



Entre algoritmos: género humano y género artificial

BETWEEN ALGORITHMS: HUMAN GENDER AND ARTIFICIAL ONE

Antonio Márquez Prieto

Catedrático de Derecho del Trabajo y de la Seguridad Social

Universidad de Málaga

amp@uma.es  0000-0002-2654-1433

Recibido: 24.01.2022 | Aceptado: 17.05.2022

RESUMEN

La inteligencia artificial, que llevaya décadas de desarrollo, supone cada vez más una realidad presente en las relaciones laborales, de forma sectorizada, pero evidenciándose al mismo tiempo en toda la realidad económica y social. Por ello aquí se aborda como fenómeno universal ante el que se intentan ofrecer medidas legales de regulación, las cuales, en la medida en que responden al binomio conducta-norma, resultan inadecuadas puesto que el desarrollo de la inteligencia artificial requiere un enfoque en el que entre en juego, como tercer elemento, la propia cultura humana de valores subyacentes. Desde este punto de vista se pretende llamar la atención sobre el hecho de que la inteligencia artificial constituye un fenómeno de enorme calado filosófico y, al mismo tiempo, supone una de las mejores oportunidades para el intento global de comprensión de lo que significa ser humano.

ABSTRACT

Artificial intelligence, which has been in development for decades, increasingly supposes a reality that is present in labor relations, in a sectorized way, but at the same time becoming evident in all economic and social reality. For this reason, here we try to address it as an universal phenomenon facing which we try to offer legal regulatory measures, which, to the extent that they respond to the behavior-norm binomial, are inadequate since the development of artificial intelligence requires a focus on the that comes into play, as a third element, the human culture of underlying values. From this point of view, it is intended to draw attention to the fact that artificial intelligence constitutes a phenomenon of enormous philosophical significance and, at the same time, represents one of the best opportunities for the global attempt to understand what it means to be human.

PALABRAS CLAVE

Algoritmos
Inteligencia artificial
Género humano
Ser humano
Digitalización
Socialidad
Humanidad

KEYWORDS

Algorithms
Artificial intelligence
Human gender
Human being
Digitization
Sociality
Humanity

SUMARIO

- I. INTRODUCCIÓN
- II. MEDIDAS NORMATIVAS DE PROTECCIÓN CONTRA ACTUACIONES ALGORÍTMICAS
 - A. EL ARTIFICIO DE LA INTERACCIÓN
 - B. MEDIDAS INSTITUCIONALES LIMITADAS
 - 1. PRÁCTICAS PROHIBIDAS
 - 2. INTELIGENCIA ARTIFICIAL CONSIDERADA DE ALTO RIESGO
 - 3. MEDIDAS NORMATIVAS INSUFICIENTES
- III. LO HUMANO Y LO ARTIFICIAL
 - A. HACIA UNA "SUPERINTELIGENCIA"
 - B. EL DESAFÍO DE "SER HUMANO" Y SEGUIR "SIENDO"
- IV. CONCLUSIÓN
- Bibliografía

I. INTRODUCCIÓN

"Género humano" es el título de un volumen de Isla Correyero¹ que incluye a su vez dos poemarios de la autora. Del primero de ellos, "Diario de una enfermera" podríamos destacar el abordaje de la crisis, la enfermedad, la ancianidad, la fragilidad del ser humano cuando se va acercando al final de su vida, en tanto que "Occidente" retrata a un mundo enfermo que adolece de falsedad, una barca que zozobra en la que aún sigue a flote una sociedad despreocupada. Es a este género humano al que en estas líneas se alude, en un tiempo en que no queremos acusar prematuramente de amenazante el desarrollo industrial de la inteligencia artificial. Pero, en la medida en que esa artificialidad no es como las anteriores, sino que es tejida con hebras de la razón, exige necesariamente un cuestionamiento humano, que en gran parte ya se viene planteando, y que nos interpela a todos y a cada uno. Intentando por ello aportar también mi visión personal, creo que es obligado considerar los escenarios posibles, las relaciones posibles, entre la racionalidad artificial –a la cual no podemos considerar sólo un autómatas o artificio mecánico, sino algo propiamente pensado para "pensar"– y la racionalidad humana. ¿Asistimos al surgimiento de otro género o especie? Y, en caso afirmativo, ¿qué tipo de repercusiones –positivas o negativas– puede ello tener para el género humano? En el marco de las relaciones laborales el asunto viene ya teniendo un importante tratamiento doctrinal, como respuesta al uso de algoritmos en la adopción de decisiones empresariales, en la organización de la propia empresa o en la configuración de plataformas cibernéticas para la prestación de servicios². Pero salta a la vista que esas posibles implicaciones en las relaciones laborales conducen directamente a las relaciones interpersonales de todo tipo, al tratarse éste de un fenómeno universal –en el que ciertamente el trabajo ha empezado a experimentar cambios singulares–.

1. Isla Correyero es el seudónimo con el que escribe su autora, Esperanza Correyero Rodríguez. La obra está publicada en 2014 por Inspirar-Expirar Ediciones.

2. Han sido abundantes e importantes los estudios que han abordado estos aspectos. Permítaseme destacar, por todos, Rivas Vallejo, M. P.: *Aplicación de la inteligencia artificial al trabajo*, Aranzadi, Cizur Menor, 2020.

Por eso, de la misma forma que en los dos poemarios del libro de Isla Correyero, se abordan a continuación, en primer lugar, las *medidas legales que atienden al ser humano* (lo que responde a la manera jurídica normal de actuar) y, en segundo lugar, la necesidad de abrir un debate serio sobre las repercusiones de este fenómeno para el *género humano* como tal (según un abordaje que va más allá de lo normativo).

II. MEDIDAS NORMATIVAS DE PROTECCIÓN CONTRA ACTUACIONES ALGORÍTMICAS

En este apartado se aborda el primer paso de los mencionados, correspondiente a la relación entre las posibles conductas de los sujetos y las consecuencias jurídicas previstas en las normas, según es habitual en un esquema jurídico de dos dimensiones (el aludido binomio conducta-norma), comenzando con la dimensión relativa a las conductas interactivas o recíprocas (en el marco de la digitalización) para a continuación abordar las correspondientes medidas legales. En un siguiente paso, como se ha dicho, se tratará de exponer cómo este enfoque bidimensional –claramente insuficiente, especialmente en el caso de la inteligencia artificial– puede ser complementado por una tercera dimensión.

A. El artificio de la interacción

Resulta interesante recurrir a modelos de comportamiento interactivo propios de la teoría de juegos, si, en vez de comentar casos reales concretos, se pretende, como en este caso, presentar de forma abstracta algunos patrones de acción recíproca entre diversos agentes –jugadores– sin que en la opción de cada uno exista influencia por parte de norma alguna, sino que la actuación o decisión de cada uno de los agentes dependa del abanico de opciones posibles de los otros sujetos intervinientes en la dinámica interactiva³. De entre múltiples modelos pueden ser mencionados dos concretos juegos: el dilema del prisionero y la caza del ciervo.

En cuanto al primero, el más famoso de los juegos de “suma no cero”⁴ su matriz de pagos, o de recompensas asignadas a los jugadores, puede cambiar según el autor que se consulte. Y, de hecho, la puntuación que marca las recompensas se corresponde con el número de años de cárcel a que pueden ser condenados dos sospechosos (prisioneros) de haber perpetrado juntos un delito. La posible condena variará dependiendo de que decidan cooperar entre sí (silenciando los dos la comisión del delito), con lo que conseguirían una pena bastante inferior, o no cooperar (confesando) viendo aumentada su condena dependiendo de que confiesen ambos o sólo uno.

3. En la medida en que responden a modelos, suponen una simplificación importante y también un enfoque significativamente reducido de lo que supondrían en realidad las posibles reacciones verdaderamente humanas (siempre de carácter muy complejo). Pero se pueden acercan, sin embargo, a la comprensión de actuaciones artificiales programadas (cuyos enunciados de partida igualmente pueden ser más simples o más complejos).

4. Vid. Martínez Coll, J. C.: *Bioeconomía*, Málaga, 1986, pp. 40-42 y Baird, D.G.; Gertner, R.H. y Picker, R.C.: *Game Theory and the Law*, Cambridge-USA y Londres, 1994, pp. 33 y ss.

Conducta de los jugadores	Número años de cárcel	Puntuación o recompensa
Ni A ni B confiesan (cooperan entre sí)	Cada uno es condenado a 1 año de cárcel	A=3, B=3
A silencia, B confiesa (B defrauda a A)	A es condenado a 10 años y B queda libre	A=1, B=4
Ambos confiesan (ambos se defraudan)	Cada uno es condenado a 6 años	A=2, B=2
A confiesa, B silencia (A defrauda a B)	A queda libre y B es condenado a 10 años	A=4, B=1

Así pues, la matriz de pagos quedaría así:

A/B	Coopera	No coopera
Coopera	3,3	1,4
No coopera	4,1	2,2

O bien, con puntuaciones alfabéticas, que permiten indicar la premisa del dilema del prisionero:

A/B	Coopera	No coopera
Coopera	a,a	c,b
No coopera	b,c	d,d

Donde **b>a>d>c**, premisa que debe cumplirse para que efectivamente se trate de un dilema del prisionero.

En este juego existe sólo una situación de equilibrio (equilibrio de Nash), en el lugar donde ambos jugadores manifiestan una actitud de no cooperar con el otro, a pesar de que, según puede verse, la recompensa conjunta sería superior si decidieran cooperar entre sí. Lo que sucede es que cada uno de los prisioneros trata de maximizar su beneficio, sin preocuparse por el resultado que el otro prisionero obtenga. Cada jugador, de hecho, ante la desconfianza de lo que el otro pueda hacer, está más incentivado a defraudar al otro, incluso tras haberse prometido entre sí la cooperación (éste es el dilema).

El segundo de los juegos aludidos, la caza del ciervo, debe su nombre a un relato de Rousseau⁵, que describe la situación en la que se encuentran dos cazadores, que deben elegir entre cazar juntos un ciervo o atrapar, por separado, cada uno una liebre. El ciervo constituye una pieza de mucho mayor interés que la liebre, pero exige coordinación entre ambos cazadores. En comparación con el dilema del prisionero, la premisa necesaria para que se trate de un juego de caza del ciervo es: **a>b≥d>c**, pudiendo representarse así su matriz de pagos:

5. Rousseau, J. J.: *Discurso sobre el origen y los fundamentos de la desigualdad entre los hombres*, Madrid, 1775 (traducción de 1982 de Ángel Pumarega).

A/B	Ciervo	Liebre
Ciervo	4,4	0,3
Liebre	3,0	2,2

Para algunos autores⁶ el juego de la caza del ciervo explica el fenómeno relacionado con la cooperación social –o la ausencia de ella– de forma más completa que el dilema del prisionero⁷ (aunque la prolongación en el tiempo de la relación entre los agentes puede hacer que la cooperación resulte explicable desde el dilema del prisionero, analizado de forma evolutiva). Sin embargo, interesa aquí resaltar una gran diferencia entre ambos juegos: el hecho de que el escenario de partida es en la caza del ciervo definible como neutro, mientras que en el dilema del prisionero es claramente de adversidad o conflicto; la decisión se toma por cada uno de los jugadores sin comunicación con el otro. Por ello constituye un escenario compatible con la existencia de un algoritmo en base al cual sea tomada la decisión por parte de uno de los dos agentes. En primer lugar, por la incomunicación (que puede también referirse a la ocultación de las razones, la expresión por sorpresa de la decisión) y en segundo lugar porque la decisión adoptada puede responder –o generar– un clima de verdadero conflicto. Ese algoritmo utilizado para decidir la opción de uno de los agentes puede tener una orientación predeterminada u otra. Así, puede citarse algún algoritmo que pretenda conducir el juego dentro de unos parámetros de cooperación y sostenibilidad de las relaciones interpersonales, como es, señaladamente “Tit for Tat”, diseñado para jugar al dilema del prisionero, que ocupa sólo 3 líneas en lenguaje Basic, y que consiste en: 1) cooperar la primera vez y 2) a partir de ahí hacer lo mismo que haga el adversario (más tarde se añadió una cierta capacidad de perdón)⁸. Frente a esa estrategia, que puede corresponder en general a una buena fe realista, puede citarse el algoritmo “MINIMAX”, que resulta muy combativo, ya que su estrategia consiste en elegir la opción más conveniente para uno mismo y peor para el adversario, sabiendo que el oponente elegirá también lo mejor para él y peor para uno mismo⁹. Pero, más allá de los algoritmos, es impactante el caso de redes neuronales como *AlphaZero*, presentada en 2017 por la empresa *DeepMind*, que, sin habersele facilitado formación sobre tácticas, sino aprendiendo ella sola desde cero, logró vencer a los ganadores mundiales en ajedrez, *shogi* y *go*¹⁰. Ante estos supuestos surge la pregunta sobre la necesaria reacción de

6. Cortázar, R., “Non-Redundant Groups, the Assurance Game and the Origins of Collective Action”, *Public Choice*, vol. 92, 1997, pp. 41-53; Janssen, M.: “On the principle of coordination”, *Economics and Philosophy*, vol. 17(2), 2001, pp. 221-234; Skyrms, B.: *La caza del ciervo y la evolución de la estructura social*, Barcelona, 2007.

7. Efectivamente es más completo el juego de la caza del ciervo en la medida en que presenta no una, sino dos situaciones de equilibrio posible: la opción por la caza de la liebre, por parte de ambos, y la decisión de ambos cazadores de mantenerse en la coordinación conjunta para la caza del ciervo, siendo esta segunda posibilidad la que mayor recompensa comporta.

8. Axelrod, R., “The evolution of cooperation”, *The Journal of Politics*, vol. 48, núm. 1, febrero-1986, pp. 234-236.

9. Beal, D., *The Nature of Minimax Search*, *Institute of Knowledge and Agent Technology*, Universidad de Maastricht. Disponible en https://project.dke.maastrichtuniversity.nl/games/files/phd/Beal_thesis.pdf (visitado el 23-1-2022).

10. <https://deepmind.com/blog/article/alphazero-shedding-new-light-grand-games-chess-shogi-and-go> (visitado el 23-1-2022).

los ordenamientos jurídicos, debiendo anticiparse que existe una gran dificultad al respecto por el poder de atracción y de convicción que juega a favor de la tecnología. Parece que se la acepta como una realidad imparable. La exagerada fascinación por el progreso de la misma se ve, sin embargo, frenada, cuando se descubre que en muchas ocasiones esta tecnología es utilizada a favor de determinados intereses –o causando importantes perjuicios y vulneraciones de derechos– por parte de instituciones públicas y grandes empresas a nivel nacional e internacional¹¹.

B. Medidas institucionales limitadas

Aparte de medidas legislativas que, a nivel nacional, aborden aspectos jurídico-laborales relacionados directamente con el fenómeno de la digitalización¹², existe en la Unión Europea una propuesta de Reglamento a modo de Ley de Inteligencia Artificial, lo cual se enmarca en toda una serie de iniciativas comunitarias relativas al desarrollo de la inteligencia artificial¹³ que, en años más recientes, ha dado lugar a la definición de una Estrategia Europea para la Inteligencia Artificial¹⁴, el establecimiento de un Plan Coordinado sobre la Inteligencia Artificial¹⁵ y la elaboración de un

11. Acerca del uso de algoritmos de impacto social significativo (con abuso de derechos de las personas) por parte de empresas y entidades públicas puede consultarse, p.e., el Observatorio de Algoritmos con Impacto Social (OASI) de Eticas Foundation: <https://eticasfoundation.org/es/weakening-of-democratic-practices-as-a-type-of-social-impact-of-algorithms-in-the-oasi-register/> (visitado el 23-1-2022).

12. En España, de forma pionera, y como resultado del Acuerdo de 10 de marzo entre Gobierno, CC.OO., UGT, CEOE y CEPYME, se han introducido, en relación al trabajo prestado a través de plataforma digital, dos modificaciones en el Estatuto de los Trabajadores por obra de la Ley 12/2021 de 28 de septiembre. En primer lugar, la introducción de una nueva letra d) en el artículo 64.4, que concede a los representantes de los trabajadores el derecho a ser informado por parte de la empresa “de los parámetros, reglas e instrucciones en los que se basan los algoritmos o sistemas de inteligencia artificial que afectan a la toma de decisiones que puedan incidir en las condiciones de trabajo, el acceso y mantenimiento del empleo, incluida la elaboración de perfiles”. En segundo lugar, la incorporación de una disposición adicional vigesimotercera, titulada “Presunción de laboralidad en el ámbito de las plataformas digitales de reparto”, con este texto: “Por aplicación de lo establecido en el artículo 8.1, se presume incluida en el ámbito de esta ley la actividad de las personas que presten servicios retribuidos consistentes en el reparto o distribución de cualquier producto de consumo o mercancía, por parte de empleadoras que ejercen las facultades empresariales de organización, dirección y control de forma directa, indirecta o implícita, mediante la gestión algorítmica del servicio o de las condiciones de trabajo, a través de una plataforma digital”. En el Preámbulo se reconoce la labor llevada a cabo por la Inspección de Trabajo en la indagación de las condiciones del trabajo prestado a través de plataformas digitales, además de mencionarse expresamente importantes Ss. del TS, como las de 26-2-1986, 20-1-2015 y 25-9-2020. Junto a lo cual sería injusto no mencionar la labor doctrinal de análisis del trabajo prestado en conexión con plataformas digitales. Es de mencionar –por todos los demás– Aragüez Valenzuela, L.: *Relación laboral «digitalizada»: colaboración y control en un contexto tecnológico*, Aranzadi, Cizur Menor, 2019.

13. En ocasiones anteriores ya la Comisión había instado a la necesidad de que la U.E. asumiera una posición de liderazgo a nivel internacional en el terreno de la inteligencia artificial. Así fue, p.e., en mayo de 2017, cuando la Comisión publicó la Revisión Intermedia de la aplicación para la Estrategia del Mercado Único Digital, COM/2017/0228 final, que puede consultarse en: <https://eur-lex.europa.eu/legal-content/ES/TXT/?uri=COM:2017:228:FIN> (visitado el 7-1-2022).

14. Inteligencia Artificial para Europa. Comunicación de la Comisión al Parlamento Europeo, al Consejo Europeo, al Consejo, al Comité Económico y Social, y al Comité de las Regiones. Bruselas, 25.4.2018 (COM) 2018 237 final, {SWD(2018) 137 final}.

15. Plan Coordinado sobre la Inteligencia Artificial. Comunicación de la Comisión al Parlamento Europeo, al Consejo Europeo, al Consejo, al Comité Económico y Social, y al Comité de las Regiones. Bruselas, 7.12.2018 (COM) 2018 795 final, que puede consultarse en: <https://eur-lex.europa.eu/legal-content/ES/TXT/HTML/?uri=CELEX:52018DC0795&from=DA> (visitado el 7-1-2022).

Libro Blanco¹⁶, además de otras medidas¹⁷. El mencionado Reglamento comunitario sobre Inteligencia Artificial, que de momento es sólo una propuesta¹⁸, clasifica las actuaciones posibles de la inteligencia artificial en función de los riesgos conectados con las mismas y, concretamente, establece tres niveles: 1) riesgo inaceptable, 2) alto riesgo y 3) riesgo bajo o mínimo. Independientemente de que la propuesta de Reglamento requeriría un comentario detallado –que no es posible llevar a cabo en estas páginas– es necesario señalar que entre el primer y el segundo nivel –en los que nos centramos aquí– la diferencia está en que las prácticas que conlleven un riesgo inaceptable quedan prohibidas, mientras que los sistemas de inteligencia artificial de alto riesgo están permitidos, si bien condicionados al cumplimiento de determinados requisitos. Analicemos, aunque sea de forma sucinta, la regulación contenida en la propuesta acerca, primeramente, de las prácticas prohibidas y, a continuación, la permisión de la llamada inteligencia artificial de alto riesgo; a lo que, en tercer lugar, se podrá añadir alguna consideración algo más general relativa al carácter limitado de estas medidas normativas.

1. *Prácticas prohibidas*

Comenzando por las conductas que según la propuesta de Reglamento quedarían prohibidas, se citan textualmente a continuación las cuatro prácticas descritas en el artículo 5, algunas de cuyas formulaciones, según ha valorado el Comité Económico y Social Europeo (CESE) “no están claras, con lo que determinadas prohibiciones pueden ser difíciles de interpretar y fáciles de eludir”¹⁹. Según la redacción del apartado 1 del mencionado artículo 5 de la propuesta de Reglamento “Estarán prohibidas las siguientes prácticas de inteligencia artificial:

16. Libro Blanco sobre la Inteligencia Artificial. Un enfoque europeo orientado a la excelencia y la confianza, Bruselas, 19.2.2020 COM(2020) 65 final, que puede consultarse en: https://ec.europa.eu/info/sites/default/files/commission-white-paper-artificial-intelligence-feb2020_es.pdf (visitado el 7-1-2022).

17. Se indica la necesidad de atender a Directrices Éticas que fomenten la confianza en la inteligencia artificial en la siguiente Comunicación de la Comisión (disponible en inglés): Building Trust in Human-Centric Artificial Intelligence. Communication from the Commission to the European Parliament, the Council, the European Economic and Social Committee and the Committee of the Regions, Bruselas, 8.4.2019, COM(2019) 168 final, que se puede consultar en: [file:///C:/Users/User/Downloads/1_en_act_part1_v8_DA596EE2-A7B1-2FF2-976724FBD96DE1F1_58496%20\(1\).pdf](file:///C:/Users/User/Downloads/1_en_act_part1_v8_DA596EE2-A7B1-2FF2-976724FBD96DE1F1_58496%20(1).pdf) (visitado el 7-1-2022). También el documento Directrices Éticas para una IA fiable, elaborado por el Grupo de Expertos de Alto Nivel sobre Inteligencia Artificial y editado por la Comisión Europea, Dirección General de Redes de Comunicación, Contenido y Tecnologías, Oficina de Publicaciones, 2019, <https://data.europa.eu/doi/10.2759/08746> (visitado el 12-1-2022).

18. Propuesta de Reglamento del Parlamento Europeo y del Consejo por el que se establecen normas armonizadas en materia de inteligencia artificial (Ley de inteligencia artificial) y se modifican determinados actos legislativos de la Unión, Bruselas, 21.4.2021 COM(2021) 206 final 2021/0106 (COD), que puede consultarse en: https://eur-lex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0008.02/DOC_1&format=PDF (visitado el 7-1-2022).

19. Dictamen del Comité Económico y Social Europeo a la Propuesta de Reglamento del Parlamento Europeo y del Consejo por el que se establecen normas armonizadas en materia de inteligencia artificial (Ley de Inteligencia Artificial) y se modifican determinados actos legislativos de la Unión, COM (2021) 206 final, ap. 4.1.

- a) La introducción en el mercado, la puesta en servicio o la utilización de un sistema de IA que se sirva de técnicas subliminales que trasciendan la conciencia de una persona para alterar de manera sustancial su comportamiento de un modo que provoque o sea probable que provoque perjuicios físicos o psicológicos a esa persona o a otra". Se trata de una indicación que, aunque se haya querido centrar en las técnicas subliminales, en realidad podría –y debería– referirse a la potencialidad de la inteligencia artificial para manipular –sea con dichas técnicas subliminales u otras por extensión–, pudiendo contemplarse los perjuicios también de manera más extensa, tanto en sentido objetivo como subjetivo. De hecho, el Dictamen mencionado del CESE (ap. 4.3) propone una modificación de la redacción de la parte que figura arriba en cursiva, quedando así: “un sistema de IA desplegado, destinado o utilizado para distorsionar sustancialmente el comportamiento de una persona de forma que pueda menoscabar los derechos fundamentales de esa persona, otra persona o de un grupo de personas, incluidas su salud y seguridad físicas o psicológicas, o la democracia y el Estado de Derecho”²⁰
- “b) La introducción en el mercado, la puesta en servicio o la utilización de un sistema de IA que aproveche alguna de las vulnerabilidades de un grupo específico de personas debido a su edad o discapacidad física o mental para alterar de manera sustancial el comportamiento de una persona que pertenezca a dicho grupo de un modo que provoque o sea probable que provoque perjuicios físicos o psicológicos a esa persona o a otra”. Se trata de una redacción un tanto ambigua que podría trasladar a los tribunales –como en otros casos– la labor de concretar más las conductas prohibidas (qué se deba entender por “vulnerabilidades de un grupo específico” o por alteración sustancial del comportamiento, qué “perjuicios físicos o psicológicos” son significativos)²¹.
- “c) La introducción en el mercado, la puesta en servicio o la utilización de sistemas de IA por parte de las autoridades públicas o en su representación con el fin de evaluar o clasificar la fiabilidad de personas físicas durante un período determinado de tiempo atendiendo a su conducta social o a características personales o de su personalidad conocidas o predichas, de forma que la clasificación social resultante provoque una o varias de las situaciones siguientes: i) un trato perjudicial o desfavorable hacia determinadas personas físicas o colectivos enteros en contextos sociales que no guarden relación con los contextos donde se generaron o recabaron los datos originalmente; ii) un trato perjudicial o desfavorable hacia determinadas personas físicas o colectivos enteros que es injustificado o desproporcionado con respecto a su

20. El Dictamen del CESE da esta explicación en su ap. 4.2: “Existen pruebas de técnicas subliminales que pueden, además de causar perjuicios físicos o psicológicos (la condición actualmente necesaria para que esta prohibición específica entre en vigor), derivar en otros efectos adversos personales, sociales o democráticos, como la alteración del comportamiento de voto, debido al entorno en el que se despliegan. Además, el elemento basado en la IA suele ser la decisión de a quién dirigirse con una técnica subliminal, en lugar de la propia técnica subliminal”.

21. El Dictamen del CESE (ap. 4.4) recomienda que, tras la palabra “provoque” se incluya “el menoscabo de los derechos fundamentales”, además de la alusión que se hace a los perjuicios físicos o psicológicos”.

comportamiento social o la gravedad de este”. La mención de la “clasificación social” y la “fiabilidad” de los ciudadanos señala, como en los otros apartados del artículo 5 de la propuesta de Reglamento, riesgos muy preocupantes; de ahí la necesidad de concretar más la redacción, que también en este caso adolece de cierta ambigüedad. El Dictamen del CESE (ap. 4.5) recomienda que esta prohibición se aplique no sólo a autoridades públicas, sino también a las semipúblicas y a organizaciones privadas, a fin de contemplar otros entornos, como el laboral, donde pueda llevarse a cabo una clasificación en base a la fiabilidad de las personas.

- “d) El uso de sistemas de identificación biométrica remota «en tiempo real» en espacios de acceso público con fines de aplicación de la ley, salvo y en la medida en que dicho uso sea estrictamente necesario para alcanzar uno o varios de los objetivos siguientes: i) la búsqueda selectiva de posibles víctimas concretas de un delito, incluidos menores desaparecidos; ii) la prevención de una amenaza específica, importante e inminente para la vida o la seguridad física de las personas físicas o de un atentado terrorista; iii) la detección, la localización, la identificación o el enjuiciamiento de la persona que ha cometido o se sospecha que ha cometido alguno de los delitos mencionados en el artículo 2, apartado 2, de la Decisión Marco 2002/584/JAI del Consejo²², para el que la normativa en vigor en el Estado miembro implicado imponga una pena o una medida de seguridad privativas de libertad cuya duración máxima sea al menos de tres años, según determine el Derecho de dicho Estado miembro”. En realidad, la posibilidad de uso legal de sistemas de identificación biométrica es más amplia de lo que parece permitir el mencionado precepto, puesto que no queda restringida a los supuestos mencionados como excepcionales²³. Así, por ejemplo, al prohibirse la identificación biométrica “en tiempo real” se entiende permitido el reconocimiento “posterior” y “próximo”; también queda permitido el reconocimiento que no tenga por finalidad identificar a una persona, lo que incluye la posibilidad de evaluar su comportamiento en base a las distintas características biométricas (que incluyen, p.e., micro-expresiones, temperatura, etc.); e incluso resulta también posible que podamos ser continuamente evaluados emocionalmente por parte de agentes públicos o privados, en lugares de trabajo o en otros ámbitos. Repárese en la clara petición de modificación contenida en el Dictamen del CESE (ap. 4.8): “Coincidiendo en líneas generales con la petición del Supervisor Europeo de Protección de Datos y el Comité Europeo de Protección de Datos, presentada el 21 de junio de 2021, de prohibir el uso de la IA para el reconocimiento automatizado de rasgos humanos en espacios de

22. Decisión Marco 2002/584/JAI del Consejo, de 13 de junio de 2002, relativa a la orden de detención europea y a los procedimientos de entrega entre Estados miembros (DO L 190 de 18 de julio de 2002, p. 1).

23. El apartado 2 del art. 5 de la propuesta de Reglamento exige además tener en cuenta otros criterios para la permisión del uso de sistemas de identificación biométrica, expresados en términos amplios. Por su parte el apartado 3 del mismo precepto exige contar con una autorización judicial, aunque también se menciona, alternativamente, una autorización de un órgano administrativo independiente, e incluso se permite que se lleve a cabo sin ningún tipo de autorización ante una “situación de urgencia debidamente justificada”.

acceso público, así como algunos otros usos de la IA que pueden dar lugar a una discriminación injusta, el CESE pide:

- la prohibición del uso de la IA para el reconocimiento biométrico automatizado, en los espacios de acceso público y privado (como los rostros, la forma de andar, la voz y otras señales biométricas), excepto con fines de autenticación en circunstancias específicas (por ejemplo, para facilitar el acceso a espacios sensibles desde el punto de vista de la seguridad);
- la prohibición del uso de la IA para el reconocimiento automatizado, en los espacios de acceso público y privado, de señales de comportamiento humano;
- la prohibición de los sistemas de IA que utilicen la biometría para clasificar a las personas en grupos en función de su etnia, género, orientación política o sexual o cualquier otro motivo de discriminación prohibido en el artículo 21 de la Carta de los Derechos Fundamentales de la Unión Europea;
- la prohibición de utilizar la IA para inferir las emociones, la conducta, la intención o los rasgos de una persona física, excepto en casos muy específicos, como determinados fines relacionados con la salud, en los que sea importante reconocer las emociones del paciente”.

2. *Inteligencia artificial considerada de alto riesgo*

El artículo 6 de la propuesta de Reglamento determina los sistemas de inteligencia artificial de alto riesgo en base dos posibles criterios. El primero, contenido en su apartado 1, consiste en el cumplimiento cumulativo de dos condiciones: a) que el sistema de inteligencia artificial esté “destinado a ser utilizado como componente de seguridad de uno de los productos contemplados en la legislación de armonización de la Unión que se indica en el anexo II, o [sea] en sí mismo uno de dichos productos”; y b) que el producto del que el sistema de inteligencia artificial sea componente de seguridad, o el propio sistema de IA como producto, deba someterse –conforme a la legislación de armonización de la UE indicada en el anexo II– “a una evaluación de la conformidad realizada por un organismo independiente para su introducción en el mercado o puesta en servicio”. Conforme al segundo criterio, mencionado en el apartado 2, “también se considerarán de alto riesgo los sistemas de IA que figuran en el anexo III”. Y, en relación a los sistemas determinados de acuerdo con ambos criterios, el artículo 8.1 dispone taxativamente que “cumplirán los requisitos que se definen en el presente capítulo”, lo cual, según una desconcertante deducción implícita, significa que el cumplimiento de los requisitos exigidos permite operar libremente a los sistemas considerados de alto riesgo²⁴. Ante ello adquiere todo el sentido la pregunta que formula el Dictamen del CESE acerca de si los requisitos impuestos reducen de forma efectiva los riesgos de los sistemas mencionados, siendo negativa su respuesta: “los riesgos de menoscabo de la salud, la seguridad y los derechos fundamentales no se reducen siempre necesariamente

24. Esos requisitos se refieren a 1) los datos y su gobernanza, 2) la documentación y el registro, 3) la transparencia y la comunicación de información a los usuarios, 4) la vigilancia humana, 5) la solidez, 6) la precisión y la seguridad.

a través del cumplimiento de los cinco requisitos establecidos para la IA de alto riesgo, en particular con respecto de los derechos fundamentales menos mencionados que podrían verse afectados por la IA, como el derecho a la dignidad humana, la presunción de inocencia, el derecho a unas condiciones de trabajo justas y equitativas, la libertad de asociación y de reunión y el derecho de huelga, entre otros” (ap. 4.14)²⁵. Pero además el CESE denuncia que la propuesta de Reglamento no ha incorporado cinco de los siete requisitos recogidos en las Directrices Éticas para una IA Fiable²⁶, recomendando que sean incorporados al texto normativo y ser tenidos en cuenta. Y, lo que es más grave, advierte (ap. 4.13) que el enfoque «basado en listas» para la IA de alto riesgo escogido en el anexo III puede llevar a la legitimación, normalización e integración de un gran número de prácticas de IA que todavía reciben muchas críticas y cuyos beneficios sociales son cuestionables o inexistentes”. En efecto, los sistemas de inteligencia artificial que son mencionados como de alto riesgo en el Anexo III de la propuesta de Reglamento (al que remite el art. 6.2) se refieren a los siguientes ámbitos:

1. Identificación biométrica y categorización de personas físicas («en tiempo real» o «en diferido»), es decir, que, con la condición de que sean cumplidos los requisitos, quedan permitidos cuando su uso se realice en términos más amplios que excedan de lo que prohíbe el artículo 5 (en espacios de acceso público, «en tiempo real» y con fines de aplicación de la ley).
2. Gestión y funcionamiento de infraestructuras esenciales (componentes de seguridad en la gestión y funcionamiento del tráfico rodado y el suministro de agua, gas, calefacción y electricidad); al igual que en el caso anterior –y en todos los posteriores–, queda clara la permisión de estos sistemas con la condición de que se cumplan los requisitos establecidos en el capítulo II de la propuesta de Reglamento.
3. Educación y formación profesional (se alude al uso de sistemas de IA que pueden llegar a ser muy invasivos tanto para determinar el acceso a la educación como para proceder a la evaluación de estudiantes)²⁷.

25. “El CESE recomienda que se incluya una disposición que contemple la situación en la que bien sea obvio o bien quede claro en la evaluación de conformidad previa que los seis requisitos no reducirán de forma suficiente el riesgo de menoscabo de la salud, la seguridad y los derechos humanos” (ap. 4.22 del Dictamen).

26. Comisión Europea, Dirección General de Redes de Comunicación, Contenido y Tecnologías, *Directrices éticas para una IA fiable*, op. cit. El Dictamen del CESE dice así en su apartado 4.10: El CESE acoge favorablemente la adaptación de estos requisitos a los elementos de las Directrices éticas para una IA fiable. Sin embargo, la Ley de Inteligencia Artificial no contempla específicamente cinco requisitos importantes relativos a la IA de alto riesgo incluidos en las Directrices, a saber: i) la acción humana, ii) la privacidad, iii) la diversidad, no discriminación y equidad, iv) la explicabilidad y v) el bienestar ambiental y social. El CESE considera que esto supone una oportunidad perdida, ya que muchos de los riesgos que plantea la IA están relacionados con la privacidad, el sesgo, la exclusión, la inexplicabilidad de los resultados de las decisiones de la IA, el menoscabo de la acción humana y el medio ambiente, aspectos todos ellos que afectan a nuestros derechos fundamentales”.

27. “Los sistemas de IA para determinar el acceso a la educación y evaluar a los estudiantes plantean una serie de riesgos de menoscabo de la salud, la seguridad y los derechos fundamentales de los estudiantes. Las herramientas de supervisión en línea, por ejemplo, destinadas supuestamente a señalar «comportamientos sospechosos» e «indicios de fraude» durante los exámenes en línea mediante el uso de todo tipo de indicadores biométricos y de

4. Empleo, gestión de los trabajadores y acceso al autoempleo (sistemas destinados a la toma de decisiones relativas a la contratación, selección, evaluación de candidatos, despidos, etc.). Ante la posibilidad de conculcación de derechos fundamentales “el CESE pide que se garantice la plena participación e información de los trabajadores y los interlocutores sociales en el proceso de toma de decisiones sobre el uso de la IA en el lugar de trabajo, así como sobre su desarrollo, adquisición y despliegue” (ap. 4.17 del Dictamen del CESE).
5. Acceso y disfrute de servicios públicos y privados esenciales y sus beneficios (sistemas de IA utilizables para evaluar el posible acceso a prestaciones y servicios de asistencia pública, mantenimiento o modificación de los mismos, evaluación de la solvencia de las personas e incluso decisión de prioridades para el envío de servicios de emergencia como bomberos y asistencia médica).
6. Asuntos relacionados con la aplicación de la ley (sistemas de IA para evaluar el riesgo que plantee una persona física que pretenda entrar en el territorio de un Estado miembro, la autenticidad de los documentos de viaje o el examen de solicitudes de asilo, visado y permisos de residencia).
7. Gestión de la migración, el asilo y el control fronterizo (entre los sistemas mencionados están polígrafos y herramientas similares, o bien instrumentos para detectar el estado emocional de las personas)²⁸.
8. Administración de justicia y procesos democráticos (a pesar de que se alude sólo a sistemas de IA para “ayudar” a los jueces en la investigación e interpretación de los hechos y de la ley, resulta algo bastante delicado; incluso se menciona esa “ayuda” para la aplicación de la ley a hechos concretos; piénsese también en los peligros de uso electoral de sistemas de IA)²⁹.

control del comportamiento, son verdaderamente invasivas y no existen pruebas científicas al respecto” (ap. 4.16 del Dictamen del CESE).

28. “La IA empleada por los servicios policiales y en la gestión de la migración, asilo y control de fronteras para realizar evaluaciones del riesgo individuales (en materia penal o de seguridad) supone un riesgo de menoscabo de la presunción de inocencia, los derechos de la defensa y el derecho de asilo de la Carta de Derechos Fundamentales de la Unión Europea. En general, los sistemas de IA suelen buscar simplemente correlaciones basadas en características de otros «casos». En estas situaciones, las sospechas no se basan en una sospecha real de un delito o una falta de la persona en cuestión, sino simplemente en características que esa persona comparte con delincuentes condenados (como la dirección, los ingresos, la nacionalidad, las deudas, el empleo, el comportamiento, tanto propio como de amigos y familiares, etc.)” (ap. 4.20 de Dictamen del CESE).

29. En efecto, como alega en su Dictamen el CESE (ap. 4.21), “el uso de la IA en los procesos democráticos y de administración de justicia es especialmente sensible, y debe abordarse con más matices y un mayor escrutinio que en la actualidad. La mera puesta a disposición de sistemas para su utilización por la autoridad judicial en la investigación y la interpretación de los hechos y de la ley, y en la aplicación de la ley a una serie concreta de hechos, no tiene en cuenta que las resoluciones judiciales van más allá de encontrar patrones en datos históricos (que es, en resumen, lo que hacen los sistemas de IA actuales). El texto también asume que estos tipos de IA solo servirán de apoyo en el ámbito judicial, con lo que la toma de decisiones totalmente automatizada se excluye del ámbito de aplicación. El CESE también lamenta que no se mencionen los sistemas o los usos de la IA en el marco de los procesos democráticos, como las elecciones”.

3. Medidas normativas insuficientes

La propuesta de Reglamento sobre inteligencia artificial regulación es claramente limitada y desenfocada. Es así en primer lugar, por lo inapropiado que resulta el enfoque de los riesgos. Por un lado parecer ser el único enfoque al que atiende el Reglamento, lo que lo convierte en un planteamiento parcial de entrada; la consideración de los riesgos es fragmentada, abordando cada sistema de IA de manera particularizada, ignorando que el fenómeno de la inteligencia artificial en su conjunto ofrece posibilidades crecientes de control dentro de una cada vez más extendida red de interacciones automatizadas, lejos de la búsqueda reflexiva del sentido correspondiente a cada una³⁰. Ciertamente se prohíben algunas posibles prácticas consideradas inaceptables –en los términos e incertidumbres expuestos–, pero, aparte de ello se establecen requisitos que –tal como se ha dicho– no están pensados para proteger de forma realmente suficiente al ser humano ante peligros considerados casi como externalidades, como efectos colaterales. Consideración fragmentada de riesgos y protección insuficiente contra los mismos que no puede redundar sino en una división y debilitamiento de la protección de la sociedad, de los derechos fundamentales.

Existe, en segundo lugar, más que un enfoque de los riesgos, una apuesta entusiasta por los beneficios de la inteligencia artificial, tal como afirma la Exposición de Motivos de la propuesta de Reglamento en sus primeras líneas, en las que se destacan, además de grandes oportunidades de mejora para la sociedad y de liderazgo de la Unión Europea a nivel tecnológico, los posibles beneficios económicos³¹. El planteamiento es claramente de mercado, como se dice claramente en el apartado 3.3 de la Exposición de Motivos: “La Comisión examinó diversas opciones para alcanzar el objetivo general de la propuesta, que es garantizar el buen funcionamiento del mercado único mediante la creación de las condiciones necesarias para el desarrollo y la utilización de una IA fiable en la Unión”. Por lo que se entiende que los requisitos cuyo cumplimiento resulta exigido para la utilización de sistemas de alto

30. Como si la IA no se refiriese a algo más ambicioso, totalizador, con capacidad para afectar al modelo social y cultural (a partir de aquí se sigue profundizando en esta llamada de atención).

31. En efecto, esto se expresa a partir de la tercera línea: “La inteligencia artificial (IA) es un conjunto de tecnologías de rápida evolución que puede generar un amplio abanico de beneficios económicos y sociales en todos los sectores y las actividades sociales. Mediante la mejora de la predicción, la optimización de las operaciones y de la asignación de los recursos y la personalización de la prestación de servicios, la inteligencia artificial puede facilitar la consecución de resultados positivos desde el punto de vista social y medioambiental, así como proporcionar ventajas competitivas esenciales a las empresas y la economía europea. Esto es especialmente necesario en sectores de gran impacto como el cambio climático, el medio ambiente y la salud, el sector público, las finanzas, la movilidad, los asuntos internos y la agricultura”. A continuación, y tras una referencia a posibles riesgos, a los que no parece darse de entrada la misma importancia que a los beneficios –porque no hay respecto de ellos la misma referencia ni detalle–, se insiste destaca el liderazgo tecnológico que la UE debería mantener, buscando un enfoque “equilibrado” respecto a valores y principios: “No obstante, los mismos elementos y técnicas que potencian los beneficios socioeconómicos de la IA también pueden dar lugar a nuevos riesgos o consecuencias negativas para personas concretas o la sociedad en su conjunto. En vista de la velocidad a la que cambia la tecnología y las dificultades que podrían surgir, la UE está decidida a buscar un enfoque equilibrado. Redunda en interés de la Unión preservar su liderazgo tecnológico y garantizar que los europeos puedan aprovechar nuevas tecnologías que se desarrollen y funcionen de acuerdo con los valores, los derechos fundamentales y los principios de la UE”.

riesgo sean contemplados ante todo como un coste económico a satisfacer, especificando además que no será muy alto³² y que su aplicación se llevará a cabo con cierta flexibilidad³³.

Y, en tercer lugar, se han de considerar medidas normativas insuficientes porque la propuesta de Reglamento de inteligencia artificial no constituye una regulación de la misma como tal. En lugar de ello, como ya se ha dicho, se someten a prohibición algunas prácticas consideradas como “riesgo inadmisibles” y se someten a cumplimiento de requisitos –para poder ser permitidos– los sistemas de IA considerados de “alto riesgo”. Es decir, en vez de una regulación de la IA, se pretende una regulación para la IA, puesto que las normas propuestas se presentan como marco legitimador de la llamada “IA fiable”³⁴. Es verdad que la Directrices Éticas para una IA Fiable habían mencionado que “es preciso encontrar un equilibrio entre lo que se puede hacer con la IA y lo que se debería hacer con ella, además de prestar la debida atención a lo que no se debería hacer con esta tecnología”³⁵. Pues bien, partiendo de que la propuesta de Reglamento atendiera dicha indicación e incluso asumiera directamente la regulación de la transformación digital como fenómeno –cosa que no se hace– habría que decir también que dicho planteamiento es erróneo: las expresiones “lo que se puede hacer” y “lo que se debería hacer”, ¿de qué sujetos se predicán? Ya que está claro el predicado, pero, hablando de la transformación digital, ¿podemos mantener –y seguir manteniendo durante mucho tiempo– que los sujetos seremos únicamente humanos? Porque, en la medida en que surjan subjetividades artificiales, ¿se podrá realmente limitar su actuación en términos de “lo que se debería hacer”, o sólo será realísticamente posible poner un límite a tiempo acerca de “lo que se puede hacer”

32. Así se expresa la Exposición de Motivos (ap. 3.3): Las empresas o las autoridades públicas que desarrollen o utilicen aplicaciones de IA que entrañen un riesgo elevado para la seguridad o los derechos fundamentales de los ciudadanos tendrían que cumplir requisitos y obligaciones específicos. El cumplimiento de estos requisitos tendría un coste aproximado de entre 6 000 y 7 000 EUR para el suministro de un sistema de IA de alto riesgo promedio con un valor de en torno a 170 000 EUR para 2025. Por su parte, los usuarios de la IA tendrían que cubrir, cuando correspondiese, el coste anual del tiempo dedicado a garantizar la vigilancia humana, en función del uso concreto que hagan. Se calcula que este coste ascendería a entre 5 000 y 8 000 EUR al año, aproximadamente. Además, los proveedores de IA de alto riesgo podrían tener que abonar unos costes de verificación adicionales de entre 3 000 y 7 500 EUR. Las empresas o las autoridades públicas que desarrollen o utilicen aplicaciones de IA no consideradas de alto riesgo únicamente tendrían que cumplir unas obligaciones mínimas de información. No obstante, podrían decidir unirse a otras y, juntas, adoptar un código de conducta para cumplir los requisitos adecuados y garantizar la fiabilidad de sus sistemas de IA. En tal caso, los costes podrían ser, como máximo, tan elevados como los de los sistemas de IA de alto riesgo, aunque lo más probable es que fueran inferiores”.

33. “Las soluciones técnicas exactas que se necesitarán para lograr el cumplimiento de tales requisitos podrían proceder de normas u otras especificaciones técnicas, o bien desarrollarse con arreglo a los conocimientos científicos o de ingeniería generales, a discreción del proveedor del sistema de IA de que se trate. Esta flexibilidad reviste una importancia especial, ya que permite a los proveedores de sistemas de IA decidir cómo quieren cumplir los requisitos, teniendo en cuenta el estado de la técnica y los avances tecnológicos y científicos en este campo” (ap. 5.2.3 de la Exp. de Mot.).

34. Así se dice en la primera línea del documento llamado Directrices Éticas para una IA Fiable de 2019 (*cit.*): “El objetivo de las presentes directrices es promover una IA fiable”. Y en el punto 11 de la Introducción se añade: “(...) identificamos que la IA fiable constituye nuestra aspiración fundamental, puesto que los seres humanos y las comunidades solamente podrán confiar en el desarrollo tecnológico y en sus aplicaciones si contamos con un marco claro y detallado para garantizar su fiabilidad”.

35. Así se afirma en la pág. 42 de dicho Documento, elaborado en 2019 por el Grupo Independiente de Expertos de Alto Nivel sobre Inteligencia Artificial (*op. cit.*).

para no progresar a partir de ahí? Ésta es una pregunta cuya respuesta puede ser muy probablemente adivinada: no es fácil un consenso para poner límites a la inteligencia artificial mientras no haya una conciencia real de hasta dónde puede aquélla llegar si no existen límites. Y, si se suman los beneficios sociales que la inteligencia artificial parece ofrecer, pero, sobre todo, los altísimos intereses en juego a favor de algunos, se entiende que la regulación del fenómeno de transformación digital no sea una prioridad, sino que dicha transformación sea presentada como un objetivo a lograr con rapidez, proponiéndolo como la solución frente a los errores y efectos perniciosos del ser humano en el medio ambiente³⁶; una “IA fiable” entendida ya no sólo como realidad en la que el ser humano puede confiar, sino como algo que es más fiable que el ser humano, capaz, por ejemplo, de evitar accidentes de tráfico, por ofrecer un mejor tiempo de reacción que los humanos y un mejor respeto de las normas de seguridad vial³⁷. Un objetivo tecnológico indiscutido ante el cual los seres humanos quedan condicionados a adaptarse continuamente y estar disponibles a adquirir las cualificaciones adecuadas que ello exige³⁸. Dentro de esta transformación digital, cuya regulación tal vez no se quiere acometer por entender que supondría un freno a algo considerado altamente benéfico, se ha expresado la propuesta de reconocer la categoría jurídica de personalidad electrónica³⁹. Este solo ejemplo, en caso de estar sometido a regulación, contribuiría a tener un marco normativo más adecuado. O bien, en caso de no regularse –en el sentido de limitar o prohibir–, operarían de facto dichos actores de forma libre, sometidos sólo a las posibilidades de la tecnología; con un importante dato añadido: que, a medida que avanza la transformación digital, habrá también que recurrir a agentes de la propia IA para encargarse del control de ella misma. Es decir, que la digitalización no sólo puede incrementar las posibilidades de control para seres humanos por parte de otros humanos, sino que los propios sistemas cibernéticos necesitarán ser controlados –independientemente

36. “(...) la transformación digital y la IA fiable ofrecen un potencial enorme para reducir el impacto que ejerce el ser humano sobre el medio ambiente y posibilitar un uso eficiente y eficaz de la energía y los recursos naturales (*Directrices Éticas...*, ob. cit., p. 42, ap. 123).

37. *Id.*, p. 43, ap. 124.

38. “Los cambios tecnológicos, económicos y medioambientales obligan a la sociedad a ser más proactiva. Gobiernos, líderes industriales, instituciones educativas y sindicatos se enfrentan a la responsabilidad de ayudar a los ciudadanos a realizar la transición a la nueva era digital, garantizando que posean las cualificaciones que requerirán los puestos de trabajo del futuro. Las tecnologías de la IA fiable podrían ayudar a predecir con mayor exactitud qué puestos de trabajo y qué profesiones se verán afectados de un modo fundamental por la tecnología, qué nuevas funciones se crearán y qué competencias se necesitarán. Esto podría ayudar a los gobiernos, sindicatos y a la industria a planificar la (re)cualificación de los trabajadores. También podría ofrecer una vía para el reciclaje profesional a aquellos ciudadanos que temen que sus cualificaciones puedan quedar obsoletas” (*Id.*, p. 46, ap. 127).

39. Resolución del Parlamento Europeo, de 16 de febrero de 2017, con recomendaciones destinadas a la Comisión sobre normas de Derecho civil sobre robótica (2015/2103(INL)), P8_TA(2017)0051, disponible en https://www.europarl.europa.eu/doceo/document/TA-8-2017-0051_ES.pdf (visitado el 19-1-2022). Apartado 59: “Pide a la Comisión que, cuando realice una evaluación de impacto de su futuro instrumento legislativo, explore, analice y considere las implicaciones de todas las posibles soluciones jurídicas, tales como: f) crear a largo plazo una personalidad jurídica específica para los robots, de forma que como mínimo los robots autónomos más complejos puedan ser considerados personas electrónicas responsables de reparar los daños que puedan causar, y posiblemente aplicar la personalidad electrónica a aquellos supuestos en los que los robots tomen decisiones autónomas inteligentes o interactúen con terceros de forma independiente”.

de su tamaño–, para lo cual serán necesarios controladores igualmente artificiales, surgiendo una vez más la pregunta: *quis custodiet ipsos custodes?*⁴⁰. Por lo que, ante el silencio de la propuesta de Reglamento al respecto, tiene toda la razón el Comité Económico y Social cuando “recomienda encarecidamente que la Ley de Inteligencia Artificial contemple que determinadas decisiones sigan correspondiendo a las personas, especialmente en ámbitos donde estas decisiones tengan un elemento moral e implicaciones jurídicas o repercusión social, como en los ámbitos judicial y policial, los servicios sociales, la sanidad, la vivienda, los servicios financieros, las relaciones laborales y la enseñanza”⁴¹.

Interesa aludir al final de este apartado a la cuestión sobre el objeto de la transformación digital, es decir, qué es lo que se está transformando, o bien, en qué terreno se está produciendo esa transformación digital. Claramente el terreno social – que corresponde al tejido humano interactivo– es el que se siente interpelado, terreno en el que lo digital adquiere cada vez más importancia, ante lo cual se entienden necesarias tanto la reflexión como la protección jurídicas. Ante la ausencia de regulación del fenómeno de la transformación digital se necesitaría una adecuada protección, no sólo del ser humano considerado individualmente, sino, de forma muy especial, del propio terreno humano. El hábitat interactivo humano necesitaría ser protegido como un medio ambiente natural ante la proliferación de especies artificiales, puesto que no asistimos al desarrollo de un tejido interactivo humano, sino al de un sistema de conexiones, seguido de conexiones de sistemas, susceptibles de funcionar unos como componentes de otros, cada vez a mayor escala, donde puede perder escala el ser humano, como elemento “organizado” insertable en una red organizativa; puesto que no puede actuar sin más como “componente”, sino que, conforme a su dignidad de persona y a su naturaleza social, necesita actuar conforme a interacciones que respondan a vínculos con sentido. No se trata sólo del efecto hueco que produce el trato cada vez más frecuente con sujetos artificiales, ni de no tener quizá la seguridad de si el interlocutor es humano o máquina⁴²; sino que el hábito de operar mediante

40. Debido claramente a las características de la IA los inspectores del cumplimiento de la ley por parte de ésta han de ser guardianes artificiales, es decir, programas para controlar el funcionamiento de esa propia IA, ya sea que luego puedan ser supervisados de alguna manera por seres humanos o directamente monitorizados también artificialmente. De ahí la pregunta sobre quién vigila a los vigilantes. Es lo que plantea Cotino Hueso, L.: “Ética en el diseño para el desarrollo de una inteligencia artificial, robótica y *big data* confiables y su utilidad desde el Derecho”, *Revista Catalana de Dret Públic*, núm. 58, 2019, p. 43.

41. Dictamen del CESE, ap. 3.7. Por su parte el art. 22.1 del Reglamento General de Protección de Datos de la UE establece que “Todo interesado tendrá derecho a no ser objeto de una decisión basada únicamente en el tratamiento automatizado, incluida la elaboración de perfiles, que produzca efectos jurídicos en él o le afecte significativamente de modo similar” (Reglamento UE 2016/679 del Parlamento Europeo y del Consejo de 27 de abril de 2016, relativo a la protección de las personas físicas en lo que respecta al tratamiento de datos personales y a la libre circulación de estos datos y por el que se deroga la Directiva 95/46/CE).

42. “Los seres humanos siempre deberían saber si están interactuando directamente con otro ser humano o con una máquina. Los responsables de garantizarlo son los profesionales de la IA. Dichos profesionales deberían, por tanto, asegurar que las personas sean conocedoras de que están interactuando con un sistema de IA o puedan preguntar y dar su aprobación al respecto (por ejemplo, mediante la publicación de cláusulas de exención de responsabilidad claras y transparentes). Obsérvese que existen casos limítrofes que complican la cuestión (por ejemplo, una voz emitida por un ser humano pero filtrada a través de un sistema de IA). Es preciso tener presente que la confusión entre humanos y máquinas puede tener múltiples consecuencias, como el establecimiento de vínculos, la influencia

intercambios y conexiones automatizadas pueden ir erosionando la interacción natural humana, reduciéndola a un modo culturalmente empobrecido y limitado a lo tecnológico, encerrando lo humano en el nivel “micro” por el protagonismo creciente en el ámbito social (“macro”) de la pseudo-cultura de procesos digitalizados. Es por tanto el modo humano de interactuar lo que debe ser protegido, como ecosistema que responda a la dignidad humana, al sentido de humanidad⁴³.

III. LO HUMANO Y LO ARTIFICIAL

Llegados ya al momento de dar el segundo paso aludido en la Introducción, es decir, ir más allá de lo estrictamente normativo, se debe señalar que el abordaje insuficiente por parte de la institucionalidad normativa podría resultar un problema tal vez similar al de la inteligencia artificial: en el caso de que, aun intentando producir una superinteligencia, se detectaran en ésta limitaciones de importancia. Se quiere apuntar a la idea de un abordaje en ambos casos reductivo; es decir, de reducción de lo humano.

A. Hacia una “superinteligencia”

En el contexto de la llamada “cuarta revolución industrial”, relacionada con la combinación de la ingente captación y utilización de datos, del uso para ello de algoritmos y de la inteligencia artificial, en definitiva, como resultado cada vez más actualizado del intento decidido de calcular, programar, automatizar y controlar, es interesante un estudio publicado recientemente que pone en evidencia el carácter nihilista de este proyecto, según ya se había anticipado⁴⁴. Sus autores afirman que “lo primero que debemos tener en cuenta para intentar comprender el significado de la inteligencia artificial (IA) es la propia noción de “Inteligencia”. Definir la inteligencia humana ya es difícil y complejo, sin embargo, cuando se habla de IA se da por sentado lo que se entiende por inteligente. La Inteligencia asociada a las máquinas sólo es un tipo de

o la reducción del valor de la persona⁷³. El desarrollo de robots humanoides debería ser objeto, por tanto, de una evaluación pormenorizada desde el punto de vista ético” (*Directrices Éticas*, ob. cit., p. 45, ap. 131).

43. Sobre la correlación entre dignidad humana y humanidad, vid., entre otros, Häberle, P.: *La libertad fundamental en el Estado Constitucional*, Universidad Católica de Perú, Lima, 1997 (cap. 2), también editado en Comares, Granada, 2003. Vid. también Monereo Pérez, J. L.: *La dignidad del trabajador: dignidad de la persona en el sistema de relaciones laborales*, Laborum, Murcia, 2019.

44. Valle Jiménez, D. y García Ramírez, D.: “Algoritmos, big data e inteligencia artificial: ¿un nihilismo anunciado?”, *Cuadernos Salmantinos de Filosofía*, núm. 48, 2021, pp. 75-103. Los autores aluden a tres principales filósofos (Nietzsche, Jünger y Heidegger) que se habían referido a este programa nihilista, del que la inteligencia artificial es un claro ejemplo. Así, por ejemplo, aluden a Heidegger, en la medida en que éste ve en la voluntad de poder y el superhombre “la representación de los absolutos que subyacen en el mundo tecnológico y en los que se manifestaría el afán de dominio y manipulación del pensar calculante propio de la técnica (...), es un hombre que debe hacerse cargo del dominio incondicionado de la tierra y del mundo en su conjunto (...). En este contexto es importante subrayar que la crítica de Heidegger a la llamada revolución tecnológica no tiene que ver con una demonización de los avances y progresos. (...) Más bien la crítica se dirige a la forma de interpretar la realidad bajo el dominio de la técnica, donde el cálculo, la planificación y programación automática, no solo adquiere un rango de privilegio, sino que se concibe como la única forma de pensar, actuar e interpretar el mundo” (pp. 88-89).

inteligencia definida por su capacidad de cálculo y procesamiento de información. Es inteligente porque puede trabajar con una cantidad de información que excede las capacidades humanas. Esta inteligencia es valorada por su velocidad y eficiencia. En consecuencia, detrás de los discursos y promesas de la IA lo que se prioriza es una idea de inteligencia limitada pero afín a los intereses de quienes la promueven y de aquella racionalidad basada en la cuantificación, cálculo y eficiencia⁴⁵. Ciertamente esa eficiencia parece verse muy acentuada por el uso masivo de datos y por la incorporación de algoritmos. “La aplicación de algoritmos e IA a todos los aspectos de la vida buscan hacer un mundo más controlable y eficiente, además de disminuir la arbitrariedad humana. Sin embargo, esta automatización o algoritmización de la realidad lo que impone es un nuevo tipo de arbitrariedad: la algorítmica. Que a diferencia de la humana se vuelve incontestable y aparentemente incontrovertible⁴⁶. Resultando, pues, que, mientras los sistemas inteligentes pueden ser considerados como la panacea para mejorar drásticamente enormes problemas sociales, pueden en realidad ser considerados inteligentes “de una forma limitada y reducida: sólo sirven para resolver los problemas para los que han sido programados, sus procesos son finitos y sus resultados discretos. Las promesas de una IA o de una superinteligencia que solucionarán los males de la humanidad hacen parte de un discurso ideológico⁴⁷.”

Debemos, pues, tener en cuenta que, en la misma medida en que la inteligencia artificial supone una reducción de la inteligencia humana –además de establecer interacciones que no responden al modo natural de la interacción humana–, también el abordaje jurídico puede ser reduccionista: limitado a la sola regulación de conductas por parte de la norma (sólo dos elementos o dimensiones: interacción de conductas y regulación normativa). Se hace necesario contar con la socialidad humana (como tercera dimensión) para una consideración jurídica completa, especialmente en el caso de la inteligencia artificial. Porque los contenidos de justicia de donde extraer inspiración para la regulación de conductas tienen que ser buscados en la comprensión de lo que es humano. De la misma forma que es necesaria la verdadera comprensión de la inteligencia humana (ya que la inteligencia artificial trata de emularla), para regular de forma sistemática el posible efecto de la inteligencia artificial, a fin de no vernos reducidos los seres humanos a un género cada vez más artificioso o artificializado, es necesario indagar en la dignidad de lo humano.

B. El desafío de “ser humano” y seguir “siendo”

Como ya se ha dicho, el modo interactivo que difunde, como dinámica suya, la inteligencia artificial, puede erosionar el modo natural humano de relaciones, pero, sobre todo, conociendo la concomitancia entre el modelo tecnológico y el nihilismo,

45. *Id.*, p. 92.

46. *Id.*, p. 94.

47. *Id.*, pp. 95-96.

tenemos ahí una pista⁴⁸ para identificar el paradigma de género humano que debe ser protegido: el de la dignidad o valor humanos, para nutrir un vacío nihilista de justicia. Un vacío que, entre otras cosas, consiste en una “inteligencia” que, aunque tiende a colonizar el espacio de la inteligencia humana, carece en realidad de verdadero pensamiento o reflexión⁴⁹. Pero son más las carencias, como plantea María Zambrano, quien diferencia entre lo que llama “etapa humana” y la del “superhombre”⁵⁰, pudiéndose destacar en la primera dos importantes elementos: la conciencia (conciencia y pensamiento, como lugar de residencia del ser humano⁵¹) y la libertad; mientras que en la segunda sobresale la voluntad de desbordar los límites humanos de la conciencia (anclada en el bien y en el mal) y una vivencia extrema libertad, desarraigada no sólo de lo sagrado sino también de la alteridad. Zambrano escribe que Nietzsche, en su propuesta de superhombre, no acogió la idea de libertad del idealismo que trascendía lo humano –por tanto, sin lucha, sin tragedia–, sino que “vivió y apuró la libertad, lejos de la impasibilidad, sólo humanamente. Hubiera sido normal que hubiese aceptado la libertad trágica de un Kierkegaard o de un Unamuno. Pero la tragedia de la libertad, o la libertad vivida trágicamente, requiere un *alguien* a quien ofrecérsela. Toda tragedia es un sacrificio; su protagonista necesita un alguien a quien ofrecerle su agonía. Y en esta soledad radical, frente a la cual la soledad de

48. En este sentido es como se puede mencionar esta parte final del artículo: sólo a modo de simples pistas iniciales que podrían ser seguidas para empezar a recorrer este camino.

49. Precisamente se ha destacado, como ejemplo del carácter nihilista de la IA, que en ella “sobresale la ausencia de meditación sobre las consecuencias humanas y sociales que conlleva esta hibris tecnológica, la medida desmesurada de los números y el cálculo. Nietzsche es quien plantea esta idea cuando habla de la historia de los próximos doscientos años, “Toda nuestra cultura europea se mueve desde hace tiempo bajo la tortura de una tensión que crece decenio en decenio como abocada a una catástrofe: inquieta, violenta, precipitada, como un río que quiere acabar, que no reflexiona ya, que tiene miedo de reflexionar sobre sí mismo”. Precisamente, en esa inquietud, en esa violencia el nihilismo (incompleto) campa a sus anchas, sustituyendo el lugar de Dios por el valor absoluto de la tecnología. Como un discurso e ideología dominante que subrepticamente ocupa el lugar de la deidad, ocultando justamente ese movimiento, en las promesas de salvación, perfección, certeza y eternidad que podrían alcanzar gracias al desarrollo de la tecnología y la inteligencia artificial, como una de las promesas más elocuentes de algunas corrientes del transhumanismo” (*id.* 98).

50. Zambrano Alarcón, M., *El hombre y lo divino*, Alianza Editorial, Madrid, 2020. Se abordan estos argumentos especialmente en sendos capítulos de este libro: “El delirio del superhombre” (pp. 184-206) y “La última aparición de lo sagrado: la nada” (pp. 207-223). Deben expresarse las disculpas correspondientes por aludir aquí de forma excesivamente escueta a nociones que María Zambrano desarrolla de forma más extensa –y que además requerirían mucho mayor comentario, dada su profundidad–. En realidad en los mencionados capítulos se alude a cuatro etapas, pues antes de lo que la autora llama “etapa humana” hay que mencionar la del “rey-mendigo” (que para ella permanece como sustrato último, haciendo las veces de entrañas del ser humano); y, tras el “superhombre”, propone, como siguiente etapa transitable, la de “la nada”, en la que el “superhombre” no ha llegado a adentrarse. Así escribe: “El superhombre, rectificación del proyecto en que el hombre de Occidente decidió su ser, no se hundió lo bastante en el oscuro seno de la vida primaria, de lo sagrado. Lo divino –descubierto por el pensamiento– le atrajo fascinándole. Quiso ser divino como lo «divino» que ya estaba pensado, descubierto... Había en realidad sacrificado al hombre ante lo divino, abismándose en él. Todo lo humano había sido destruido implacablemente, menos el tiempo. Y, más allá del tiempo, le hubiera esperado una última resistencia: la nada” (pp. 205-206).

51. “Si alguien quiere sorprender al hombre como tal, en el punto exacto de su humanización, habrá de dirigirse a este período del pensamiento europeo que va desde Descartes hasta el fin del idealismo, es decir, hasta el momento en que frente a Hegel se levantan las réplicas de Marx, Kierkegaard y más tarde Nietzsche. Los tres, desde ángulos tan diferentes, alzan frente a «lo humano» lo no-humano: lo no-humano económico que deshumaniza –enajena– la vida pretendidamente humana y su libertad, logro último del vivir desde y según la conciencia. La precariedad e indigencia del ser mismo del hombre –otra vez el mendigo frente a Dios–, mas cargado por el fardo de una culpa en Kierkegaard, y la continuación «lógica» del idealismo ávido de deificación en el superhombre de Nietzsche” (*id.*, p. 197).

la conciencia cartesiana se borra, Nietzsche apuró la tragedia de la libertad humana. No tuvo a quien dedicar su sacrificio. Y comenzó la destrucción”⁵². Y, ante todo ello, surgen las cuestiones sobre la forma de encontrar nociones de lo humano que vayan más allá de la conciencia y la libertad, que se correspondieron claramente con la Modernidad (como modo de *ser* humano) y que pueden verse amenazadas, entre otros posibles embates, por el proyecto de una inteligencia artificial reductiva de lo humano, pudiendo ayudar, como posible esbozo de respuesta, estas dos consideraciones finales:

La primera viene ofrecida por María Zambrano, quien, a continuación de su tratamiento sobre el superhombre, plantea la etapa de “la nada”, del no ser, de la total soledad incluso (en referencia entre otros a Sartre, una vez rota la relación con lo sagrado). La oportunidad a la que se refiere la filósofa es una forma de ser compartida, en alteridad: “Y la cuestión que se presenta es si el hombre puede, en verdad, estar entera, absolutamente, solo. A su lado va «el otro», el otro sombra de sí mismo, como Unamuno alumbrara en esa su tragedia –una de las raras tragedias logradas– *El otro*. ¿Quién es «el otro»? El hermano invisible, o perdido, aquel que me haría ser de veras si compartiera su existir conmigo; si nos integráramos en un ser único, a quien ya no le podría ser dirigida la pregunta terrible: «¿Qué has hecho de tu hermano?»”⁵³.

La segunda consideración, más claramente en el ámbito jurídico, es también una llamada, desde la socialidad humana, a una conducta de reciprocidad en correspondencia coherente con el valor, la dignidad, el *ser* de lo humano, que palpita en la cultura jurídica –no sólo de los juristas, sino del pueblo– independientemente de sus atributos –conciencia, libertad, autonomía...–, a modo de fundamento o referencia de humanidad –aunque no exista una explicación más completa–. Esa consideración nos la ofrece, por ejemplo, la Sentencia de 14 de octubre de 2004 del Tribunal de Justicia de la Unión Europea (caso *Omega*)⁵⁴, que considera conforme al Derecho Comunitario la prohibición por parte de las autoridades alemanas de un juego de entretenimiento que incluye simulaciones de acciones homicidas entre jugadores. Lo destacable de este pronunciamiento es que son los valores principios de la cultura jurídica⁵⁵ los que sustentan una reacción contra actuaciones –aun siendo simuladas o virtuales– que fomentan una actitud contraria al respeto que merece la dignidad humana: “El órgano jurisdiccional remitente expone que la dignidad humana es un principio constitucional que puede verse vulnerado mediante el trato degradante a un adversario, que no

52. Se refiere María Zambrano a la destrucción de la “etapa humana”. Como si Nietzsche reprochara a Sócrates, que había padecido para hacer nacer al hombre, el tipo de criatura a que había dado lugar. Sigue diciendo: “El motivo de la disidencia tiene relación con lo enunciado en la acusación contra Sócrates. Era la vieja cuestión habida entre la filosofía y la piedad, en la que se jugaba no sólo la relación con los dioses del Olimpo griego, sino algo más terrible: el desarraigo del hombre de lo sagrado” (*id.*, p. 202).

53. *Id.*, p. 216.

54. STJUE 14-10-2004, Asunto C-36/02, *Omega Spielhallen- und Automatenaufstellungs-GmbH* contra *Oberbürgermeisterin der Bundesstadt Bonn*.

55. “En el litigio principal, las autoridades competentes estimaron que la actividad objeto de la orden de prohibición amenazaba el orden público debido a que, según la concepción predominante en la opinión pública, la explotación comercial de juegos de entretenimiento que impliquen la simulación de acciones homicidas menoscaba un valor fundamental consagrado por la constitución nacional, como es la dignidad humana”.

concorre en el presente caso, o creando o fomentando en el jugador una actitud que niegue el derecho fundamental de toda persona a la consideración y al respeto, como sucede en el presente caso, mediante la representación de actos ficticios de violencia con fines lúdicos. Un valor constitucional superior como es la dignidad humana no puede quedar excluido en el marco de un juego de entretenimiento”.

IV. CONCLUSIÓN

Con el fenómeno de la inteligencia artificial tal vez no estemos asistiendo al surgimiento de un nuevo género o especie artificial, al margen del género humano, aunque sí existen motivos que requieren la protección jurídica de éste (contra la vulneración de derechos que conlleven las prácticas de IA o la erosión o empobrecimiento del modo humano de interacción).

La propuesta de Reglamento de Inteligencia Artificial de la UE resulta claramente insuficiente: se consideran de forma muy fragmentada los riesgos posibles y no se adoptan en general medidas verdaderamente eficaces para frenarlos. No existe, además, un intento de regular de forma completa y directa la inteligencia artificial como tal, pues dicha regulación parece más bien responder a la elaboración de un marco jurídico que legitime el uso de la IA en el mercado y refuerce el liderazgo europeo a nivel tecnológico y económico.

Pero, además, el abordaje jurídico basado sólo en la regulación de conductas (en base a sólo dos dimensiones: comportamiento interactivo e institucionalidad normativa), que resulta incompleto en términos generales, es especialmente limitado en el caso de la IA, que hace más necesario ir más allá de lo normativo. Precisamente porque tanto la limitación de dicho enfoque como el efecto posible de la IA pueden redundar en el mismo problema: la reducción del sentido humano, especialmente a nivel interactivo y colectivo. Es preciso recurrir a la socialidad humana (como tercera dimensión) para profundizar en los contenidos de justicia que, inspirados en la dignidad de lo humano, puedan orientar la regulación de las conductas. Al final de este trabajo se alude a algunas nociones clásicas del ser humano, como la conciencia y la libertad que, ciertamente, pueden estar sometidas a crisis culturales, tecnológicas o de otra índole que conduzcan a una verdadera pérdida o anulación. A pesar de lo cual el género humano parece tener oportunidad de resurgir redescubriendo desde la nada eslabones perdidos de alteridad y reciprocidad y dando sentido a la cultura social de reacción a favor de la dignidad humana.

Bibliografía

- Aragüez Valenzuela, L.: *Relación laboral «digitalizada»: colaboración y control en un contexto tecnológico*, Aranzadi, Cizur Menor, 2019.
- Axelrod, R.: “The evolution of cooperation”, *The Journal of Politics*, vol. 48, núm. 1, febrero-1986, pp. 234-236.
- Baird, D.G.; Gertner, R.H. and Picker, R.C.: *Game Theory and the Law*, Cambridge-USA y Londres, 1994.

- Beal, D.: The Nature of Minimax Search, Institute of Knowledge and Agent Technology, Universidad de Maastricht. Disponible en https://project.dke.maastrichtuniversity.nl/games/files/phd/Beal_thesis.pdf.
- Comisión Europea: Libro Blanco sobre la Inteligencia Artificial. Un enfoque europeo orientado a la excelencia y la confianza, Bruselas, 19.2.2020 COM(2020) 65 final, que puede consultarse en: https://ec.europa.eu/info/sites/default/files/commission-white-paper-artificial-intelligence-feb2020_es.pdf.
- Correyero Rodríguez, E.: *Género Humano*, Inspirar-Expirar Ediciones, 2014.
- Cortázar, R.: "Non-Redundant Groups, the Assurance Game and the Origins of Collective Action", *Public Choice*, vol. 92, 1997, pp. 41-53.
- Cotino Hueso, L.: "Ética en el diseño para el desarrollo de una inteligencia artificial, robótica y *big data* confiables y su utilidad desde el Derecho", *Revista Catalana de Dret Públic*, núm. 58, 2019, pp. 29-48.
- Grupo de Expertos de Alto Nivel sobre Inteligencia Artificial, *Directrices Éticas para una IA fiable*, Comisión Europea, Dirección General de Redes de Comunicación, Contenido y Tecnologías, Oficina de Publicaciones, 2019, <https://data.europa.eu/doi/10.2759/08746>.
- Häberle, P.: *La libertad fundamental en el Estado Constitucional*, Universidad Católica de Perú, Lima, 1997 (cap. 2), también editado en Comares, Granada, 2003.
- Janssen, M.: "On the principle of coordination", *Economics and Philosophy*, vol. 17(2), 2001, pp. 221-234.
- Martínez Coll, J. C.: *Bioeconomía*, Málaga, 1986.
- Monereo Pérez, J. L.: *La dignidad del trabajador: dignidad de la persona en el sistema de relaciones laborales*, Laborum, Murcia, 2019.
- Rivas Vallejo, M. P.: *Aplicación de la inteligencia artificial al trabajo*, Aranzadi, Cizur Menor, 2020.
- Rousseau, J. J.: *Discurso sobre el origen y los fundamentos de la desigualdad entre los hombres*, Madrid, 1775 (traducción de 1982 de Ángel Pumarega).
- Skyrms, B.: *La caza del ciervo y la evolución de la estructura social*, Barcelona, 2007.
- Valle Jiménez, D. y García Ramírez, D.: "Algoritmos, *big data* e inteligencia artificial: ¿un nihilismo anunciado?", *Cuadernos Salmantinos de Filosofía*, núm. 48, 2021, pp. 75-103.
- Zambrano Alarcón, M.: *El hombre y lo divino*, Alianza Editorial, Madrid, 2020.