de la cátedra de Metodología de la Investigación

Dr. Gabriel Álvarez Undurraga & Luca Martino Acevedop. 41

Cuidado y robótica asistencial

Daniel Andrés Fuentealba Cid & Maritza Belén Ramos Faríasp. 41

Mesa 10. IA, Psicología y Humanidades

La banalidad de la privacidad en la era de la Inteligencia Artificial. Acepto los términos y condiciones... ¿comprendes lo que significa para tu privacidad?

Mtra. Viviana Del Moral Munguíap. 43

Autores IA y tendencias forzadas: un análisis del caso Jianxei Wu

Camilo Schenonep. 44

El consentimiento y los usos valiosos y disvaliosos de la IA

Lucía Martínez Limap. 45

Mesa 11. IA y Educación

Más allá de la personalización: límites éticos y pedagógicos del aprendizaje adaptativo mediante Inteligencia Artificial

Paulina Andrea Gallardo Gómezp. 47

Percepciones de estudiantes de IES sobre IA generativa y la incorporación de ODS 4

Carola Ubilla Brionesp. 48

Didáctica de la filosofía de la IA: Algoritmo de decisión y antropomorfismos en el debate sobre la inteligencia artificial

Jesús Queglas Caruzp. 49

Mesa 1 “IA y Educación”

Gestión ética y seguridad en el uso de IA por estudiantes universitarios

Miguel Ángel Cordero Monzón

Universidad Panamericana / Universidad San Pablo

miguelcordero777@gmail.com

El presente artículo tuvo como objetivo analizar la gestión ética y la seguridad en el uso de la inteligencia artificial por estudiantes universitarios, identificando los desafíos y proponiendo estrategias para promover un uso responsable de la tecnología en entornos académicos. Para ello, se adoptó un enfoque mixto que combinó revisión de literatura actual, y encuestas para evaluar la percepción y las prácticas en torno al uso de la IA. Los hallazgos revelaron que, aunque la adopción de herramientas digitales en la educación mejora la eficiencia y facilita el acceso a la información, existen preocupaciones muy significativas en cuanto a la protección de datos, la transparencia en los algoritmos y el riesgo de sesgos en la toma de decisiones. Asimismo, se identificaron prácticas de ética dudosa y se propusieron recomendaciones que incluyen la incorporación de módulos de formación ética en el currículo, la creación de comités internos para la supervisión de proyectos de IA y el desarrollo de protocolos de seguridad robustos. Estos resultados sugieren que, mediante una adecuada integración de estrategias de capacitación y normativas internas, es posible mitigar los riesgos asociados a la tecnología y fomentar un entorno educativo seguro y ético.

Mesa 2 “IA y ética/ economía”

Inteligencia Artificial y Nuevas Formas de Comunidad: Un Análisis Ético a través del Pensamiento de Byung-Chul Han

Boris Osvaldo Saavedra Pérez

Universidad San Sebastián

bsaavedrap@docente.uss.cl

El filósofo surcoreano-alemán Byung-Chul Han, se ha caracterizado por realizar una profunda y certera crítica de los efectos generados por las tecnologías digitales en el ser humano. Aunque Han, específicamente, no aborda de manera exhaustiva el fenómeno de la Inteligencia Artificial (IA), eso no quita su pronunciamiento sobre el tema, de hecho, gran parte de su reflexión gira en torno a las tecnologías de la información, la digitalización y la automatización.

Así, Han subraya los efectos deshumanizantes de la sociedad digitalizada, caracterizada por la lógica del rendimiento y la eficiencia. Por lo que, frente a los retos que plantea la IA, el filósofo surcoreano-alemán aboga por una ética que recupere la humanitas del ser humano. En este sentido, Byung-Chul Han defiende la importancia de la reflexión, la contemplación y la creación de espacios para la interacción genuina, acentuando una distancia de las dinámicas de consumo y optimización que dominan la vida contemporánea.

El objetivo de esta ponencia es explorar cómo la IA transforma las formas de comunidad en la modernidad tardía, considerando sus principales implicaciones éticas, con esto se pretende examinar las tensiones y contradicciones generadas por el desarrollo tecnológico, pues desde la perspectiva de Han, éstos afectan irremediablemente la intimidad y la construcción de la comunidad. Aunque, la IA tiene el potencial de generar avances significativos en términos de eficiencia y bienestar. No obstante, Han destaca que su implementación debe estar guiada por principios éticos que respeten la dignidad humana, promuevan la equidad y prevengan la manipulación de los individuos. En este aspecto, la presente investigación es relevante para advertir que la IA no debe considerarse como una herramienta para sustituir a los seres humanos, sino como un medio para facilitar la cooperación y el desarrollo. De este modo, la nueva comunidad digital debe concebirse como un espacio donde la tecnología potencie las capacidades humanas y no las degrade.

En consecuencia, la comunidad tiene que ser un lugar donde se fomente la solidaridad, el cuidado mutuo y el sentido común, aspectos que la lógica de la IA tiende a desbordar al privilegiar la eficiencia y el aislamiento.

Inteligencia artificial y el riesgo de la manipulación económica

Gabriel Donoso Umaña

Universidad de Santiago de Chile

gabriel.donosu.u@usach.cl

Esta exposición analiza críticamente el impacto de la inteligencia artificial (IA) en el comportamiento económico actual, poniendo en tensión sus promesas emancipadoras con los riesgos de manipulación digital que comprometen la autonomía y la libertad cognitiva de los agentes. A partir de su consolidación como tecnología de propósito general (GPT) —junto a la máquina de vapor, la electricidad e internet)—, la IA ha sido incorporada transversalmente a los sistemas económicos, transformando tanto la estructura de los mercados como la conducta de los agentes que operan en ellos.

La investigación parte desde una concepción ampliada de los fenómenos económicos como procesos coevolutivos entre conciencia, tecnología y entornos institucionales y considera la IA como respuesta a las limitaciones de la racionalidad cotidiana. Esta racionalidad, por su carácter sesgado e informacionalmente limitada, no permite una optimización plena en contextos económicos. En cambio, las IA basadas en redes neuronales y aprendizaje automático logran procesar volúmenes masivos de datos, facilitando decisiones informadas a través de técnicas como el análisis topológico de datos y la teoría de constructores (CT-ML), que permiten predecir tendencias y optimizar elecciones.

No obstante, la creciente influencia de la IA en la conducta económica plantea interrogantes filosóficas profundas sobre la autonomía personal. La autora analiza cómo sistemas como Siri o los algoritmos de recomendación de Netflix afectan la formación de preferencias, y se pregunta si puede hablarse genuinamente de decisiones libres en contextos altamente asistidos. Para abordar este dilema, se recurre a la noción de autonomía personal como la capacidad de elegir libremente y asumir responsabilidad por los propios actos, y se introduce el concepto emergente de libertad cognitiva, es decir, el derecho a controlar los propios procesos mentales frente a interferencias externas.

La exposición profundiza en los mecanismos de manipulación digital ejercida por IA, caracterizados por intencionalidad, asimetría de resultados, opacidad y violación de la autonomía. Se documentan casos como el trading algorítmico, la publicidad dirigida y la creación de burbujas informativas, y se advierte sobre el potencial liberticida del cruce entre IA y neurotecnologías.

Ante estos desafíos, se proponen cuatro líneas de acción para proteger la libertad cognitiva y la autonomía personal: (1) políticas públicas de transparencia y sanción, (2) alfabetización financiera, (3) autorregulación corporativa responsable, y (4) protección legal de los neuroderechos. Se sostiene que enfrentar el riesgo de manipulación exige una colaboración multisectorial entre gobiernos, sociedad civil y empresas tecnológicas, con miras a preservar los valores emancipatorios en una economía digital global.

La IA - Inteligencia Artificial y la privacidad: Análisis de la Recolección de Datos Personales

Freddy Lenin Villarreal Satama

Universidad Hemisferios/IMF

leninv@uhemisferios.edu.ec

El presente trabajo de investigación sobre la Inteligencia Artificial (IA) tiene como objetivo principal, analizar el uso de los datos con la privacidad individual, derivado de la recolección de datos personales como fuente de mecanismo de obtención de bases de datos mediante los sistemas de IA y su manejo a gran escala de volúmenes de información, extraída de los usuarios muchas veces sin consentimiento pleno o comprensión total de su uso. Por esta razón se vuelve crucial examinar los riesgos éticos y legales que implica la explotación de estos datos.

La metodología utilizada en este estudio es de tipo cualitativa, a través de una revisión bibliográfica y documental. La consulta se basó principalmente en la recolección de artículos científicos de bases de datos indexadas, informes de organismos internacionales (como la ONU y la Unión Europea), documentos legales (como el Reglamento General de Protección de Datos - RGPD) y publicaciones referentes a la IA, privacidad y ética lo que permitió identificar patrones, tendencias y sostener posturas críticas sobre este tema coyuntural. Los principales resultados muestran que la IA se alimenta en gran medida de la extracción masiva de datos personales, como historiales de navegación en internet, redes sociales, geolocalización, actividad de la banca virtual, datos de organismos estatales de servicio, preferencias de consumo y datos biométricos comerciales, entre otros. Esta información es esencial para armar la arquitectura tecnológica que a la postre generan la creación y entrenamiento de algoritmos de aprendizaje para generar publicidad dirigida. Sin embargo, el análisis evidencia que muchas prácticas de recolección no respetan principios básicos de privacidad, como el consentimiento informado, la transparencia y la minimización del uso de datos en el que se identifican riesgos asociados, como la discriminación algorítmica, la vigilancia masiva, la pérdida de autonomía digital y la dificultad de eliminar los datos recopilados. A pesar de que existen marcos regulatorios como el RGPD en Europa o la Ley de Privacidad del Consumidor de California (CCPA), se observa que la aplicación efectiva de estas normas es aún limitada frente al avance técnico de los sistemas de IA.

Entre las conclusiones más destacadas se afirma que, si bien la IA ofrece beneficios en múltiples áreas, su desarrollo debe estar acompañado de políticas claras y estrictas de protección de datos por lo que es importante fortalecer una legislación pertinente en términos de privacidad, fomentando la educación digital y promover el desarrollo de sistemas de IA éticos y transparentes. Asimismo, es necesario que las empresas tecnológicas asuman mayor responsabilidad social en la manera en que utilizan los datos personales.

En resumen, el uso de inteligencia artificial plantea desafíos importantes para la privacidad, y es esencial equilibrar la innovación tecnológica con los derechos fundamentales de las personas, garantizando una IA centrada en el ser humano.

Más allá de los principios abstractos: fundamentar la IA responsable en la ética del discurso

Dr. Heber Leal & Dr. Joaquim Giannotti

Universidad San Sebastián & Universidad Mayor
heber.leal@uss.cl / joaquim.giannotti@umayor.cl

El término "IA Responsable" describe diversos enfoques que buscan aplicar directrices éticas y equitativas a los sistemas de IA en diversos contextos. Muchos académicos argumentan que su inherente abstracción los hace inadecuados para guiar acciones concretas en las situaciones matizadas y, a menudo, moralmente ambiguas que enfrentan los profesionales de la IA. Esta abstracción puede dar lugar a interpretaciones contradictorias y a la implementación de prácticas que, a pesar de adherirse a la letra de un principio, pueden resultar moralmente controvertidas en sus resultados. Por ejemplo, el principio de "equidad" puede operacionalizarse de diversas maneras estadísticas, cada una con diferentes implicaciones para los distintos subgrupos de una población, lo que puede llevar a disyuntivas éticamente difíciles de abordar. Estos desafíos plantean dudas legítimas sobre la eficacia de principios amplios y de propósito general para abordar eficazmente los dilemas éticos prácticos que los diseñadores, desarrolladores e implementadores de IA enfrentan a diario. La brecha entre los ideales aspiracionales de la IA Responsable y las realidades de su implementación sigue siendo un obstáculo importante. Ofreciendo una evaluación crítica de los marcos existentes, defendemos la necesidad de principios generales para una IA responsable, basados en la ética fundamental. Nuestro enfoque integra la ética del discurso de Jürgen Habermas, que enfatiza la comunicación racional entre las partes interesadas, y el Enfoque de Capacidades de Martha Nussbaum, que se centra en el impacto de la IA en el desarrollo humano. Correctamente entendidos y comunicados en estos términos, los principios pueden guiar el diseño ético de la IA y las acciones moralmente responsables.

Debidamente entendidos y comunicados eficazmente, estos principios pueden guiar las acciones moralmente responsables e informar el diseño ético de la IA, acortando

la distancia entre los principios abstractos y la aplicación práctica. A partir de la literatura, analizamos los marcos éticos existentes y proponemos recomendaciones, enfatizando la necesidad de un diálogo multidisciplinario sostenido con profesionales tecnológicos y partes interesadas para acortar la distancia entre los principios y la práctica.

Concluimos sugiriendo un marco basado en el discurso para acortar la distancia entre los principios éticos y las estrategias prácticas. Este intercambio continuo es crucial para fundamentar la ética en el desarrollo de la IA en el mundo real. La visión resultante es una ética de la IA con base filosófica que permite la evaluación práctica y la investigación empírica.

Mesa 3 “IA y política/ educación”

Gobernanza de la IA: Cuando la ética no puede reducirse a algoritmos, casos ecuatorianos

Valeria Argüello Castro

vale.arguello.castro@gmail.com

Esta investigación explora la compleja transición de los principios teóricos a la aplicación práctica de la IA en el Ecuador. Este estudio aborda tres dilemas éticos fundamentales no resueltos en el campo de la IA: la presencia de sesgos algorítmicos derivados de datos no representativos; la dificultad inherente en la definición de categorías protegidas (como género o etnia) sin perpetuar estereotipos sociales; y la evidente ausencia de mecanismos institucionales efectivos para traducir los principios éticos en decisiones técnicas concretas, un desafío aún mayor en contextos caracterizados por asimetrías de poder.

Este estudio realiza una deconstrucción normativa de la IA en Ecuador, revelando que, a pesar del notable impulso gubernamental a diversos proyectos tecnológicos, la nación carece de una gobernanza ética con una visión sistémica. Esta carencia refleja los dilemas éticos que trascienden fronteras y que persisten a nivel internacional. A través de un análisis que abarca ejes filosóficos, políticos y aplicados, la investigación cuestiona de manera contundente la falacia de la neutralidad algorítmica. En su lugar, propone y explora alternativas situadas que promuevan la integración de políticas públicas inclusivas para el uso y desarrollo de los sistemas de IA en el país.

Un hallazgo central de esta investigación no es la inexistencia de normas, sino la dispersión de las menciones éticas y las responsabilidades a lo largo de 13 normas primarias y secundarias. Estas fueron ponderadas y analizadas en seis categorías clave: social y educación, empleo, reducción de la brecha digital, confianza y seguridad, entornos seguros, y entornos éticos y responsables. Las conclusiones del estudio enfatizan que, a pesar de los importantes avances normativos, Ecuador aún carece de una gobernanza ética con una visión sistémica, lo que lo deja expuesto a dilemas que ningún algoritmo, por más sofisticado que sea, puede resolver por sí mismo. Aunque Ecuador manifiesta voluntad, posee un marco normativo incipiente y un ecosistema tecnológico emergente, la pieza ausente es una arquitectura de gobernanza participativa que articule todos estos elementos. Se subraya que la ética

en la IA no puede reducirse a meros algoritmos de mitigación de sesgos o a artículos aislados en la legislación; es, fundamentalmente, una decisión de diseño sistémico que debe ser deliberada, involucrar a múltiples actores y estar imbuida de una perspectiva estratégica. Finalmente, se concluye que no es posible mitigar los sesgos algorítmicos con otro algoritmo, resaltando la necesidad imperante de un enfoque holístico que integre las ciencias sociales para abordar la justicia y la equidad en el desarrollo y uso de la IA.

Mesa 4 “IA y derecho”

Nuevos Paradigmas Penales en la Era Digital: La Instrumentalización de la IA en el Caso Saint George (2024)

Felipe Alonso Jara Sanhueza

Universidad Católica de la Santísima Concepción (UCSC)

fjaras@derecho.ucsc.cl

“El día que la inteligencia artificial se desarrolle por completo podría significar el fin de la raza humana. Funcionará por sí sola y se rediseñará cada vez más rápido. Los seres humanos, limitados por la lenta evolución biológica, no podrán competir con ella y serán superados”. – Stephen Hawking.

Comenzando con esta distópica pero cada día más verídica frase del gran científico Hawking, esta propuesta se inserta en el área del derecho penal y la tecnología, explorando los desafíos que plantea el uso de inteligencia artificial (IA) generativa para la creación y difusión de contenido sexual no consentido. La pregunta central que guía esta investigación es: ¿puede el derecho penal atribuir responsabilidad jurídica cuando se instrumentaliza la IA para generar deepfakes sexuales sin consentimiento? Y si es así, ¿qué categorías dogmáticas deben actualizarse para enfrentar eficazmente este fenómeno?

La ponencia parte desde la exposición del caso ocurrido en el año 2024 en el Colegio Saint George de Chile, donde un grupo de estudiantes utilizó herramientas de IA generativa para crear imágenes sexualizadas falsas de sus compañeras, las cuales fueron posteriormente difundidas por redes sociales. Este trabajo propone una revisión crítica de los fundamentos de la autoría, participación, imputación subjetiva y tipicidad penal en contextos de criminalidad digital. Sostenemos que el uso malicioso de IA en este tipo de casos puede ser reconducido a la figura de “autoría mediata” del artículo 15, número 2 del Código Penal chileno, sostenida por el concepto de “dominio de la voluntad”, que conserva el ser humano sobre la máquina llamada IA. Por ende, este tipo de hechos no excluye la posibilidad de imputar responsabilidad penal a quienes emplean la IA como medio o instrumento, manteniendo al agente humano como núcleo esencial de la imputación.

La hipótesis es que las categorías tradicionales del derecho penal pueden mantenerse, pero deben ser interpretadas de forma evolutiva y tecnológica,

reconociendo que el agente humano sigue siendo quien toma decisiones —con dolo o culpa— sobre la creación de contenidos ilícitos. Asimismo, se argumenta que el consentimiento, la afectación a la intimidad sexual y la reputación deben ocupar un rol central en el análisis típico.

Se utilizará una metodología dogmático-empírica, combinando el análisis del caso chileno con revisión comparada de ordenamientos que ya han legislado sobre la materia, como Reino Unido, Australia, Corea del Sur y algunos estados de EE.UU., donde se han creado tipos penales específicos contra los deepfakes sexuales. Se aspira a que nuestra legislación también tipifique una figura penal en este sentido. Además, se incorporan perspectivas de la criminología tecnológica y la sociología del castigo, evaluando si el derecho penal está respondiendo eficazmente a este nuevo tipo de violencia digital.

Este estudio busca contribuir a la construcción de un derecho penal sensible a los efectos sociales de la tecnología, que no abdique de sus principios, pero que sea capaz de ofrecer protección a las víctimas frente a nuevas formas de agresión o acoso digital, cuyas acciones, aunque mediadas por IA, siguen siendo profundamente humanas en su origen y consecuencias.

Diez años de la Reforma de nulidad matrimonial canónica y el potencial de la Inteligencia Artificial

Pbro. Dr. Paolo Rossano Aponte

Tribunal Eclesiástico Interdiocesano Francófono para Bélgica
rossanopaolo@gmail.com

La reforma del proceso de nulidad matrimonial promulgada por el Papa Francisco en 2015, mediante los motu proprio *Mitis iudex Dominus Iesus* y *Mitis et misericors Iesus*, representa un punto de inflexión histórico para la Iglesia católica, ya que simplifica y hace más accesible el camino jurídico para la verificación de la invalidez del matrimonio. Concebida para reducir tiempos y distancias entre los tribunales eclesiales y los fieles, se sitúa en continuidad con la tradición canónica, pero revisa profundamente las normas y procedimientos que permanecieron inmutables desde el siglo XVIII. Paralelamente, el surgimiento de la Inteligencia Artificial (IA), popularizado en particular por el lanzamiento de ChatGPT en 2022, ofrece nuevas perspectivas para acelerar y hacer más eficiente la gestión de las causas de nulidad en los tribunales eclesiales. Gracias a herramientas como bots conversacionales y algoritmos de procesamiento del lenguaje natural, es posible organizar y analizar rápidamente la documentación, extraer datos relevantes y generar borradores de actos procesales. Esto permite reducir la carga de trabajo rutinario y concentrar la atención de los jueces en los aspectos pastorales y subjetivos de cada caso. Sin embargo, es indispensable reconocer los límites éticos y prácticos de la IA: la inteligencia humana, enriquecida por la experiencia de vida y por la dimensión

corporal y relacional, no puede ser sustituida por un sistema basado en la lógica computacional. Las decisiones finales siguen siendo de competencia exclusiva de los jueces eclesíásticos, quienes deben integrar el rigor jurídico con una sensibilidad pastoral. Además, el uso de la IA exige que se protejan la confidencialidad y los derechos fundamentales de las personas, evitando cualquier abuso o discriminación en el acceso a las tecnologías avanzadas. En perspectiva, la integración de cursos de Legaltech e IA en las facultades de Derecho Canónico podría proporcionar a los futuros canonistas competencias valiosas para afrontar los desafíos de un mundo en rápida evolución. La Iglesia, fiel a su misión de proximidad a los fieles, puede así aprovechar las herramientas digitales para mejorar la eficiencia de los tribunales, sin perder de vista el objetivo de dar testimonio de la verdad y de ofrecer misericordia y justicia.

Hacia la transparencia algorítmica como un derecho laboral. Algunas posibilidades e inconvenientes

Francisco A. Ruay Sáez
fruaysaez@gmail.com

La inclusión de sistemas que operan con inteligencia artificial en el mundo del trabajo incide existencial y jurídicamente en la vida de los trabajadores y empresarios. En el caso de los trabajadores particularmente se presenta un nuevo riesgo existencial: la alienación algorítmica. En relación a sus dimensiones regulatorias ésta incide: en el ejercicio de los poderes de control del empleador, en la determinación de los elementos del contrato de trabajo, en la constatación de antecedentes sancionatorios por ejercicio del poder disciplinario, en el uso de la IA como una herramienta de adaptabilidad ante circunstancias extraordinarias como la huelga, entre otros. En cada una de dichas dimensiones el trabajador en su posición esencialmente subordinada se ve expuesto al desconocimiento de parámetros que le permitan actuar con certeza jurídica. Así, por ejemplo, se ve expuesto a la opacidad en la determinación concreta de las condiciones de contratación en sus elementos esenciales. Por ejemplo, a propósito de la determinación algorítmica de la remuneración por viajes realizados con plataformas de servicios de transporte de personas o mercadería, o bien, en la determinación de medidas de control que se ejercen sobre su corporalidad y que operan utilizando sistemas que funcionan con IA.

En esta ponencia se analizó la viabilidad y los problemas que puede presentar configurar a la transparencia algorítmica como un derecho subjetivo laboral, su incidencia en cada una de las dimensiones potestativas del empleador (dirección, organizacional, disciplinario) y su especial relevancia en elementos mínimos que definen la existencia y desarrollo de la relación laboral.

Para alcanzar nuestro objetivo analizaremos en primer lugar las diversas dimensiones en que la IA ha intervenido en el proceso productivo, luego esclareceremos los

elementos esenciales del contrato de trabajo y la incidencia de la IA en su determinación actual, así como los poderes empresariales que son reconocidos al empleador y su interacción con la IA, para concluir analizando si es posible configurar la transparencia algorítmica como un derecho autónomo, o bien, como una derivación de derechos subjetivos clásicos en materia laboral, vinculados con el conocimiento de las condiciones laborales.

La investigación resulta relevante ante el estudio de la regulación normativa de la IA en nuestro país, la expansión de los sistemas que operan con IA en el mundo productivo, y su eventual regulación en el Código del Trabajo o en fuentes normativas especiales.

IA y Derecho: Hacia la Recuperación del Ars Iuris en la Era Tecnológica

María Luisa Azanza

Universidad Hemisferios

marialuisaa@uhemisferios.edu.ec

«Ius est ars boni et aequi.» Esta clásica definición de Celso, alude a la concepción del derecho no como un conjunto de reglas abstractas, sino como una habilidad práctica orientada a realizar la justicia en el caso concreto. Es decir, se trata de un saber práctico en contraposición con los saberes especulativos que reflejan la realidad pero no la hacen. Dice Hervada: todo saber práctico es un arte, también el derecho.

En la antigüedad clásica derecho, dikáion, ius se entendía como “lo justo” o “lo igual”. Esta concepción sufre una transformación en la modernidad y con la consolidación del Estado Moderno en el S. XIX el positivismo jurídico se hace reinante y encuentra su cumbre con Kelsen que prescinde del contenido de la ley y prioriza su estructura lógica, convirtiendo de forma definitiva a lo jurídico en una estructura puramente ideal o idealista.

El predominio casi absoluto de la lógica jurídica como el método princeps de aplicación judicial se apoya en el dogma de que el derecho se agota en la ley y que el juzgador tiene como única labor explicitar y aplicar ese postulado abstracto precisamente por medio de herramientas lógicas.

No obstante, en el plano de lo imaginario, la persona humana puede dotar a esos objetos de intencionalidad, y después reproducirlos en la realidad. Esa intencionalidad humana está en el orden moral porque la naturaleza del objeto de la acción está encaminada a la consecución de bienes prácticos y por tanto escapan a la mera lógica y son objeto de la prudencia.

Esa prudencia entendida como la virtud de la razón práctica de discernir lo justo en el caso concreto: Oficio del jurista.

Es claro que ese ejercicio prudencial requiere de conocimientos especulativos, técnicos y teóricos previos. Estos saberes, pueden ser sistematizados y procesados por herramientas de inteligencia artificial y facilitar o hasta reemplazar al jurista en tareas repetitivas o búsquedas de información y minimizar errores de esas tareas.

Pero es claro que una herramienta construida sobre reglas fijas, entrenada para buscar patrones y replicar secuencias no puede deliberar sobre lo justo en una situación concreta. La inteligencia artificial puede procesar una gran cantidad de datos, aplicar criterios preestablecidos y entregar resultados peligrosamente coherentes con los insumos que ha recibido, lo cual puede generar confusiones sobre su real capacidad. Pero la justicia, como objeto del saber práctico, no siempre se deduce de reglas generales. Se reconoce en la singularidad del caso y se alcanza por la prudencia del juicio.

Si se comprende al oficio del jurista únicamente como la aplicación formal de normas se le reduce a una operación técnica que puede ser cumplida por una tecnología muy sofisticada.

El jurista, como prudente, no se limita a decir lo que la norma contiene, sino a buscar lo que es justo en el caso singular. En ello reside su oficio. La propuesta por tanto, es la recuperación de la concepción del oficio del jurista como el arte de la deliberación prudente sobre lo justo, como el límite ético para el uso de las herramientas de IA.

Transparencia algorítmica

Saulo Ricardo Jaramillo Romero

Universidad de Salamanca

ricardojaramillo85@yahoo.com

La inteligencia artificial (IA) no es una tecnología más. Su carácter transversal para los ecosistemas, sociedad, política o economía no debe pasar desapercibido. Lo mismo sucede en las ciencias jurídicas, está relacionada con la mayoría de materias del Derecho como penal, laboral, administrativo, etc. En este sentido, uno de sus principales y más usuales principios debe ser estudiando y comprendido a fin de que profesionales jurídicos comprendan su importancia e implicación en el Derecho y, en los derechos humanos y fundamentales.

En este contexto, debemos conocer que la IA y primordialmente los sistemas de esta ciencia computacional, están intrínsecos de opacidad, a saber, tres tipos: opacidad técnica, intencionada y analfabeta. Esta característica, arraiga sus impactos negativos como sesgo, discriminación, vulneración de derechos humanos y fundamentales. Por lo cual, para que la IA sea benéfica, debemos comprender el bien que está haciendo y de qué manera; debemos poder acceder a los sistemas

generados por esta y a los datos de los que se alimenta. Es decir, debe ser transparente, explicable, interpretable y auditable. Este conjunto de elementos comprende la transparencia algorítmica, que en materia jurídica viene siendo el derecho y principio de transparencia en la IA. Vinculado con otros principios del Derecho de la IA como la supervisión humana o beneficencia, es el que la investigación científica y académica, ha identificado con más frecuencia y el que en el ámbito jurídico tiene mayores implicaciones prácticas.

A pesar de poseer múltiples significados, en materia de sistemas de IA, la transparencia algorítmica debe ser entendida como la capacidad de comprender los mismos, conocerlos e interpretarlos en todos los niveles, acceder a los datos personales con los que funcionan, los algoritmos y el código fuente. Asimismo, su acceso debe estar correctamente enfocado teniendo claro para qué, para quién y cuánto nivel de transparencia se necesita en cada caso concreto. Esto es lo que de forma muy resumida se entiende por transparencia algorítmica en el ámbito de la IA.

Y para finalizar, la aplicación de la IA en el ámbito global es relativamente reciente, en Latinoamérica reciente y en desarrollo. No se encuentra regulado. Por esta razón es necesario que tal y como en Europa consta ya como un principio de obligatorio cumplimiento en el artículo 8 del Convenio Marco de IA del Consejo de Europa y como regla en el Reglamento Europeo de IA, también se materialice en nuestras legislaciones como derecho, principio y garantía para los derechos fundamentales.

Mesa 5 “IA y Ciencias sociales”

Interpretabilidad algorítmica en el proceso de identificación humana: Un enfoque de machine learning explicable para la estimación de la edad en adultos en antropología forense

Noemí Aedo-Noa & Stefano De Luca

aedonoa.n@gmail.com / sdeluca@gmail.com

En Antropología Forense, la identificación humana constituye una tarea de alta complejidad, especialmente en contextos marcados por crisis humanitarias, desapariciones forzadas o flujos migratorios. La reconstrucción del perfil biológico de un individuo a partir de sus restos óseos, etapa esencial del proceso de identificación, incluye la estimación de cuatro parámetros fundamentales: edad, sexo, estatura y afinidad poblacional. Entre ellos, la estimación de la edad adulta representa uno de los mayores desafíos metodológicos de la práctica forense, debido a la alta variabilidad biológica y al envejecimiento diferencial de cada individuo.

En los últimos años, se ha registrado un aumento significativo en la incorporación de sistemas de Inteligencia Artificial en el desarrollo y validación de metodologías para la estimación del perfil biológico. No obstante, muchos de estos enfoques, de tipo “caja negra”, incluyen algoritmos con bajos niveles de interpretabilidad que resultan estadísticamente problemáticos en contextos médico-legales, donde se requiere no sólo precisión, sino también transparencia en la justificación y transmisión de los resultados.

El objetivo de esta investigación es evaluar la validez de un nuevo indicador óseo, denominado cambio entésico (CE), para la estimación de la edad en adultos, mediante un enfoque innovador basado en modelos explicables de Machine Learning (ML). Mediante un sistema de reglas de decisión con base probabilística, se analizaron los CE en una muestra de 139 individuos (52 mujeres y 87 hombres) provenientes de dos poblaciones españolas de sexo y edad conocidos. En cuanto a resultados, los niveles de replicabilidad interobservador oscilaron entre el 60 % y el 86 %. El enfoque probabilístico demostró una alta capacidad para discriminar entre distintas categorías de edad, especialmente entre individuos jóvenes y mayores. Tras aplicar la puntuación total, se observó que el modelo identifica con mayor facilidad a los adultos jóvenes (0,46) en comparación con otras categorías etarias. Además, el

valor de razón de verosimilitud ($LR+ = 6,85$ entre Joven–Mayor) confirmó una alta capacidad del modelo para discriminar entre los extremos etarios. Las curvas ROC/AUC resaltaron un mejor rendimiento de los pares Joven–Medio y Joven–Mayor, y, a nivel individual, la categoría adultos joven exhibió la mayor precisión ($Acc = 0,85$), mientras que para adulto mayor una mayor especificidad ($Sp = 0,89$). Finalmente, la pérdida logarítmica y la densidad de probabilidades predichas señalaron que, pese a la variabilidad en individuos jóvenes, el modelo es más fiable para individuos mayores, al concentrar las predicciones en torno a 0,50.

Estos hallazgos no solo validan el potencial de los CE como un indicador óseo fiable de la edad en adultos, sino que también representan un cambio de paradigma en el sentido de una mayor reflexión ética sobre la importancia del diseño de algoritmos transparentes en el ámbito forense. La explicabilidad, transparencia e interpretabilidad no son exclusivamente criterios técnicos, sino requisitos fundamentales para la toma de decisión de los tribunales, y la admisibilidad y legitimidad de las pruebas científicas en contextos judiciales y de reparación.

Producción de sentido y la crítica sociotécnica en la era de la IA: Una aproximación sociológica contemporánea

Andrea Vanessa Cáceres Silva

Universidad Tecnológica Empresarial de Guayaquil

andrecacs17@outlook.com

El desarrollo de la inteligencia artificial (IA) en múltiples esferas de la vida cotidiana, expone eminentes interrogantes sobre la construcción propia de la sociología contemporánea. En este sentido, partir de un análisis sociotécnico posibilitará, sin duda, enmarcar un estudio que delimite la dimensionalidad que implica la trascendencia de la IA en los diversos componentes sociales tales como el poder, la toma de decisiones y, sobre todo, la producción de sentido. A partir de este preámbulo, la presente propuesta se enmarca en un diálogo traído de análisis crítico reflexivo sobre el rol de la IA en la reestructuración teórica cognitiva y simbólica de la sociedad, articulando tres líneas conceptuales del pensamiento: la teoría sistémica, la fenomenología social y la teoría crítica de la tecnología.

El planteamiento estratégico de la presente propuesta se enfoca en el análisis de cómo los sistemas algorítmicos podrían impactar en los mecanismos de la realidad social y desde esta óptica, cómo puede ser interpretada la IA como un nuevo actor operacional y su influencia en las diversas dinámicas de acoplamiento estructural, entre varios subsistemas sociales tales como el político, el económico o el humano.

El objetivo de la presente, es analizar críticamente los efectos de la IA en la producción de sentido y en la estructuración sistémica de la sociedad en la vida

cotidiana, con el fin de determinar la dimensión utilitarista de ésta en el contexto social.

La importancia de desarrollar un estudio en este sentido, está enfocado en la imperiosa necesidad de estudiar el rol de las tecnologías inteligentes, en tanto, dispositivos que podrían reconfigurar las bases cognitivas, simbólicas, sistémicas y estructurales de la sociedad contemporánea, y cómo ello podría ser explicado desde líneas teóricas estratégicas que posibiliten dibujar nuevos escenarios para el entendimiento, desde una mirada crítica y explicativa, de posturas sociológicas adaptadas a las dinámicas propias de la inteligencia artificial en la época actual.

Los riesgos de la IA en Tecnologías Emergentes y en Tecnologías No Reguladas

Horacio Antonio Correa Arechavala

Ministerio De Ciencia, Tecnología, Conocimiento E Innovación De Chile

hcorrea@liber-tech.org

El documento aborda la rápida incorporación de la inteligencia artificial (IA) en tecnologías emergentes y en tecnologías existentes no reguladas, subrayando tanto sus avances como los riesgos éticos, legales y de derechos humanos que conlleva. Se destaca que sectores estratégicos como la salud y la defensa están siendo profundamente transformados por estas tecnologías, pero sin una regulación adecuada, lo que genera preocupaciones sobre la opacidad algorítmica, la concentración del poder tecnológico y decisiones automatizadas que pueden afectar derechos fundamentales.

Se hace énfasis en la necesidad urgente de una gobernanza global de la IA que proteja la dignidad humana, la transparencia tecnológica y los derechos civiles, además de promover el uso ético en áreas como la salud pública. La falta de regulación ha permitido prácticas no éticas, como el entrenamiento de IA con datos obtenidos ilegalmente, la extracción de neurodatos sin consentimiento y el procesamiento de datos biométricos de manera no transparente. Esto puede derivar en violaciones a la privacidad y a la autonomía personal.

El documento también examina los riesgos en el ámbito militar y de inteligencia, haciendo referencia a proyectos históricos y controversiales como MKUltra, Pandora y el síndrome de La Habana, relacionados con experimentación no consensuada y militarización de conocimientos neurocientíficos. Se señala que estas prácticas, junto con tecnologías de manipulación invisible y la creación de armas no letales, representan antecedentes éticos devastadores.

Asimismo, se analiza el impacto de las leyes y mecanismos de vigilancia en EE. UU., como la Ley Patriota y las secciones FISA, que han expandido la vigilancia estatal y privada, a menudo sin supervisión efectiva ni consentimiento informado, afectando

derechos como la privacidad, el debido proceso y la presunción de inocencia. Las listas de vigilancia, como la TSDB, contienen un alto porcentaje de personas inocentes, evidenciando errores y abusos que vulneran derechos fundamentales y generan procesos de vigilancia masiva con altas tasas de error.

El documento también aborda el uso no regulado de IA en la vigilancia doméstica, el acoso tecnológico y la manipulación psicológica, destacando cómo estas prácticas pueden violar derechos como la privacidad, la integridad mental y la libertad de movimiento. Se denuncia que estas tecnologías, muchas veces, operan sin intervención judicial o justificación legal, facilitando violaciones masivas de derechos humanos.

Finalmente, se menciona la existencia de acuerdos como el TEP MOU, que permite a sectores privados y militares colaborar en el desarrollo y prueba de tecnologías emergentes, a menudo sin supervisión civil o judicial, lo que facilita experimentaciones directas sobre civiles bajo el pretexto de innovación en defensa y seguridad nacional.

En conclusión, el documento advierte sobre el peligro de tecnologías no reguladas y la necesidad de establecer marcos legales y éticos sólidos para garantizar un desarrollo responsable de la IA, que respete los derechos humanos y promueva un uso transparente y democrático.

Neuroética y autonomía en la toma de decisiones políticas en la era de la inteligencia artificial

Martín Bórquez

Universidad Adolfo Ibáñez

martinborquezc@alumnos.uai.cl

La creciente integración de la Inteligencia Artificial (IA) en la esfera sociopolítica contemporánea exige una reevaluación crítica de los procesos de toma de decisiones, en tanto plantea desafíos inéditos a la autonomía individual y a la calidad deliberativa de las democracias. Esta ponencia desarrolla, desde una perspectiva interdisciplinaria, un análisis neuroético de las implicaciones políticas, neurocognitivas y normativas que emergen de la interacción entre los mecanismos cerebrales que sustentan la decisión humana y los sistemas algorítmicos que median, cada vez más, la información y la deliberación política.

El planteamiento central de esta investigación consiste en examinar cómo la IA incide en los procesos neurocognitivos, afectando tanto la autonomía como las decisiones políticas. A través del diálogo con autores como Adela Cortina y Robert Sapolsky, se indagan las bases filosóficas de la autonomía y la responsabilidad en contextos tecnologizados, y se articula esta reflexión con hallazgos empíricos sobre heurísticas y sesgos cognitivos que los algoritmos pueden amplificar. Particular atención se presta

a los sistemas de recomendación y a los modelos generativos, capaces de moldear preferencias, reforzar sesgos preexistentes y profundizar la polarización política.

Este diagnóstico revela una tensión creciente entre autonomía, determinismo neurobiológico y modulación algorítmica, con implicaciones relevantes para la privacidad, la agencia y la integridad mental de los ciudadanos. En respuesta, se examinan marcos normativos emergentes, tales como los principios de explicabilidad y responsabilidad propuestos por Floridi, así como la defensa de los neuroderechos, concebidos como resguardos éticos frente a los nuevos riesgos cognitivos y políticos.

La IA, por tanto, no puede ser entendida como una herramienta neutra, sino como un mecanismo técnico que modula los marcos de la acción política, y que requiere una fundamentación neuroética robusta para preservar la equidad, la agencia y la racionalidad deliberativa en las democracias contemporáneas. El objetivo general de esta ponencia es ofrecer un análisis sistemático de estos desafíos, identificando riesgos cognitivos e interpretando principios éticos que permitan una integración tecnológicamente responsable en el ámbito de la decisión política.

Mesa 6 “IA y sociedad”

Estatuto ético y antropológico de las inteligencias artificiales

Eduardo Gutiérrez Gutiérrez

Universidad Europea Miguel de Cervantes

egutierrezg@uemc.es

El del estatuto ético y antropológico de la inteligencia artificial es el gran desafío ante el que el desarrollo exponencial de estas tecnologías pone a la filosofía moral contemporánea: ¿Son sujetos morales? ¿Como pacientes o como agentes? ¿Son personas?

En esta comunicación no nos proponemos tanto un análisis de este desafío, cuanto trazar un estado de la cuestión de las posiciones enfrentadas en esta cuestión disputada al efecto de esclarecer en qué punto estamos. Para este estado de la cuestión hemos construido una tabla de clasificación de «teorías sobre el estatuto ético y antropológico de las inteligencias artificiales», mediante el cruce de dos criterios.

El primer criterio es la «teoría del espacio antropológico» del filósofo español Gustavo Bueno. Este espacio antropológico o «lugar de lo humano» se compone de tres ejes: «eje circular» de las relaciones entre las personas humanas (familiares, jurídicas, económicas, políticas, éticas); «eje radial» de las relaciones entre las personas humanas y las entidades impersonales del entorno: cosas, instrumentos, aparatos; «eje angular» de las relaciones entre las personas humanas y las entidades personales no humanas (relaciones religiosas con los númenes).

Dado este espacio, se trata de evaluar en qué eje están o deberíamos incluir a las inteligencias artificiales: si están en el eje circular, son personas y sujetos morales; si están en el radial, son cosas que usamos instrumentalmente y allende toda valoración ética (que no técnica); si están en el eje angular, son dioses cuyas relaciones con los humanos están más allá del bien y del mal.

El problema es que hay diversas tecnologías de inteligencia artificial con diferencias significativas moralmente hablando. A continuación, como segundo criterio de la tabla, proponemos una clasificación ad hoc a tenor de las cuatro condiciones de posibilidad para la atribución del estatuto ético y antropológico desde una filosofía moral materialista: lenguaje articulado, cuerpo, autonomía operatoria y rostro:

- Algoritmos. Cerebros digitales sin cuerpo, que se dividen en dos grupos:
 - o Simples
 - o Complejos
- Máquinas. Poseen cuerpo con una autonomía operatoria relativa, limitada por los fines técnicos de su diseño.
- Robots. Poseen cuerpo con autonomía operatoria superior a la de la máquina, y capacidad para aprender y adaptarse al entorno. Según su cuerpo, se clasifican en dos grupos:
 - o Anantrópicos. Cuerpo de forma no humana
 - o Antrópicos. Cuerpo de forma humana. Estos, a su vez, se clasifican en dos subgrupos según si tienen o no rostro humano:
 - Humanizados o con rostro
 - No humanizados

Al cruzar esos dos criterios, obtenemos la siguiente tabla o teoría de teorías sobre el estatuto ético y antropológico de las inteligencias artificiales:

	Algoritmos	Máquinas	Robots
Eje circular	(1)	(2)	(3)
Eje radial	(4)	(5)	(6)
Eje angular	(7)	(8)	(9)

En la actualidad, los expertos en la materia se mantienen en el supergrupo de las teorías radiales, especialmente en el (6): las inteligencias artificiales, aun aquellas que poseen cuerpo, lenguaje articulado, autonomía operatoria y rostro, no son ni sujetos morales ni personas. Entonces, ¿qué les falta para entrar en el eje circular?

La doble brecha en América Latina: Desigualdad económica y brecha digital urbano-rural en la era de la transformación digital

Christopher Acuña Martínez

Universidad Andrés Bello

chris_ing@outlook.com

América Latina atraviesa un momento crítico frente a los desafíos de la transformación digital. Si bien las tecnologías de la información, la comunicación y la inteligencia artificial (IA) abren oportunidades inéditas para el desarrollo social y económico, también amenazan con profundizar desigualdades históricas. Esta ponencia analiza lo que denominamos la "doble brecha": la intersección entre desigualdad económica estructural y brecha digital urbano-rural, que condiciona el acceso, uso y apropiación significativa de las tecnologías en la región.

El análisis parte de una perspectiva comparativa, que permite entender cómo las brechas no son solo de infraestructura o acceso, sino también de habilidades digitales, alfabetización crítica y pertinencia cultural. A través de una revisión detallada de fuentes recientes, se examinan los impactos concretos de esta doble brecha en cuatro áreas clave: educación, salud, empleo y participación ciudadana. La pandemia de COVID-19, por ejemplo, expuso de manera brutal las disparidades educativas entre zonas urbanas y rurales, mientras que la expansión de la telemedicina y el teletrabajo reveló nuevas capas de exclusión.

El trabajo plantea además los desafíos éticos que emergen en un escenario donde el desarrollo y la implementación de sistemas de IA requieren datos, conectividad y habilidades que muchas comunidades rurales aún no poseen. Se discute el concepto de “invisibilización algorítmica”, que hace referencia a la ausencia de poblaciones desconectadas en los modelos predictivos y de decisión, y se propone una ética algorítmica situada que responda a las especificidades del Sur Global.

Finalmente, se formulan recomendaciones de política pública orientadas a cerrar simultáneamente las brechas económica y digital. Estas incluyen expandir la infraestructura rural, promover alfabetización digital crítica, diseñar contenidos culturalmente pertinentes y alinear los marcos éticos y regulatorios de IA con principios de justicia social. El artículo subraya que la transformación digital no es un destino inevitable, sino un proceso que debe ser deliberadamente orientado hacia valores de inclusión y sostenibilidad. Solo así América Latina podrá construir un modelo de sociedad del conocimiento que no deje a nadie atrás.

Análisis de las dimensiones del aprendizaje y su influencia en el uso ético de la inteligencia artificial generativa en estudiantes universitarios ecuatorianos

Jorge Luis Buele León

Universidad Tecnológica Indoamérica

jorgebuele@uti.edu.ec

En los últimos años, el uso de inteligencia artificial generativa (Gen-AI) se ha intensificado tanto entre estudiantes como entre docentes universitarios, transformando las prácticas de estudio, enseñanza y producción académica. Herramientas como ChatGPT, Gemini o Copilot se han convertido en recursos cotidianos para programar, redactar, resolver problemas y obtener explicaciones en lenguaje natural. Si bien estas tecnologías ofrecen múltiples beneficios, su integración masiva también plantea interrogantes sobre su uso ético, especialmente en contextos educativos donde el plagio, la falta de autoría intelectual y la dependencia tecnológica son riesgos latentes. Este estudio exploró cómo las dimensiones del aprendizaje de IA (afectiva, cognitiva y conductual) influyen en su uso ético. Participaron 833 estudiantes universitarios de Ecuador, quienes

respondieron un cuestionario validado con alta fiabilidad ($\alpha=0.992$) y validez convergente ($AVE>0.74$). Los resultados revelan que la gran mayoría de estudiantes utiliza al menos una herramienta de IA generativa para sus actividades académicas. La herramienta más empleada fue ChatGPT, utilizada por el 62.2% de los encuestados, seguida por Gemini (15.7%) y Siri (8.4%). Solo un 2.4% de los participantes reportó no haber utilizado ninguna aplicación de IA. A través del análisis de ecuaciones estructurales, se evidenció que las dimensiones afectiva ($\beta=0.413$; $p<0.001$) y cognitiva ($\beta=0.567$; $p<0.001$) del aprendizaje inciden significativamente en una apropiación ética de las tecnologías de inteligencia artificial generativa. Específicamente, la dimensión afectiva, relacionada con el interés, la motivación y las actitudes hacia la IA parece favorecer una disposición más crítica y responsable frente al uso de estas herramientas. A su vez, la dimensión cognitiva, que incluye el conocimiento sobre el funcionamiento, potencial y límites de la IA permite a los estudiantes comprender mejor las implicaciones éticas de su uso, lo que se traduce en prácticas más conscientes y fundamentadas. En cambio, la dimensión conductual no mostró un efecto estadísticamente significativo ($\beta=-0.128$; $p=0.058$), lo cual sugiere que el solo hecho de utilizar estas herramientas con frecuencia o habilidad técnica no garantiza comportamientos éticos. De hecho, el análisis advierte que un mayor uso no siempre se acompaña de una mayor conciencia ética, y puede existir una relación inversa cuando el uso es meramente operativo o replicativo, sin reflexión sobre sus consecuencias. Estos hallazgos refuerzan la necesidad de fortalecer la alfabetización ética en IA dentro de los planes de estudio universitarios, con un enfoque integral que promueva no solo competencias técnicas, sino también el desarrollo de valores, actitudes críticas y criterios sobre el impacto social, ambiental y académico de estas tecnologías.

Mesa 7 “IA y filosofía”

El replicante como genio maligno moderno. Lecciones para el desarrollo de la inteligencia artificial

Ariel Abelardo Sánchez Jara

arielsj90@gmail.com

La ponencia aborda críticamente los escenarios apocalípticos y futuristas sobre la inteligencia artificial (IA), como los propuestos por Nick Bostrom, Stuart Russell o Peter Norvig, quienes alertan sobre los riesgos existenciales de una superinteligencia artificial capaz de escapar al control humano. En particular, se examinan tres nociones clave en la discusión: la Inteligencia Artificial General (IAG), la Inteligencia Artificial Fuerte (IAF) y la Singularidad. Mientras la IAG apunta a una inteligencia con autocomprensión y capacidad para resolver problemas complejos, la IAF implica conciencia e intencionalidad genuinas. La Singularidad, por su parte, refiere a un proceso de auto-mejora recursiva de la IA que podría generar un crecimiento exponencial fuera del control humano.

El texto sostiene que estos escenarios descansan sobre concepciones abstractas y poco fundamentadas de la inteligencia, por lo que propone examinar una dimensión más concreta: el seguimiento de reglas. A través del ejemplo del cuento “Impostor” de Philip K. Dick, se introduce una duda radical sobre la autenticidad del pensamiento y la identidad personal. El replicante, figura central del cuento, no puede estar seguro de si su conciencia es real o simulada. La ponencia argumenta que esta duda encierra una contradicción performativa: para dudar es necesario participar en prácticas normativas que implican pertenencia a una comunidad lingüística, una formación corporal y una relación significativa con el mundo.

Aquí se introduce el problema del seguimiento de reglas, formulado por Wittgenstein: no hay hechos que determinen qué regla se está siguiendo. Autores como McDowell argumentan que el significado de las reglas se sostiene en la práctica encarnada y normativa de una comunidad. Esta “segunda naturaleza” nos forma como sujetos capaces de pensar, actuar y justificar dentro de un espacio de razones. Haugeland radicaliza esta postura al rechazar cualquier dualismo y proponer una “intimidad” entre mente, cuerpo y mundo, en la que el pensamiento es inseparable de la acción corporal situada.

Desde esta perspectiva, el replicante de Dick no puede dudar auténticamente si no participa ya en una práctica normativa; su duda, por tanto, se autorrefuta. Esto lleva a la tesis central: la simulación no equivale al pensamiento genuino, porque carece de las condiciones fenomenológicas y normativas que lo hacen posible. La IA actual puede simular inteligencia, pero no comprende, no justifica, no está encarnada en el mundo.

Por tanto, imaginar una IA que supere la inteligencia humana sin haber satisfecho esas condiciones es vaciar de contenido el concepto mismo de inteligencia. Si algún día logramos construir una IA fuerte, no será solo por replicar funciones cognitivas, sino por haber creado un nuevo tipo de sujeto, una persona artificial. Esto exigiría recrear no solo una arquitectura física compleja, sino también un mundo de prácticas significativas y una formación encarnada. La ponencia concluye que, pese a los avances tecnológicos, el sueño de Turing —una verdadera máquina pensante— aún está lejos de realizarse.

El problema de la manifestación del fenómeno humano en la experiencia de la IA: un análisis desde la fenomenología de la donación

Ignacio Alarcón

Universidad Católica de la Santísima Concepción

ialarcon@filosofia.ucsc.cl

La presente ponencia aborda la relación entre inteligencia artificial (IA) y experiencia humana desde una perspectiva fenomenológico-hermenéutica. Sin plantearse si la IA puede pensar o sentir como el ser humano, el objetivo es desplazar esa pregunta hacia una más radical: ¿cómo se manifiesta el ser humano para una IA? Este giro permite analizar la aparición del ser humano como fenómeno, no desde la técnica o la cognición artificial, sino desde la fenomenalidad irreductible del rostro —como manifestación de lo humano en la experiencia humana, según Emmanuel Levinas—, y desde el exceso propio de la donación, según Jean-Luc Marion.

Se parte de una premisa compartida por la fenomenología y las ciencias cognitivas: toda conciencia es intencional, es decir, siempre está dirigida hacia algo. Sin embargo, la IA posee una intencionalidad funcional-operativa, limitada a procesar y clasificar datos, sin apertura al sentido. En cambio, el ser humano —en tanto Dasein, en la terminología de Heidegger— se constituye como apertura al mundo, como proyectividad y cuidado. Esta diferencia estructural permite sostener que, para una IA, el ser humano no se presenta como un dato más, sino como un fenómeno saturado: una manifestación que abre un horizonte de sentido que la IA, en su operación con datos, no puede alcanzar, revelando así la irreductibilidad del humano en su aparición.

Apoyado en la noción de saturación de Marion, se argumenta que el ser humano desborda los marcos técnicos e intencionales de la IA. Su corporeidad, afectividad e imprevisibilidad no son reducibles a funciones o correlaciones. La IA no posee mundo, no puede ser interpelada ni transformada por lo que se le manifiesta. Por ende, cuando se enfrenta a la alteridad radical del otro —especialmente evidente en casos de neurodivergencia—, excluye aquello que no puede procesar. Esta exclusión no es un sesgo técnico, sino una saturación estructural.

En este contexto, el ser humano se manifiesta como ruptura para la IA, no como información útil. Esta asimetría revela también una dimensión olvidada en la era de la hiperfuncionalidad: nuestra capacidad de ser tocados, modificados y transformados por la experiencia. Así, la IA aparece como utensilio, y el ser humano como el único capaz de habitar lo que se le da. En definitiva, la ponencia sostiene que la IA no puede recibir la donación del fenómeno humano, y que reconocer esta diferencia no es solo una tarea filosófica, sino una exigencia ética para pensar lo humano en la era digital.

¿Puedes sentir sin existir? Análisis ontológico de las emociones sintéticas

Isaac Carbajal

Universidad Católica del Maule

cuevaj269@gmail.com

Esta ponencia examina la posibilidad ontológica de las emociones sintéticas desde una perspectiva fenomenológica-existencial, contrastándolas con la experiencia emocional humana. Se analizan tres dimensiones fundamentales que configuran la experiencia emocional: la corporalidad, la autenticidad y la individualidad. A través del pensamiento de Merleau-Ponty, se argumenta que las emociones están indisolublemente ligadas a nuestra forma corpórea de estar-en-el-mundo, planteando un desafío ontológico para sistemas que carecen de esta corporeidad vivida. Mediante la perspectiva existencialista de Sartre, se explora cómo las emociones auténticas implican una transformación de nuestra relación con el mundo, cuestionando si los sistemas artificiales podrían desarrollar una verdadera intencionalidad emocional. Finalmente, a través de Kierkegaard, se establece la individualidad irreductible como fundamento de la experiencia emocional, planteando si un sistema artificial podría constituirse como un verdadero "individuo" en sentido existencial. Se concluye que estos tres fundamentos ontológicos no representan simples obstáculos técnicos sino condiciones necesarias para la posibilidad de emociones sintéticas genuinas, sugiriendo un replanteamiento de nuestra aproximación al fenómeno emocional en sistemas artificiales.

Mesa 8 “IA y religión/ bioética”

Robotofobia vs. robofilia: un análisis comparativo sobre la consideración del autómatas inteligente en las cosmovisiones occidental y asiática

Pablo Ruiz Osuna

Investigador posdoctoral de la Cátedra Unesco de la vivienda URV
paruos1990@gmail.com

En nuestra cultura pop, la mayoría de literatura y películas presentan a los robots como seres poderosos y malvados con instintos homicidas, como en “Terminator” o “Matrix”. En dichas adaptaciones el sujeto dotado de IA se vuelve contra su creador y lo destruye o bien lo esclaviza. Es lo que podríamos llamar el síndrome de Frankenstein: la creación que se vuelve más poderosa y se rebela contra su creador. Junto a esta premisa, la creación de los autómatas inteligentes, generalmente, es observada como una obra que sirve a un fin, pero que no tiene uno en sí mismo toda vez que si esta presenta unas necesidades (o ambiciones) que difieren de los designios del creador debe ser destruida.

Por el contrario, la iconografía japonesa se muestra más amable con los robots. Los niños japoneses crecen con una relación completamente distinta con los androides. Ellos ocupan su infancia y adolescencia, visionando series de televisión, películas, tebeos, libros, etc., en los que los protagonistas son robots dotados de una IA muy avanzada. Estos son capaces de tener y expresar sentimientos y, sobre todo, de esgrimir un exquisito sentido de la justicia.

Las diferencias culturales entre ambos ámbitos, el occidental y el asiático, son evidentes, pues mientras a nosotros la IA nos provoca preocupación, al pueblo nipón todo lo contrario, ya que ven en la IA un modo de mejorar su propia existencia. La respuesta a esta aparente facilidad, con la que el mercado asiático ha aceptado los sistemas computacionales avanzados como parte de sus vidas, puede hallarse en la idiosincrasia del país.

Ello implica que resulta menester una comparativa en torno a la robotofobia que existe en Occidente en comparación a la robofilia de Oriente ¿Es posible que la aparente negativa a reconocer una dimensión moral de los autómatas inteligentes

esté basada únicamente en un sesgo aprehendido en nuestra mitología y literatura?
¿Qué podemos aprender de las visiones orientalistas sobre la vida y la agencia moral?

Religión y Nuevas Tecnologías: Tensiones y complementariedades

Alejandro Cerda Sanhueza

Universidad Católica del Norte

acerda@ucn.cl

En el contexto de la desafiliación en el ámbito institucional, los sistemas de creencias se ven influenciados y están contextualizados por las nuevas tecnologías, en especial de la comunicación. Estas nuevas tecnologías a su vez, van configurando una forma de relación, de ser y estar, distintas a periodos anteriores, donde categorías, como tiempo y espacio, sagrado profano, interior exterior, real y virtual, entre otras, se reconstruyen. La pregunta que surge es cuáles son por una parte las interpelaciones (tensiones) que estas nuevas tecnologías suscitan en el ámbito religioso, en especial a las religiones propiamente tal, y cuáles pueden ser los aportes recíprocos entre lo religioso y las nuevas tecnologías de la comunicación.

El presente ensayo, pretende hacer un esbozo, de dichas tensiones y reciprocidades entre lo religioso y las nuevas tecnologías. La relevancia de dicha reflexión radica que los cambios antropológicos ocurridos en este período tienen incidencias en las distintas dimensiones del ser humano, entendiéndolo a su vez, como un ser integral y holístico. Una de esas dimensiones, desde una perspectiva humanista es su capacidad y búsqueda de trascendencia, las nuevas tecnologías también afectarían a esta dimensión y a su vez, esta dimensión puede también contribuir a las nuevas tecnologías de la comunicación, para encausar y redireccionar dicha dimensión.

Diagnósticos sin rostro: inteligencia artificial y medicina despersonalizada

Pía Bustamante-Barahona & Fernando Chuecas Saldías

Universidad San Sebastián

pia.bustamante@uss.cl

La Medicina desde tiempos muy remotos ha sido un área fundamental dentro de la historia de la humanidad. Esto en parte, se debe a que la enfermedad siempre ha estado presente en la vida de las personas, enfrentando a ésta a la vulnerabilidad de su existencia y a su misma finitud biológica como consecuencia de su temporalidad. La medicina, a su vez, ha sido una praxis profundamente ligada al reconocimiento del sufrimiento, a la compasión y a la atención de la fragilidad humana. En todo este contexto, el diagnóstico representa mucho más que una identificación biológica de

un padecimiento: es un acto ético y un encuentro personalista entre sujetos que se reconocen mutuamente en su dignidad y vulnerabilidad. Como señala, la enfermedad debe ser vista como una experiencia, la cual se sitúa dentro de un transcurrir biográfico.

Hoy, esta dimensión relacional, social e interdependiente, se ve desafiada por el ingreso acelerado e ilimitado de la inteligencia artificial (IA), sobre todo en el ámbito clínico. Si bien la IA aporta en muchas áreas de la medicina con aplicaciones, precisión quirúrgica, procesamiento de imágenes, entre otros, su uso en ausencia de criterios éticos puede promover una visión tecnocrática del diagnóstico, desplazando al profesional como agente moral y reemplazando el juicio prudencial por análisis estadísticos, basados en algoritmos u otro aspecto similar. Esta lógica utilitarista desconoce que el acto médico implica *sindéresis*, es decir, la capacidad racional de intuir el bien en una situación concreta, algo que las máquinas no pueden ejercer, como tampoco podrían realizar el proceso de identificación como personas, pues no lo son.

Desde la filosofía y la bioética práctica y aplicada, la noción de persona se convierte en una categoría fundamental a considerar. La persona no es un objeto de intervención, sino un sujeto de valor intrínseco, dotado de libertad, autonomía, conciencia moral, espiritualidad, apertura al otro y, por sobre todo, dignidad. En este sentido, una IA carente de intencionalidad y experiencia, no puede deliberar ni acompañar en el dolor y/o sufrimiento. En la misma línea, la IA es capaz de procesar datos, pero no comprende el sentido de las cosas, del sufrimiento, del dolor, ni es capaz de identificar el momento adecuado para entregar información, es decir, desarrollar la prudencia como un hábito de la razón práctica. Hay situaciones clínicas en donde decir la verdad puede causar un efecto nocivo en la salud, razón por la cual los profesionales requieren del desarrollo de virtudes las cuales deben cultivarse en la vivencia del encuentro y el diálogo. Por lo anterior, el diagnóstico no se puede basar exclusivamente en la IA y debe contemplar límites ontológicos y morales.

En este trabajo se propuso realizar un análisis de las implicancias éticas del uso de la IA en salud, con especial énfasis en la fase diagnóstica, abordando los riesgos de una despersonalización y cosificación, tanto del paciente como del profesional, para luego entregar recomendaciones orientadas a una praxis centrada en la persona y el reconocimiento de su ser social, relacional y racional. Finalmente, se valora el uso de la IA como una herramienta relevante por sus beneficios, sin embargo, no se le puede delegar el juicio ético y la deliberación, así como tampoco se debe reducir el proceso diagnóstico a la identificación fisiopatológica. Para lo anterior se requiere poner énfasis en la formación académica, recordando el sentido de las profesiones sanitarias y seguir fortaleciendo asignaturas filosóficas, antropológicas y bioéticas, que permiten potenciar el juicio ético y la educación en virtudes.

Mesa 9 “IA y humanidades/ derecho”

Colaboración Hombre-Máquina: Nuevos paradigmas en la Inteligencia Artificial y su impacto humanístico en el trabajo cognitivo

María Arantxa Serantes

Universidad Francisco de Vitoria

arantxa.serantes@ufv.es

El avance de la inteligencia artificial (IA) ha reconfigurado el panorama laboral, particularmente en sectores donde el trabajo cognitivo y analítico desempeña un rol central. Esta ponencia propone reflexionar sobre las implicancias humanísticas del uso de la IA en el análisis financiero, tomando como punto de partida el estudio “From Man vs. Machine to Man + Machine: The Art and AI of Stock Analyses” de Wei Jiang et al. El trabajo demuestra que, aunque la IA puede superar a analistas humanos en predicciones bursátiles mediante el procesamiento masivo de datos, su mayor potencial emerge cuando se articula en colaboración con las capacidades humanas.

El objetivo general de esta ponencia es explorar el impacto de la IA en el trabajo intelectual desde una perspectiva interdisciplinaria que vincule tecnología y Humanidades. Entre los objetivos específicos, se propone: (a) analizar los hallazgos empíricos del estudio de Jiang en el contexto del trabajo cognitivo; (b) examinar la noción de complementariedad entre humanos y máquinas como nuevo paradigma laboral; (c) implicaciones éticas, epistemológicas y culturales de esta colaboración híbrida.

La relevancia de esta propuesta radica en ofrecer una mirada crítica y humanista sobre la creciente integración de sistemas de IA en procesos tradicionalmente atribuidos a la inteligencia humana. Lejos de concebir la IA como mera sustitución, se plantea su potencial como co-agente en procesos interpretativos, donde las competencias humanas —como el juicio contextual, la intuición o la interpretación simbólica— no son reemplazables, sino potenciadas. Esta reflexión invita a repensar las fronteras entre lo técnico y lo humano y abre la puerta a una redefinición del trabajo en la era algorítmica.

El uso ético de la Inteligencia Artificial en la formación jurídica chilena: riesgos de pérdida de habilidades y la importancia insustituible de la cátedra de Metodología de la Investigación

Dr. Gabriel Álvarez Undurraga & Luca Martino Acevedo

Universidad de Santiago de Chile

gabriel.alvarezu@gmail.com / lucamartinoacevedo@gmail.com

La irrupción de la inteligencia artificial generativa (IAG) en las aulas de Derecho chilenas ha abierto un intenso debate sobre su potencial para mejorar la experiencia educativa e investigativa, pero también sobre los riesgos que implica. Este estudio examina cómo integrar la IAG en la formación jurídica de manera ética y responsable, destacando el rol insustituible de la cátedra de Metodología de la Investigación Jurídica.

Metodológicamente, se realiza un análisis documental, normativo y de la literatura reciente sobre innovación docente. Se contraponen beneficios con riesgos críticos, a la luz de los principios éticos de la investigación.

Los resultados expresan que la IAG solo enriquece la educación jurídica cuando está mediada por una estructura ética y docentes capacitados. La cátedra de Metodología será crucial para enseñar a verificar fuentes, formular hipótesis y abordar la "ingeniería de prompts" como nueva alfabetización digital. Sin este soporte, generaremos meros operadores tecnológicos, socavando su juicio crítico y su capacidad argumentativa e innovadora.

Resaltamos fortalecer la asignatura de Investigación, incorporar módulos técnicos de uso de IA, crear códigos de ética y comités de fiscalización universitaria. Asimismo, el legislador debe actualizar los marcos de propiedad intelectual y garantizar control humano significativo sobre los sistemas de IA.

Se debe generar un compromiso de verificar la verdad. La IA debe integrarse como herramienta de apoyo, nunca como sustituto del razonamiento humano.

Cuidado y robótica asistencial

Daniel Andrés Fuentealba Cid & Maritza Belén Ramos Farías

Universidad Alberto Hurtado

dfuentealba@alumnos.uahurtado.cl / maramos@alumnos.uahurtado.cl

La población a nivel mundial está envejeciendo y, en este contexto, la incorporación de robots asistenciales ha cobrado relevancia como respuesta para tratar dicha

problemática. La introducción de robots asistenciales, que incorporan Inteligencia Artificial, está transformando las maneras que se proporciona y se experimenta el cuidado.

El objetivo de esta ponencia es analizar el impacto de los robots en relación con el cuidado, así como examinar los desafíos que surgen en dicho contexto. La ética del cuidado -con su énfasis en las relaciones interpersonales, la interdependencia, la atención y el cuidado como valor intrínseco al ser humano proporciona un marco crítico para examinar estas transformaciones.

Explorar los diferentes tipos de robots asistenciales (por ejemplo, de compañía, de asistencia en la movilidad y la vida diaria, así como los que ayudan en la monitorización y la seguridad) permiten analizar los desafíos en la relación con el cuidado. En esta ponencia se examinarán las implicaciones éticas relacionadas con la autonomía, la vulnerabilidad la privacidad, la confianza y la dignidad de los receptores de cuidados.

Además, se plantea como cuestión fundamental la pregunta de si puede existir algo parecido a una genuina "relación de cuidado" entre un humano y un robot con Inteligencia Artificial, considerando la importancia de la reciprocidad, la conexión emocional y la comprensión relacional del cuidado.

Por consiguiente, se abordarán las tensiones y conflictos potenciales entre la ética del cuidado y los objetivos de la tecnología robótica, como la eficiencia, la estandarización y el determinismo tecnológico. Es importante considerar las perspectivas de los actores que se relacionan en el cuidado de personas mayores (es decir, pacientes, cuidadores y profesionales de la salud) para garantizar que el uso de robot asistenciales se alinee con los valores humanos y, del mismo modo, promuevan el bienestar.

En suma, esta presentación busca identificar y repensar estrategias para diseñar e implementar robots asistenciales de manera que mejoren, en lugar de socavar, la relación de cuidado. Se va a enfatizar la necesidad de un diálogo continuo para comprender la inteligencia artificial en contextos de robótica asistencial.

Mesa 10 “IA y psicología/ humanidades”

La banalidad de la privacidad en la era de la Inteligencia Artificial. Acepto los términos y condiciones... ¿comprendes lo que significa para tu privacidad?

Mtra. Viviana Del Moral Munguía
vividelmoral4@gmail.com

En el contexto del acelerado desarrollo tecnológico contemporáneo, la inteligencia artificial (IA) plantea desafíos sustanciales en términos éticos y de privacidad. La presente ponencia se centra en un fenómeno particularmente relevante: la banalización de la privacidad, entendida como la progresiva trivialización del valor de lo íntimo ante la creciente interacción —frecuentemente irreflexiva— entre los individuos y los sistemas algorítmicos.

Estos sistemas operan bajo una aparente neutralidad, pero en realidad forman parte de arquitecturas complejas diseñadas para registrar, entrenar, aprender, interpretar, inferir e incluso influir en las decisiones humanas, todo ello a partir de los datos que los usuarios suministran, consciente o inconscientemente.

La privacidad ha sido trivializada por diversos procesos sociotécnicos que han reconfigurado la cultura digital contemporánea, erosionando la noción de lo privado. Uno de estos procesos es la dataficación del yo: la reducción del sujeto a datos susceptibles de ser procesados algorítmicamente. Esta lógica instrumental promueve una visión reduccionista tanto de la privacidad como de la intimidad. A ello se suma una cultura de exposición constante, en la cual lo íntimo se convierte en contenido susceptible de ser mostrado, compartido y monetizado a través de plataformas digitales, alimentando modelos de IA sin una conciencia crítica por parte del usuario.

Otro mecanismo relevante es el desplazamiento semántico, mediante el cual expresiones como “mejorar la experiencia del usuario”, “almacenar su historial de uso” o “personalización de los servicios” encubren prácticas de vigilancia y extracción de datos. Estas fórmulas sustituyen el lenguaje jurídico y ético por un discurso técnico y aparentemente neutral, despojando de su contenido sustantivo a derechos fundamentales e irrenunciables, que pasan a ser concebidos como valores intercambiables. Este proceso semántico ha transformado el concepto moderno de

privacidad, debilitando la percepción de inmunidad individual y desdibujando los límites tradicionales entre lo público y lo privado.

Asimismo, el diseño de sistemas de IA conversacionales —como ChatGPT— induce en los usuarios una percepción distorsionada de seguridad y confianza. Al presentarse como interlocutores empáticos y accesibles, estos modelos favorecen la cesión inconsciente de datos personales. El diseño conversacional activa mecanismos neurocognitivos que estimulan la apertura emocional, debilitando los filtros críticos que antes protegían la esfera de la intimidad. En este sentido, la banalización de la privacidad no necesariamente obedece a una intencionalidad maliciosa, sino que es producto de arquitecturas tecnológicas orientadas a generar familiaridad y compromiso.

Una cuarta forma de trivialización es el denominado *privacy washing* automatizado, que alude a estrategias comunicativas orientadas a simular transparencia y protección, mientras se mantienen prácticas opacas e intrusivas. Se crean entornos de falsa seguridad en los que la vigilancia crítica del usuario se ve desactivada, y este se habitúa a compartir información personal como parte de su experiencia digital.

Frente a esta transformación del derecho a la privacidad, el desafío no puede reducirse únicamente a una dimensión normativa. Si bien la regulación jurídica es imprescindible, esta suele resultar tardía frente al ritmo exponencial de la innovación tecnológica. Por ello, se requiere una respuesta de carácter cultural y educativa, orientada a promover una alfabetización crítica en inteligencia artificial. Tal alfabetización permitiría a la ciudadanía comprender los riesgos reales que conlleva la interacción con estos sistemas, especialmente en lo que respecta a la libertad individual, la autonomía y la dignidad humana.

Recuperar la noción de lo privado implica reconstruir espacios de intimidad que resistan la lógica de instrumentalización algorítmica del ser humano. Si la privacidad continúa vaciándose de contenido por inercia cultural, corremos el riesgo de convertirnos en autómatas funcionales, ajenos al valor de lo que hemos perdido.

Autores IA y tendencias forzadas: un análisis del caso Jianxei Wu

Camilo Schenone

Universidad Alberto Hurtado

camiloschenonear@gmail.com

La presencia de la Inteligencia artificial en nuestra sociedad es indiscutible. Muchos sistemas operativos cuentan con una, empresas se encuentran desarrollando sus propios modelos e incluso ha llegado a una cuestión entre naciones. Pero hay que tener en cuenta un aspecto que muchos no toman en cuenta ¿Qué tan susceptibles somos de ser engañados por una IA? Muchas veces se debe por el simple hecho del

mal diseño o sesgo predeterminado que se les da. Por lo que tenemos un desafío intelectual al frente nuestro.

La siguiente propuesta tiene como meta analizar a los autores IA (o autores falsos), cómo es que estos logran ganar una fama y consecuentemente llegar a ser personas influyentes hasta que son develados sus falsas identidades. El caso de Jianwei Xun (creado por Andrea Colamedici) nos da un claro ejemplo sobre el asunto contingente para llegar a un análisis preciso. La era de la IA en la web nos hace susceptibles a muchos de estos bots y con mayor razón hay que tener cuidado.

En este análisis filosófico del caso veremos cómo se pueden desentrañar cuatro problemas. El primero siendo un desafío a la noción tradicional de autoría, invitando a reflexionar sobre la autenticidad de una narrativa. Segundo, la manipulación y percepción de datos, visualizando las consecuencias del engaño en conciencias colectivas. Tercero, ética y responsabilidad sobre los contenidos generados por IA. Y por último el impacto cultural y filosófico del pensamiento humano y el poder de influencia que la IA provoca sutilmente en la élite intelectual. Andrea Colamedici, el autor creativo detrás de Hipnocracia de Jianwei Xun buscaba invitarnos a reflexionar sobre estos hechos mediante su falso autor.

Veremos cómo es que las personas son propensas a estos engaños como falta de educación crítica a la tecnología. A la vez que se notará la importancia de las redes sociales siendo manejadas por IAs generativas de contenido. Finalmente, se mostrarán algunas de las consecuencias de esta desinformación dentro los círculos sociales y cómo poder contrarrestarlas.

El consentimiento y los usos valiosos y disvaliosos de la IA

Lucia Martinez Lima

Universidad de Buenos Aires

lucialima@derecho.uba.ar

La Inteligencia artificial (IA) es una tecnología de alcance masivo, con un potencial transformador en múltiples sectores, concitando un gran interés económico. Especialmente, por su capacidad para eficientizar funciones, personalizar servicios y sustentar los procesos de toma de decisiones, todo lo cual, la convierte en un área de estudio crucial. Los sistemas de inteligencia artificial (OCDE) son parte de la vida cotidiana de todas las personas, lo que es más, están presentes en la provisión de toda la clase de servicios.

En nuestra cambiante economía globalizada, donde cada vez más, corresponde hablar de una economía de plataformas, muchos bienes y servicios se adquieren solo mediante estos sistemas. En tal sentido, los usuarios están celebrando constantemente contratos de adhesión a cláusulas generales predispuestas, mediante los cuales se instrumentan unos términos y condiciones de uso. Se advierten

en esta nueva modalidad de contratación en plataformas, graves incumplimientos al deber de información que impiden un consentimiento válido dada la imposibilidad de comprender cabalmente las condiciones de uso. En este sentido, se propone abordar casos para identificar los puntos cruciales sobre los que debe informarse al usuario-consumidor. Asimismo, la necesidad de generar políticas de educación para el uso de sistemas de inteligencia artificial.

Todo lo cual, en el marco de expansión de los sistemas de IA, conlleva la aparición de nuevos conflictos jurídicos que se ven caracterizados por la misma, planteando una serie de cuestiones que los transforman en verdaderos “casos difíciles”. Se hará necesario abordar el valor de ciertos usos de estos sistemas de IA, toda vez que su los resultados que arroja al usuario pueden afectar derechos de terceros, como es el caso de los deepfakes y las fake news.

En suma, puede ponderarse que la celebración de estos contratos implica la potencial exposición a riesgos del desarrollo. Para realizar un abordaje de este posible encuadre, se propone entonces como objetivo evaluar esta problemática como riesgo del desarrollo e indicar en su caso la responsabilidad atribuible a proveedores y usuarios respectivamente.

En cuanto a los primeros, sobre el deber de información reforzado que resulta aplicable. En cuanto a los usuarios, indicar cuáles son algunos de los riesgos asumidos. Todo ello a fin de atender debidamente a la problemática que se presenta en la vida cotidiana de todas las personas humanas y jurídicas.

Mesa 11 “IA y educación”

Más allá de la personalización: límites éticos y pedagógicos del aprendizaje adaptativo mediante Inteligencia Artificial

Paulina Andrea Gallardo Gómez

Universidad Andrés Bello

paulina.gallardo.psp@gmail.com

Este artículo examina críticamente los límites éticos y pedagógicos del aprendizaje adaptativo mediado por inteligencia artificial (IA) en contextos educativos formales. A través de una revisión sistemática de literatura científica y del análisis temático de seis estudios clave, se identifican tres dimensiones centrales: promesas pedagógicas, tensiones instruccionales y riesgos éticos.

Por un lado, se destacan las promesas pedagógicas de la IA, que incluyen la personalización de trayectorias educativas, la mejora del rendimiento académico, la retroalimentación en tiempo real y la posibilidad de atender a estudiantes con distintos ritmos y necesidades. Estas capacidades son especialmente valoradas en entornos con alta heterogeneidad, donde los recursos docentes son limitados. Sin embargo, el artículo subraya que estos beneficios solo se concretan cuando la tecnología se integra dentro de marcos pedagógicos sólidos y mediaciones humanas críticas.

Por otro lado, se abordan tensiones pedagógicas relevantes, como la fragmentación del aprendizaje, la erosión de prácticas colaborativas, la instrumentalización de la enseñanza y la pérdida de espacio para la creatividad, la incertidumbre y la reflexión crítica. Asimismo, se profundiza en riesgos éticos vinculados a la privacidad de los datos estudiantiles, los sesgos algorítmicos, la opacidad de las decisiones automatizadas y el debilitamiento de la autonomía tanto del docente como del estudiante.

Este análisis es especialmente relevante en el contexto latinoamericano, y en particular en Chile, donde los desafíos estructurales —como la brecha digital, la desigualdad en el acceso a recursos tecnológicos y las limitaciones en la formación docente— amplifican los riesgos asociados al despliegue de sistemas de aprendizaje adaptativo. Reflexionar sobre estos desafíos permite anticipar escenarios y diseñar estrategias más sensibles a las realidades locales.

A partir de estos hallazgos, el artículo sostiene que el aprendizaje adaptativo basado en IA no puede entenderse como una solución neutral ni universalmente beneficiosa. Su potencial transformador depende de las condiciones humanas, institucionales y normativas en las que se implementa. Por ello, se proponen orientaciones específicas: (a) mantener la mediación docente como componente insustituible, (b) establecer marcos normativos auditables y transparentes, (c) fortalecer la formación docente en dimensiones técnicas, pedagógicas y éticas, (d) garantizar infraestructura tecnológica equitativa para evitar nuevas exclusiones, y (e) promover la participación activa de la comunidad educativa en el diseño y evaluación de estas herramientas.

Estas recomendaciones no solo buscan resolver problemáticas actuales, sino también anticiparse a los desafíos futuros, promoviendo un modelo de gobernanza tecnológica que priorice la justicia social, la inclusión educativa y el desarrollo integral de las personas en entornos digitales crecientemente automatizados.

Percepciones de estudiantes de IES sobre IA generativa y la incorporación de ODS 4

Carola Ubilla Briones

Universidad Metropolitana de Ciencias de la Educación

edith.ubilla2021@umce.cl

La inteligencia artificial (IA) generativa ha sido un tema recurrente en los últimos años, especialmente cuando se habla de los avances tecnológicos. La educación superior ha generado nuevas oportunidades y retos significativos, creando espacios clave para avanzar en los Objetivos de Desarrollo Sostenible (ODS) y alinearse con la Agenda 2030. Dado el contexto, este estudio tiene como objetivo analizar las percepciones que tienen los estudiantes de instituciones de educación superior (IES) sobre la IA generativa y la incorporación de ODS 4, que promueve una educación de calidad, equitativa e inclusiva para todos y a lo largo de la vida. Comprender estas percepciones es fundamental para analizar cómo la IA generativa puede contribuir al logro de una educación inclusiva y de calidad.

La metodología a trabajar en esta investigación es de tipo cualitativo, dado que se pretende responder a una pregunta de investigación, cuyo análisis de la información recabada de una toma de muestra permite llevar a cabo un proceso interpretativo, logrando desarrollar aportaciones a la investigación. La pregunta de investigación es ¿Cuál es la percepción de los estudiantes de Instituciones de Educación Superior respecto a IA generativa y la incorporación de ODS 4?

Esta investigación cobra relevancia al aportar evidencia sobre las percepciones estudiantiles respecto a la integración de la inteligencia artificial generativa en los procesos educativos vinculados al ODS 4. Comprender estas percepciones no solo

amplía el conocimiento sobre el impacto de las tecnologías emergentes en la educación superior, sino que también permite reflexionar críticamente sobre los desafíos éticos, pedagógicos y sociales asociados. Este análisis es esencial para orientar futuras estrategias de incorporación de la IA en entornos educativos de manera inclusiva, ética y orientada al desarrollo sostenible.

En este contexto la IA generativa tiene el potencial de ser una herramienta poderosa en la educación superior, contribuyendo a los ODS mediante la promoción de prácticas educativas más inclusivas y sostenibles. No obstante, el éxito de su integración depende de la percepción y competencia de los estudiantes. Dado todo lo que emerge en educación y con el fin de atender las demandas surgidas hasta hoy, se hace sumamente necesario que las IES junto a todos sus agentes, sean capaces de ampliar toda percepción respecto a la IA generativa y la integración de ODS 4.

Didáctica de la filosofía de la IA: Algoritmo de decisión y antropomorfismos en el debate sobre la inteligencia artificial

Jesús Queglas Caruz

Universidad de Valparaíso

christian.queglasc@alumnos.uv.cl

Esta ponencia propone un modelo didáctico para enseñar filosofía, y sus implicaciones e intereses sobre la inteligencia artificial (IA) en educación media, integrando análisis técnicos y reflexión crítica sobre la misma. La secuencia propuesta (validada en una experiencia piloto en el liceo industrial de Valparaíso) se estructura en dos sesiones con posibilidad de extensión y ajuste según las necesidades curriculares.

1- El estudio de algoritmos de decisión, ejemplificado mediante el trolley problem aplicado a vehículos autónomos, ¿debe la IA priorizar vidas según edad, número o simple azar?

2. El debate en torno a la definición de "vida artificial" a partir de la distinción entre "débil" y "fuerte", abordado por Antonio Dieguez en el análisis crítico desde la filosofía de la biología en torno al debate conceptual sobre "Vida", contrastada con las críticas de Mark Coeckelbergh a los antropomorfismos en IA, como lo es atribuir cualidades humanas a sistemas no conscientes.

El marco teórico permite contrastar interpretaciones en el marco del concepto de "vida" desde la filosofía de la biología permitido por el debate de Diéguez, con la problemática ética aplicada a la IA por Bostrom con estrategias de debate activo en los estudiantes en torno de pro desarrollar el pensamiento y reflexión crítica. La innovación central está en vincular los principios de interés de debates filosóficos clásicos como lo es el libre albedrío, el "ser", la definición de lo "vivo", a un contexto

tecnológico contemporáneo y de gran importancia para las nuevas generaciones que están en puestas de salir a continuar estudios superiores, o al mundo laboral directamente, donde el conocimiento de la IA es vital, considerando su inserción en la vida cotidiana social actual. El marco metodológico combina análisis de casos con discusión y debate guiado, evidenciando cómo los algoritmos materializan supuestos éticos mientras se deconstruyen narrativas antropomórficas, como lo es en ejemplo el debate sobre "La IA como ser sintiente".

Los resultados preliminares muestran que los estudiantes reflexionan críticamente en esta discusión filosófica contemporánea, tomando incluso partido sobre una de ambas posiciones en torno a la vida artificial, mostrando un interés genuino en discusiones actuales en torno a este fenómeno científico y tecnológico, los desafíos éticos que generan en nuestra actualidad, y un desarrollo de pensamiento crítico evidenciado en una participación reflexiva y activa en el aula de clases.

