



Venditio fumi: Autorregulación Empresarial e Inteligencia Artificial.

Venditio fumi: Business Self-regulation and Artificial Intelligence.

DOI: 10.32870/sincronia.axxvi.n81.12a22

Jonathan Piedra Alegría

Universidad Nacional de Costa Rica. (COSTA RICA)

CE: jonathan.piedra.alegria@una.cr / ID ORCID: 0000-0003-4532-4415

Esta obra está bajo una [Licencia Creative Commons Atribución-NoComercial 4.0 Internacional](https://creativecommons.org/licenses/by-nc/4.0/)

Recibido: 24/09/2021

Revisado: 05/10/2021

Aprobado: 03/11/2021

RESUMEN

En la primera parte de este artículo se analizará la utilización reduccionista e instrumental de la ética en temas relacionados con la regulación privada de la Inteligencia Artificial. Principalmente en lo relacionado con la autorregulación ética empresarial. Se parte de la hipótesis de acuerdo con cual, las propuestas éticas utilizadas por las grandes corporaciones de la Inteligencia Artificial son parte de una estrategia que busca solamente crear una apariencia de responsabilidad empresarial y social, que tiene como propósito, confundir a los legisladores y a los encargados de políticas nacionales, así como a los consumidores finales, con tal de mantener sus prácticas empresariales y ganancias intactas. Se proponen como ejemplos paradigmáticos los casos de Google y Amazon con tal de mostrar esta situación, para concluir como los desafíos éticos y legales de la implementación de la IA, son reinterpretados bajo una "gestión de riesgo éticos", en la cual se trivializa y se vacía de contenido el discurso ético.

Palabras clave: Inteligencia artificial. Autorregulación. Ética. Riesgo ético.

ABSTRACT

In the first part of this article, we will analyze the reductionist and instrumental use of ethics in issues related to the private regulation of Artificial Intelligence. Mainly to ethical business



self-regulation. It is based on the hypothesis according to which, the ethical proposals used by the large Artificial Intelligence corporations are part of a strategy that seeks only to create an appearance of business and social responsibility, whose purpose is to confuse legislators, as well as well as end consumers, to keep their business practices and profits intact. We will use the cases of Google and Amazon as paradigmatic examples to show this situation. This part concludes how companies reinterpret the ethical and legal challenges of the implementation of AI, towards an "ethical risk management" that trivializes and empties the ethical discourse of content.

Keywords: Artificial Intelligence. self-regulation. Ethics. Ethical risk.

Introducción

El avance y desarrollo en la Inteligencia Artificial (IA) transformará nuestra vida de manera incuestionable. Asistentes virtuales *-reconocimiento de lenguaje natural-*, chatbots *-procesamiento de lenguaje escrito-* o sistemas que nos recomiendan una película *-algoritmos de recomendación-* ya se han vuelto parte de nuestras vidas. Cambios profundos están ocurriendo en áreas como la salud y la sanidad, la educación o en la movilidad ciudadana. Su impacto potencial es tan grande que algunos académicos consideran que prácticamente todos los temas de la vida cotidiana podrían ser mejorados por la IA (Morte & Monasterio, 2017), llegando inclusive a resolver (en algún momento) los problemas mundiales más apremiantes. Esto hace creer a ciertos especialistas que la IA es la tecnología más floreciente de esta década, denominando así el tiempo actual como la "Era de la IA" (Mondal 2020). Esta capacidad transformativa, así como el ritmo creciente con que se está implementando en las actividades humanas, no dejan duda respecto a que los avances en IA tendrán consecuencias sociales de gran alcance. Los posibles beneficios económicos de su implementación han hecho que tanto el sector privado como el público se interesen cada vez más en su investigación y puesta en funcionamiento.

Un estudio de *Price Waterhouse Coopers* citado por el BID (2020) señala que para el año 2030 la adopción de la IA podría hacer crecer el PIB mundial cerca de un 14 por ciento (equivalente a 16,5 trillones de dólares americanos), debido, fundamentalmente, al aumento de la productividad



(i.e robots industriales, vehículos autónomos o inteligencia asistida). Por mencionar un ejemplo, algunos datos sobre Europa muestran que "Del 2010 al 2014 las ventas de robots aumentaron una media del 17 por ciento por año y es la industria electrónica la que está generando el mayor número de ofertas de empleo" (*Normas de Derecho Civil sobre Robótica*).

Dado este contexto, así como la amplia gama de usos potenciales de la IA (en particular el aprendizaje automático -*machine learning*-, la computación cognitiva o uso de los macrodatos -*Big data*) la conciencia sobre una serie de problemas sobre su uso ha aumentado. Temas como la privacidad, el uso ético de los datos, la responsabilidad, autonomía, la sostenibilidad ambiental o la transparencia ahora no solo se discuten en espacios especializados sino también en los principales medios de comunicación. Un sentimiento de desconfianza y preocupación nace no solo de los riesgos inherentes de este tipo de la tecnología, sino además de la incuestionable acumulación de poder (económico, tecnológico y social) que las empresas líderes en IA están logrando.

Esto plantea dudas sobre el control y su posible regulación. Cuestiones como: ¿Cuál sería el punto de partida para la regulación? O ¿Qué se debería tener en cuenta para esto? Son preguntas que las empresas, así como los Estados y comunidades supranacionales están considerando actualmente. Aspectos éticos y políticos relacionados con la aplicación de ciertas tecnologías de IA son comunes los debates sobre este tema.

A saber, el tema de la autonomía (ie. responsabilidad imputable a los sistemas autónomos y semiautónomos, así como su inclusión en ámbitos médico-asistenciales o su interacción con los seres humanos -robots sociales-) es uno de los focos de discusión en los documentos que buscan su regulación. Otros temas como el reemplazo del trabajo humano por robots con IA, el manejo de los datos o la creación de algoritmos sesgados o discriminatorios tampoco son dejados de lado, siendo estos los principales riesgos y temores que genera la expansión y utilización de la IA. El Parlamento Europeo expresó que:

Considerando que, ahora que la humanidad se encuentra a las puertas de una era en la que robots, bots, androides y otras formas de inteligencia artificial cada vez más sofisticadas parecen dispuestas a desencadenar una nueva revolución industrial —que probablemente



afecte a todos los estratos de la sociedad—, resulta de vital importancia que el legislador pondere las consecuencias jurídicas y éticas, sin obstaculizar con ello la innovación. (EU Commission, 2017).

Como cabría esperar, estos aspectos plantean cuestiones importantes. Si tenemos presente que los sistemas y las tecnologías de IA son desarrolladas fundamentalmente por empresas privadas (Kai-Fu Lee, 2020), y que, además, no necesariamente basan su operación en algún sustento ético; ¿Como definir el *approach* más idóneo para que se "no obstaculice la innovación"? o ¿Cuál sería el acercamiento más adecuado para ponderar estas consecuencias?

¿Regular o no regular?

Las respuestas a estas preguntas se pueden ubicar, grosso modo, en dos grandes grupos: Podemos encontrar en primer lugar (I) a los "progresistas económicos" que consideran que la IA nos puede llevar a una bonanza económica y social, en el cual esta tecnología se convertiría en una guía para la riqueza y el trabajo (Kaplan, 2016). Así pues, lo más adecuado sería fomentar la investigación, el desarrollo e implementación de la IA, manteniendo las regulaciones estatales al mínimo necesario. Este vendría a ser, concretamente, el enfoque estadounidense sobre la IA. Este se centra en la innovación y en el desarrollo de tecnologías (que no podrían darse si hay regulación) más que en la defensa de las personas o la protección de los Derechos Humanos.

Esta perspectiva no implica el establecimiento de una regulación particular, ni un posicionamiento claro con respecto al tipo de marco normativo necesario y menos aún con postura ética deseada, aunque si destaca la necesidad de ciertos mínimos que deben ser respetados por las empresas por la vía de la autorregulación.

Luego encontramos aquellos que plantean que debería tenerse cautela con la puesta en marcha acríticamente de esta tecnología. En el mismo Estados Unidos, el *Berkman Klein del Center for Internet & Society*, considera que debe existir una deliberación exhaustiva y a largo plazo de los impactos de la utilización de tecnologías relacionadas con la IA, y no únicamente una visión



cortoplacista como la que comúnmente se presentan en los discursos sobre este tipo de tecnologías. Usualmente estas narrativas se dirigen hacia (II) la creación de un marco legal que permita regular de algún modo las complejas interacciones de una IA con su entorno.

Esta posición contempla la necesidad de su regulación desde el punto de vista de las políticas públicas, así como desde la creación de normas, leyes y reglamentos claros que definan el espacio de acción de las tecnologías de IA. Una crítica importante a esta perspectiva proviene de las grandes compañías privadas que dominan el campo de la IA. Muchas de ellas consideran que la regulación es inadecuada (carece del "*on-the-ground approach*"), limitando la innovación debido a una falta de conocimiento de la realidad tecnológica y de la dinámica del mercado actual.

Un punto medio.

Precisamente esto ha provocado que en muchos casos se opte por un punto de vista intermedio: Las *Guías éticas* sobre la IA. Según algunos estudios, solo en el año 2019 podían encontrarse 89 documentos de este tipo (Jobin et al, 2019). Llegando casi duplicarse (más de 160 *Guías éticas* sobre IA) su número en el año 2020 (Algorithm Watch, 2020). Estas *Guías* son resoluciones no-vinculantes de organismos internacionales o dictámenes emitidos por grupos interdisciplinarios de expertos (tanto públicos como privados). Todas ellas tienen la intención de servir como pautas que permitan crear lineamientos básicos sobre el tema, así como situar algunos de los dilemas de la IA en un marco de análisis crítico, al mismo tiempo que presentan algunas alternativas éticas frente a estos desafíos.

No obstante, tampoco escapan a las críticas. Los inconvenientes más importantes de estas *Guías* son varios: Para comenzar, son bastante imprecisas en cuanto a los detalles técnicos y conceptuales. También es normal que carezcan de mecanismos de implementación o supervisión y, por el contrario, presenten una gran cantidad de descripciones ambiguas, confusas o claramente problemáticas. Por lo general se basan en una lista muy limitada de principios, que comúnmente se pueden reducir a cinco de ellos: transparencia, justicia y equidad, un uso no-dañino, responsabilidad e integridad (Jobin et al, 2019). Principios que pocas veces son operativizados y que usualmente



carecen de elementos para su aplicación real. Tampoco es extraño encontrar referencias vagas a lugares comunes como el respeto al derecho civil o a las normas del ordenamiento jurídico correspondiente.

Estos documentos presentan pautas no-obligatorias que pueden servir como un paso inicial para la posterior regulación. Desde esta perspectiva, es común encontrar propuestas que buscan la *autorregulación ética-normativa* de las empresas (i.e Google, Amazon, Microsoft) frente a los retos sociales y políticos de la IA. En términos generales es complicado discernir claramente como este tipo de propuestas aportan a la discusión sobre la regulación, en vista de que no existe ninguna pauta obligatoria, postura ética específica o un sistema de valores que pueda orientar el debate¹. Por lo que son necesarios más estudios sobre este tema con tal de dimensionarlas adecuadamente.

Aun así, los gobiernos de diferentes países, organismos supranacionales y las mismas empresas están impulsando propuestas en este sentido. Tomemos como caso de referencia la propuesta de la OCDE a este respecto. Este organismo creó una lista de cinco principios (y cinco recomendaciones)² basados en valores éticos para una administración responsable de la IA (OECD, 2019). Lo que ahí se presenta, *grosso modo*, son premisas poco desarrolladas en cuanto a los sistemas de IA. Repitiendo *locus communes* tales como el respeto a las leyes nacionales, los DDHH o a la diversidad humana. Aspectos que desde luego son de gran importancia y sin duda alguna, deberían tenerse en cuenta, pero no tienen ni un solo elemento de política pública que intente abordar realmente los desafíos ético-jurídicos que presentan estas tecnologías.

Lo mismo sucede con el tema de los DDHH. A pesar de que se menciona -con mayor o menor énfasis- el respeto a estos derechos (bajo la premisa del cumplimiento de los instrumentos internacionales), en casi todos los documentos, lo cierto es que siempre aparecen sin ninguna referencia directa a alguno de ellos -más allá del respecto a la dignidad o la no-discriminación- y en las raras ocasiones que sí lo hacen siempre están totalmente carentes de contenido normativo. Por

¹ Esto no quiere decir que la Guías éticas no sean importantes. Sin duda alguna son un paso inicial para una futura regulación. Nuestra crítica se dirige hacia ciertos elementos de estos documentos, como la autorregulación empresarial, así como un acercamiento *light* en ciertos temas de la IA.

² Hoy 42 países han adoptado estos principios.



el contrario, se omiten groseramente los derechos sociales, laborales o los relacionados con el ambiente. Manteniendo incuestionable la estructura de poder o las desigualdades estructurales propias del sistema económico en el cual lucran. Es decir, este tipo de propuestas crean un ambiguo espacio de interpretación ético-normativo que podría ser utilizado por aquellos que quieren utilizar de las *Guías éticas* como una forma de evadir la regulación o simplemente como un *ethical washing* (Floridi, 2019).

La narrativa ético-regulatoria en la Inteligencia Artificial

Para nadie es un secreto que los principales avances en el ámbito de la IA son producto de las empresas privadas. Más específicamente, son dos grandes grupos de empresas quienes manejan el desarrollo y la innovación de la IA: El grupo conocido como GAFAM: Alphabet, Amazon, Facebook, Apple y Microsoft, (Prisma European Network, 2021) así como el conjunto de empresas chinas incluidas bajo el acrónimo de BATX: Baidu, Alibaba, Tencent y Xiaomi (Prisma European Network, 2021) controlan el mercado de la IA. Esta concentración del poder económico y tecnológico no solo crea una mayor desigualdad social (Kai-Fu Lee, 2020), sino que además hace más notorias las transgresiones a la ley, a los derechos fundamentales -o incluso en temas de DDHH³- por parte de estas empresas.

Como respuesta a estos problemas, muchas de estas empresas han creado "Comités Éticos" o "Guías éticas de comportamientos"⁴ sobre la IA con la finalidad de autorregularse en su prácticas empresariales y tecnológicas. Sin embargo, esta práctica difiere mucho de lo que se esperaría.

AI for Good?

³ El caso DeepMind de Google en Reino Unido, Cambridge Analytica y Facebook o el software de reconocimiento facial discriminatorio en Amazon o Microsoft, son una muestra directamente relacionada con la IA. Pero no se reducen a esto. Por ejemplo, la manera en Amazon abusa y explota laboralmente a sus empleados (Adler & Schneider, 2020) o las prácticas anti-sindicales de Google (Scheiber & Wakabayashi, 2019) son bien conocidas

⁴ Vid. "Perspectives on Issues in AI Governance" de Google (2019), "Operationalizing responsible AI" de Microsoft o la propuesta "Partnership on AI" (PAI) (2018).



En los primeros meses del año 2019, Google anunció de una forma llamativa la creación de un *Comité Asesor a nivel global sobre temas de IA* (Advanced Technology External Advisory Council). Con esto, la empresa californiana se unía a un grupo selecto de empresas "comprometidas" (como Amazon o Microsoft) con el "bienestar social" que también habían creado grupos asesores para fines parecidos. Este comité estaba formado por especialistas en tecnología, ética digital y personas con experiencia en políticas públicas. El objetivo principal de este grupo era proporcionar recomendaciones de carácter ético (en general) para el uso de IA en Google, así como para los investigadores que trabajan en áreas como el software de reconocimiento facial, manejo de datos o temas relacionados con la privacidad. Sin embargo, una semana después de su conformación el *Comité* fue disuelto, debido a los insistentes reclamos de sus miembros para fuera excluida una de sus integrantes, en vista de que era bien conocida su postura en contra de la población LGBTI y anti-migrantes (Joshi, 2019).

Este comité surgió, en alguna medida, como una respuesta al *Proyecto Maven* -2018-, en el cual Google le vendió al Ejército de los Estados Unidos un software de IA para los drones no tripulados. Cuando los empleados de Google se enteraron, muchos de ellos se pronunciaron en contra del proyecto, llegando incluso a renunciar a la empresa. Esta oposición finalmente llevó a Google a no renovar con el Ejército este contrato⁵.

A raíz de esto, el CEO de Google emitió en ese año -2018- una especie de lista de principios éticos⁶ en los que se mencionaba que la empresa no diseñaría o desplegaría IA para tecnologías que «puedan causar un daño general», para armas, para vigilancia o para tecnologías "whose purpose

⁵ No parece que esto fuera del todo cierto. En vista de un enigmático Tweet (cerca de un año después) de Donald Trump (presidente de Estados Unidos en ese momento) en el cual mencionaba, como luego de una reunión con Pinchai, tenía claro que Google estaba comprometido con el ejército de USA. "Just met with @SundarPichai, President of @Google, who is obviously doing quite well. He stated strongly that he is totally committed to the U.S. Military, not the Chinese Military..." — Donald J. Trump (@realDonaldTrump) March 27, 2019

⁶ Microsoft también emitió una serie de principios ("Operationalizing responsible AI") aún más generales que Google (<https://www.microsoft.com/en-us/ai/our-approach?activetab=pivot1:primaryr5>) en ese mismo año. De igual manera en el 2021, publicó una entrada en su blog oficial en la que se hace referencia a "The building blocks of Microsoft's responsible AI program" pero si precisar ningún aspecto en particular. (<https://blogs.microsoft.com/on-the-issues/2021/01/19/microsoft-responsible-ai-program/>).



contravenes widely accepted principles of international law and human rights". (Pichai, 2018). Lugares comunes, que como ya hemos indicado, se repiten a lo largo de la mayoría de las *Guías Éticas sobre la IA*. A pesar de esta publicación, Google, en ese mismo año, reconoció haber probado una versión de su motor de búsqueda diseñado únicamente para cumplir con la censura que China aplica a sus habitantes (Simonite, 2020) en abierta oposición con todos con los principios que ellos indicaron.

Frente a esto, debemos decir que este tipo de prácticas no son necesariamente contradictorias. Esto se debe, ante todo, a que los *Principios* de Google son un texto bastante superficial que no aborda ningún problema concreto. En este documento no se menciona ni un solo mecanismo o alguna forma para rendirle cuentas al Estado, inversionistas o público-usuario. Solo aparecen unos cuantos principios genéricos - privacidad, la responsabilidad, un uso no-dañino o la seguridad, entre otros- que no llegan a ningún lado. Lo mismo sucede con *Perspectives on Issues in AI Governance* (Google, 2019) un texto que se publicó en enero del año 2019 cuya finalidad consiste en informar a la opinión pública y a los gobiernos -al estilo de un *White Paper*- sobre cuál sería la «ruta» idónea respecto a la IA. Es necesario decir, que estas *Perspectivas* están planteadas de una manera más formal que los *Principios* del 2018. Lo que hace que su propuesta sea considerablemente más seria y amplia que las pautas previas, aunque ni así su finalidad.

Su contenido principal se resume al desarrollo de cinco áreas en las que los gobiernos (no necesariamente las empresas) pueden (no deben) trabajar en conjunto con la sociedad civil y los profesionales de la IA para proporcionar una guía importante sobre el desarrollo y uso responsable de la IA. Si bien el documento señala algunas consideraciones ético-políticas, no podría decirse que haya un desarrollo de ningún problema concreto. Por el contrario, una parte importante de este *White paper* es un proto panfleto que busca mostrar la importancia de la IA para nuestras vidas, así como lo inoportuno de una regulación estricta. El texto se toma la molestia, incluso, de utilizar toda una sección para señalar lo inadecuado de un enfoque regulatorio como el de la UE⁷, ya que una regulación como esa, provocaría una grave "inseguridad jurídica" a las empresas, creando una

⁷ Por ejemplo, The European Commission Staff Working Document on liability for emerging digital technologies.



«responsabilidad fuera de lugar», "sobrecargando" a los fabricantes de sistemas de IA, lo que crearía sin lugar a duda "un efecto paralizador en la innovación y la competencia".

Claramente este tipo de propuestas éticas lo único que buscan es evitar la verdadera regulación, que se muestra como inoportuna, innecesaria, ignorante de las dinámicas tecnológicas o simplemente equivocada. Para evitar esto, la solución más adecuada es la autorregulación. Prácticas de autoevaluación o enfoques regulatorios propios son la clave para dinamizar el sector y beneficiar a todos. "A self-regulatory or co-regulatory set of international governance norms that could be applied flexibly and adaptively would enable policy safeguards while preserving the space for continued beneficial innovation" (Google, 2019). No obstante, todo indica que el propósito final de estas posiciones es evadir un compromiso real con la regulación de las actividades comerciales de estas empresas⁸.

Un tema insignificante

Una situación similar ocurre con Amazon. El del 22 de mayo del año 2018, diferentes medios de comunicación reportaron que *Amazon Web Services*, estaba vendiendo su tecnología de reconocimiento facial, llamada *Rekognition* a una serie de departamentos policiales⁹ en Estados Unidos (Oregon y Florida¹⁰).

Rekognition puede rastrear a personas en tiempo real en la transmisión de un video de lugares públicos o de multitudes. Además, es capaz de buscar en bases de datos con millones de imágenes e identificar 100 rostros (personas) en una sola imagen (Ng, 2018). Amazon promovía el uso de esta tecnología para utilizarla en las cámaras corporales que llevan los policías (*American Civil Liberties Union Foundations of California*, 2018), lo cual transformaría totalmente el concepto

⁸ Lo mismo podría decirse de iniciativas como "Partnership on AI" (2018) que agrupa a empresas líderes en la IA como Amazon, Apple, Baidu, Facebook, Google, IBM o Intel, en la cual se busca resaltar más la afiliación a esta asociación que brindar alguna solución concreta o una propuesta de regulación.

⁹ No obstante, se viene utilizando por distintas agencias de Gobierno de los Estados Unidos desde el año 2016. Por ejemplo, ya era utilizada por la U.S. Immigration and Customs Enforcement (ICE).

¹⁰ A raíz de esto muchos otros Estados también estaban interesados. Como fue el caso de California y Arizona.



de intrusión a la vida privada, vigilancia policial y responsabilidad del Estado¹¹. Por si esto fuera poco, este software tiene un sesgo discriminatorio contra las personas no-caucásicas (i.e población afrodescendiente o latina).

En una prueba realizada por la *American Civil Liberties Union* (UCLU) utilizando el servicio de Amazon para el reconocimiento facial, esta organización descubrió que el software coincidía incorrectamente los rostros de 28 miembros del Congreso de los Estados Unidos (congresistas con diferentes tonalidades de piel morena) con las fotos de una serie de delincuentes arrestados (Armasu, 2018).

En el año siguiente, un estudio del *Massachusetts Institute of Technology* (Singer, 2018) mostró que este software identificó de manera errónea a las mujeres de piel oscura (darker-skinned women) como hombres el 31 por ciento de las veces, pero no cometió ningún error con los hombres de piel clara (light-skinned men)¹². Frente a esto, Amazon comentó que los miedos derivados de estos resultados eran "an insignificant public policy issue." (Pasternack, 2019), restándole importancia al tema.

Esta situación no es una cuestión publica banal. Su uso abiertamente discriminatorio hacia ciertos grupos poblacionales, junto a su implementación en las fuerzas policiales, generaban un reclamo legítimo frente a las posibles violaciones de los DDHH derivadas de su utilización. Un miedo que no es nada insignificante. La respuesta de Amazon fue lo que podríamos esperar: Discursos vacíos que evadían cualquier tipo de responsabilidad. La empresa se empeñó en criticar en primer lugar a la UCLU por utilizar una "configuración incorrecta" o diferente a la recomendada de *Rekognition* (Singer, 2019), Incluso, criticaron a los Gobiernos por el mal uso de su software, ya que Amazon posee una «Política de uso aceptable» que prohíbe el uso de sus servicios para actividades ilegales, que violen los derechos de los demás o que puedan ser dañinos en general.

¹¹ Un control y vigilancia de esta clase solo lo hemos visto en la serie de televisión británica *Black Mirror* o en las peores distopías literarias cyberpunk.

¹² Este tipo de sesgos raciales y por género, ya habrían sido previamente mostrados por las investigadoras del MIT, Joy Buolamwini y Timnit Gebru en el caso de los softwares de IBM y Microsoft. Así como la plataforma Face++ de la compañía china Megvii (Vid. *Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification*)



Estas declaraciones nunca hicieron referencia al algoritmo en sí mismo, ni a los ejemplos discriminatorios utilizados por *Rekognition* para entrenarse, a los sesgos por género o raciales encontrados en los estudios, ni muchos menos a la responsabilidad social de Amazon cuando vende su tecnología para fines represivos o bélicos. La culpa es de quién usa la tecnología, no de ellos. Con esta alusión (ya clásica) a la neutralidad tecnológica no solo evaden su responsabilidad social, sino que además le atribuyen totalmente a los Estados los posibles malos usos, así los problemas éticos y legales correspondientes.

Para sustentar esta tesis, la empresa a través de su Gerente General de IA indicó en una publicación oficial -junio del 2018- del Blog de Amazon que no habían recibido ni un solo reporte sobre algún tipo de abuso policial con la utilización de *Rekognition*. A pesar de esto, también en una entrada del mismo Blog – febrero del 2019- el Vicepresidente de políticas públicas globales de *Amazon Web Services* escribió que "We support the calls for an appropriate national legislative framework that protects individual civil rights and ensures that governments are transparent in their use of facial recognition technology." y para esto señaló -de nuevo- una especie de mini-guía ética a manera de manual sobre la IA en general (Punke, 2019), centrada en los softwares de reconocimiento facial.

En esa entrada se indicaron un conjunto de principios como los siguientes: En el caso del reconocimiento facial siempre debía existir revisión humana. Utilizar softwares con un umbral de 99 por ciento de confianza cuando se use para fines que pueden lesionar los DDHH (¿Contradicción con la *Política de uso aceptable*?) y en los siguientes puntos más de lo mismo: Respeto a la ley. Transparencia por parte de los Gobiernos y avisos para que los ciudadanos sepan que estas tecnologías se están utilizando en espacios públicos o comerciales. Orientaciones que, aun siendo válidas en general, son intencionadamente superficiales, vagas y sin ninguna referencia a la responsabilidad empresarial, al desarrollo de tecnologías intrusivas, las lesiones presentes a los derechos (cometidas actualmente por Amazon), mecanismos de corrección, políticas concretas o formas específicas para impedir su uso «ilegitimo».



Inyectar ética.

Como podemos ver esto es parte de una instrumentalización que distorsiona el discurso ético para sus propios fines. Tomemos como ejemplo el caso de Reid Blackman. Un filósofo estadounidense, fundador y director ejecutivo de *Virtue*, una consultora de riesgos éticos (gestión del riesgo ético) que trabaja con empresas líderes en IA¹³ para integrar la ética y la mitigación de riesgos éticos en la cultura empresarial, así como el desarrollo, despliegue y adquisición de productos de tecnología emergente.

Un riesgo ético es una consecuencia negativa inesperada, resultado de acciones no éticas. En este caso, estos riesgos se asocian a problemas reputacionales, disminución de capital de inversión, detrimento de credibilidad pública y obviamente, pérdida de clientes. Es por esto por lo que en la página web oficial destinada a promocionarse, presenta el slogan: "How do you inject ethics into culture and product?"

¿Qué quiere decir exactamente inyectar ética en un producto? No queda claro. Lo cierto es que según Blackman, los consumidores quieren comprar productos o servicios en "empresas éticas", por lo que las consecuencias de no seguir un "enfoque ético empresarial" son graves.

These companies are investing in answers to once esoteric ethical questions because they've realized one simple truth: [...] Missing the mark can expose companies to reputational, regulatory, and legal risks [...] leads to wasted resources, inefficiencies in product development and deployment, and even an inability to use data to train AI models at all. (Blackman, 2020).

Posiblemente estas "preguntas éticas esotéricas" a las cuales se refiere Blackman, son preguntas básicas sobre la responsabilidad empresarial, justicia social o respeto a los derechos humanos. Preguntas que, a pesar de lo que el estadounidense pueda decir, aún ahora son ignoradas, en vista de que una propuesta como esta (*inject ethics into culture and product*), lo único que busca es

¹³ Una de las empresas con las que colaboró Blackman fue Amazon, tal y como lo indica en su página web. (<https://reidblackman.com>)



establecer una apariencia de corrección y responsabilidad empresarial y social, que tiene el propósito de confundir a los legisladores y a los encargados de políticas nacionales, así como a los consumidores finales, con tal de mantener las prácticas empresariales intactas, en las cuales se mantiene incuestionado el sistema económico subyacente, las malas prácticas empresariales con sus empleados o las transgresiones flagrantes a los DDHH. Todo esto cubierto con un hábito de conciencia y preocupación social.

Blackman señala que si bien, introducir declaraciones o planteamientos éticos no es necesario (un argumento bastante dudoso) para la operación de las empresas, debería hacerse ya que son "herramientas útiles para lograr sus objetivos" (Tardif, 2021). Es decir, desde esta lógica utilitarista, la única intención de las propuestas éticas es la obtención de fines comerciales (basados en el lucro y la optimización de recursos). Por lo tanto, evadir la regulación o evitar una disminución en su reputación, no es en sí mismo algo éticamente indebido.

Es así como esta "inyección de ética" viene acompañada de conferencias, conformación de grupos de expertos o con la creación de centros de investigación sobre Ética de la IA -como el caso de Facebook en Munich (Quiñonero, 2019)- que nunca llegan a ningún punto. Tampoco es extraño encontrar reflexiones sobre el futuro de la IA, que lo único que hacen es evadir las consideraciones sobre los temas urgentes en el presente, manteniendo el *status quo* empresarial intacto.

Estas propuestas simplemente cortinas de humo que sirven para exponer a los legisladores que el autogobierno empresarial y los mecanismos autorregulatorios en la industria son suficientes, por lo que no se necesita ninguna regulación específica para reducir los posibles riesgos tecnológicos, eliminar los escenarios de abuso o impedir las constantes violaciones a los DDHH (en temas como privacidad, discriminación, libertad de expresión etc.). A pesar de en algún caso concreto, estas empresas puedan solicitar leyes más concretas sobre los sistemas de IA, estas peticiones son vagas y superficiales.

Esto nos plantea un escenario muy distinto al que nos podríamos imaginar cuando una empresa quiere resolver los problemas éticos que se le presentan. En este contexto el debate sobre una IA ética va en paralelo con una resistencia a toda regulación. Cualquier idea de (auto)



regulación empresarial de las tecnologías tiende a dejar de lado un papel activo de Estado y, en cambio, busca aumentar la influencia del sector privado, llegando incluso a ver a las leyes como «un simple obstáculo» para la innovación tecnológica. Según Calo "Ethics guidelines of the AI industry serve to suggest to legislators that internal self-governance in science and industry is sufficient, and that no specific laws are necessary to mitigate possible technological risks and to eliminate scenarios of abuse" (Hagendorff, 2020). La finalidad última no es aplicar principios éticos para crear una tecnología menos intrusiva, más accesible o para erradicar los sesgos discriminatorios de sus algoritmos, sino para la mitigación de riesgos éticos que puedan traducirse en perjuicios empresariales.

Desde luego, esto no quiere decir que la ética no sea importante en las empresas, sino solamente que existe una tendencia que utiliza el discurso ético de manera floja, simplona o que lo instrumentaliza como una gestión de riesgo que solo busca seguir generando beneficios económicos. Si, además, a esto le agregamos la imposibilidad de crear una regulación homogénea para el uso de las tecnologías (mucho menos aún para la IA, debido a sus aspectos cambiantes y disruptivos) se crea un terreno fértil para que las corporaciones y empresas utilicen a su antojo la ética, seleccionando aquellos principios que más les conviene (*ethical shopping*), vaciándola de contenido o trivializándola (*ethical washing*). Generando, de manera creciente, la apariencia de que la toma de decisiones políticas y legales sobre estos temas están descontextualizadas o carentes de asidero tecnológico.

Conclusiones

En estas páginas se sostiene que existe un uso instrumental de la ética (al menos de las grandes corporaciones) que va más allá del *ethical washing*. Una gestión (ética) del riesgo que solo busca de mantener los beneficios económicos de las empresas. El discurso ético se vuelve una fachada para evitar la verdadera regulación. Pero incluso va más allá, la autorregulación ética, en concreto lo que busca es la desregulación normativa, promoviendo como la única solución la gobernanza impulsada por el mercado, y no por los Estados. Es una instrumentalización del lenguaje ético que genera una



aparición de una empresa social y éticamente responsable. Simplificando el verdadero trabajo ético, creando comités solo para la opinión pública o contratando a filósofos especialistas en ética, sin ningún poder real en modificar las políticas de empresa o sin siquiera tener una participación en el diseño de propuestas. El uso autorregulatorio de la ética es solo parte de una estrategia comunicativa que busca crear y consolidar una imagen de responsabilidad en los clientes y en la sociedad. En el mejor de los casos la ética se convierte en un *risk managing* en el que no importa la justicia social, ni la protección o resguardo de los seres humanos. Pareciera que este tipo de empresas no solo buscan el control sobre los consumidores, sino además sobre los políticos. Como resultado de esto, muchas de estas iniciativas, en particular las patrocinadas por la industria, no pasan de ser (en el mejor de los casos) señales de buenas intenciones y de conductas adecuadas con el único propósito de retrasar la regulación, centrar el debate en problemas abstractos (lo cual retrasa más aún la discusión) y evitar cualesquiera soluciones técnicas.

Referencias

- Adler, D. & Schneider, J. (2020). Amazon workers are fighting for their rights. This holiday season, think of them. Retrieved 6 May 2021, from <https://www.theguardian.com/commentisfree/2020/dec/01/amazon-workers-fighting-for-their-rights>
- Algorithm Watch. (2020, April 30). AI Ethics Guidelines Global Inventory. *Algorithm Watch*. <https://inventory.algorithmwatch.org/>
- Armasu, L. (2018). Amazon 'Rekognition' Falsely Identifies 28 Congress Members as Criminals (Updated). *Tom's Hardware*. (2021). <https://www.tomshardware.com/news/amazon-rekognition-congress-members-criminals,37515.html>
- Blackman, R. (2020). A Practical Guide to Building Ethical AI. *Harvard Business Review*. <https://hbr.org/2020/10/a-practical-guide-to-building-ethical-ai>
- Buolamwini, J. & Gebru, T. (2018). Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification. *Proceedings of Machine Learning Research*.



- Clarke, R. (2019). Principles and business processes for responsible AI. *Computer Law and Security Review*. <https://doi.org/10.1016/j.clsr.2019.04.007>
- Coeckelbergh, M. (2019) Artificial Intelligence: Some ethical issues and regulatory challenges. *Technology and Regulation*, 31-34.
- EU Commission (2017) 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))
- Floridi, L. (2019). Translating Principles into Practices of Digital Ethics: Five Risks of Being Unethical. *Philosophy & Technology*, s13347-019-00354-x. <https://doi.org/10.1007/s13347-019-00354-x>
- Google. (2019). *Perspectives on issues in AI governance* (pp. 1–34). <https://ai.google/static/documents/perspectives-on-issues-in-ai-governance.pdf>
- Hagendorff, T. (2020). *The Ethics of AI Ethics: An Evaluation of Guidelines*. *Minds and Machines*. doi: 10.1007/s11023-020-09517-8
- Jobin, A., Ienca, M., & Vayena, E. (2019) The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9): 389-399.
- Joshi, A. (2019) Google dissolves its Advanced Technology External Advisory Council in a week after repeat criticism on selection of members. *Packt*. <https://hub.packtpub.com/google-dissolves-its-advanced-technology-external-advisory-council-in-a-week-after-repeat-criticism-on-selection-of-members/>
- Morte, R & Monasterio, A (2017). Entrevista a Ramón López de Mántaras. *Dilemata* (24), 301–309.
- Ng, A. (2018) Amazon is selling facial recognition technology to law enforcement. *C/Net*. <https://www.cnet.com/news/amazon-is-selling-facial-recognition-technology-to-law-enforcement/>
- OECD. (2019). *AI Principles*. <https://oecd.ai/ai-principles>
- Partnership on AI. (2018). *About us*. <https://www.partnershiponai.org/about/>.



- Pasternack, A. (2019) Amazon says face recognition fears are "insignificant." The SEC disagrees. *Fast Company*. <https://www.fastcompany.com/90329464/amazon-cant-block-investor-vote-on-face-recognition-says-sec>
- Pichai, S. (2018). *AI at google: Our principles*. <https://www.blog.google/technology/ai/aiprinciples/>
- Prisma European Network. (2021). *Digital Sovereignty: GAFAM, BATX, and ... the European Union?* <https://prisma-network.eu/our-work/DIGITAL-SOVEREIGNTY-GAFAM-BATX-EUROPEAN-UNION>
- Punke, M. (2019) Some Thoughts on Facial Recognition Legislation. *Amazon Web Services*. <https://aws.amazon.com/es/blogs/machine-learning/some-thoughts-on-facial-recognition-legislation/>
- Scheiber, N. & Wakabayashi, D (2019) Google Hires Firm Known for Anti-Union Efforts. *New York Times*. Tecnología. <https://www.nytimes.com/2019/11/20/technology/Google-union-consultant.html>
- Simonite, T. (2020). Google Offers to Help Others With the Tricky Ethics of AI. *Wired*. <https://www.wired.com/story/google-help-others-tricky-ethics-ai/>
- Singer, N. (27 de julio 2018). ¿La tecnología de reconocimiento facial de Amazon puede ser racista? *New York Times*. Tecnología. <https://www.nytimes.com/es/2018/07/27/espanol/amazon-rekogniton-aclu.html>
- Singer, N. (2019). Amazon is pushing facial technology that a study says could be biased. *MIT Media Lab*. from <https://www.media.mit.edu/articles/amazon-is-pushing-facial-technology-that-a-study-says-could-be-biased/>
- Tardif, A. (2021). Reid Blackman, Ph.D, Founder and CEO of Virtue Consultants. *Interview Series*. <https://www.unite.ai/reid-blackman-ph-d-founder-and-ceo-of-virtue-consultants-interview-series/>