

# ¿Regular o esperar en inteligencia artificial?



Por el momento, el mundo busca partir de principios éticos, y establecer una regulación que no ahogue la innovación en IA, pero asuma sus peores riesgos.

REDACCIÓN ISTMO

Una mala regulación sería más perjudicial para el ecosistema de la inteligencia artificial, que la carencia de ella. Sin embargo, no sólo es tarea del legislador entender a fondo las implicaciones de esta tecnología, que podría extenderse a toda la economía. El empresario mexicano, aún el que se incorpore como usuario, debe conocer los límites éticos y regulatorios que el uso de ésta implicará en el futuro.

Sobre este tema charlaron para *istmo*, los profesores Alejandro Salcedo Romo, de las áreas de Factor Humano y Empresa-Familia, y Philippe Prince-Tritto, investigador docente de la Universidad Panamericana y fundador del Laboratorio de Derecho e Inteligencia Artificial de la misma institución. Ambos académicos están relacionados con el Derecho y describen algunos de los aspectos sobre la regulación de la IA, y las implicaciones que esto tendría para las empresas en México.

### Alejandro:

De formación inicial soy abogado; mi carrera de licenciatura fue Derecho, y a nivel posgrado me enfoqué en Filosofía. Lo que se busca en esta edición de *istmo* es explorar desde distintos lentes la inteligencia artificial (IA) y sus implicaciones con el mundo de los negocios. En ese sentido es muy relevante contar con la perspectiva jurídica. En ello quisiera adentrarme en esta charla, pero me gustaría contextualizar primeramente con tu perfil. ¿Cuál es tu ocupación actual y cómo llegaste a los temas de la IA?

### Philippe:

Soy francés, tengo 35 años y hace cuatro años entré en la Universidad Panamericana como investigador docente de tiempo completo. Vengo del sector de la Ingeniería en Ciencias de la Computación. Cursé mis estudios en un momento en el cual estábamos saliendo de lo que podemos llamar «el invierno de la IA».

En ese entonces no se hablaba mucho de ella. Las técnicas ya estaban aquí, pero no necesariamente se usaba tanto el término. En 2010 hubo un *boom* de la IA, porque empezaron a funcionar muy bien ciertas técnicas, principalmente el Deep Learning.

Yo estaba trabajando en bases de datos, en software de gestión y me empecé a preguntar qué quería hacer de mi vida. Como todo el mundo, buscaba un sentido y no lo encontraba tanto en la ingeniería pura y dura. Necesitaba algo más, que pudiera ayudar a la sociedad, que tuviera más en cuenta el factor humano, y pensé en varias opciones: una de éstas era el Derecho. Empecé a estudiar Derecho nuevamente desde licenciatura y luego Derecho Digital. Después me fui a trabajar en *legaltech*.

Mi objetivo era realmente explorar desde un punto de vista de Ciencias de la Computación qué vínculos podía haber con el Derecho. Muy ingenuamente pensaba que con mi *background* en ingeniería iba a poder automatizar todo el Derecho y resolver este gran problema. Obviamente me di cuenta que la cuestión es mucho más compleja. Hubo muchos cambios de opinión y de perspectiva. Después de mi maestría trabajé en Protección de Datos personales en París, en una consultoría, principalmente para hacer un software que permitiera cumplir con el reglamento europeo de protección de datos.

Luego regresé a la UP y el proyecto que estamos llevando a cabo desde hace cuatro años aquí es construir un Laboratorio de Inteligencia Artificial y Derecho, que va a explorar el tema en sus dos vertientes. La primera vertiente va a ser la regulación de la IA. Es un aspecto meramente jurídico donde nos preguntamos cuáles son las mejores técnicas para regularla. Analizamos lo existente y hacemos Derecho Prospectivo también, Propiedad Intelectual, etcétera.

La otra vertiente, lo más más peculiar de este laboratorio, es que hacemos experimentos para poner a prueba las hipótesis acerca de la digitalización del Derecho. Es decir, la automatización de procesos jurídicos, con preguntas que hace unos años podían parecer fantasía y fuera de lo real: si la IA va a reemplazar a los abogados. Obviamente no nos planteamos esas preguntas de una manera tan trivial, sino que vamos a hacer experimentos específicos para ver hasta dónde podemos llegar con la IA y con la automatización de ciertas tareas jurídicas. Para esto usamos técnicas de ciencia muy novedosas.

Al mismo tiempo estoy haciendo el doctorado en Inteligencia Artificial. Regrese a mi primera

carrera, también en la UP. Mi tesis de doctorado trata de la causalidad y el análisis de textos jurídicos, porque hay una rama de la IA que se llama Inteligencia Artificial Causal y estoy explorando estas técnicas para ver cómo pueden aplicarse al procesamiento de textos jurídicos.

### Alejandro:

Tienes un *background* que me parece idóneo. Déjame conjugar tres puntos de lo charlado hasta ahora: por un lado tu formación en la parte técnica: ingeniería, y ahora el doctorado en Inteligencia Artificial. Por otro, esta pregunta que se convierte en un parteaguas: con qué sentido, con qué propósito, con qué finalidad. Es una pregunta que me parece fundamental al momento de plantearnos cómo queremos desarrollar e implementar este tipo de tecnologías. El tercer punto es la parte jurídica.

Déjame arrancar por la relación entre IA y Derecho pues, lo sabrás muy bien, ordinariamente el Derecho va detrás de los cambios sociales. Esto típicamente se debe a la velocidad que tiene la evolución de las prácticas, de las nuevas técnicas, aún aquellas no tan sofisticadas como la inteligencia artificial. ¿Cómo se da este primer encuentro con la IA? ¿Cuál es el momento en que el Derecho dice: «esto sí va a tener un impacto, tengo que atenderlo»? Además, hablando del *timing*, ¿qué tan tarde llega el Derecho cuando se empieza a entablar esta discusión o diálogo con la inteligencia artificial?

### Philippe:

Hace mucho que el Derecho se interesa por la tecnología. No es la primera tecnología que trata. Por ejemplo, está la firma electrónica. Desde hace mucho el Derecho se interesa por la protección de datos personales, que deriva de la protección de la vida privada. Incluso fue una construcción en buena medida jurisprudencial, antes de ser reconocida en los textos fundamentales. ¿El Derecho va detrás de las tecnologías? sí, pero esa es su naturaleza.

Podemos siempre hablar de Derecho Prospectivo, una vez que ya tenemos un contexto social que nos permite vislumbrar algún riesgo o cambio importante. Existe un libro muy recomendable de Nassim Taleb, que se llama *El cisne negro*, en donde se plantea un evento que va a aportar





## ¿El Derecho va detrás de las tecnologías? Sí, pero esa es su naturaleza.

un cambio social, un impacto muy importante que no se puede vislumbrar con antelación. Es decir, trata de la aleatoriedad de la sociedad.

En materia de tecnología los Cisnes Negros toman un poco de tiempo para establecerse, pero no dejan de ser fenómenos que van a cambiar completamente ciertas partes de nuestra actividad, sin que lo podamos ver con antelación, a menos que se trate de una especulación. Sin embargo, el Derecho no se fundamenta en especulación: va a seguir alguna brecha cuando ésta se pueda ver.

Ahora bien, para la IA los problemas han empezado a ocurrir con la protección de datos personales y los legisladores del mundo han podido interesarse en el tema bastante pronto. En la Unión Europea desarrollaron una directiva y luego un reglamento en 2016 sobre la protección de datos personales, que si lo aplicamos al pie de la letra, ya tenemos ciertas barreras o cierto marco para el desarrollo de la IA.

La inteligencia artificial va a utilizar datos; es una constante que debemos tener en cuenta. Si no hay datos, no hay IA, pues ésta nace con la acumulación y la recolección masiva de datos. Si no existieran esos datos masivos, no necesitaríamos IA. Muchos países se han alineado a esta tendencia planteada por el Parlamento Europeo y han empezado a legislar sobre el tema.

Otros también han empezado a interesarse por la IA, porque puede cambiar muchos aspectos de nuestras vidas, como es el trabajo: automatizar lo más rutinario, incluso ciertas tareas que pensábamos eran solamente propias de la humanidad, como hacer tareas de redacción. Entonces, está principalmente el interés de proteger el factor humano.

Varias organizaciones internacionales han empezado a interesarse. El proceso de legislar es lento, porque si hay algo peor que no tener regulación es tener una mala regulación. Efectivamente, una vez que la regulación entra en vigor, si está mal hecha, no solamente va a impedir tutelar los derechos que busca proteger, sino que además va a retrasar las innovaciones y los beneficios que podría traer cierta tecnología.

Si hablamos de regulación tecnológica, el proceso para encontrar los términos que capten realmente el arreglo correcto, sin inhibir que las empresas innoven e inviertan es una

tarea muy difícil. Muchas veces las organizaciones internacionales van a funcionar haciendo consultas públicas, ya sea con la academia o con el sector privado de cada país miembro. Todo eso toma tiempo.

Entonces el encuentro entre Derecho e Inteligencia Artificial es algo constante, que realmente no va a ser de un día para otro. Es algo que se lleva a cabo desde la invención de la técnica, como el libro, por ejemplo. Desde que el humano es humano, creamos técnicas. Cuando hay técnica hay novedad, y cuando hay novedad hay riesgos y entonces tendrá que haber Derecho.

### Alejandro:

Mencionabas que el desarrollo de la regulación ha sido más bien jurisprudencial, atendiendo precisamente casos que permiten voltear a ver esta realidad y tener una respuesta en tiempo real. En este sentido, ¿cuáles son los casos icónicos que llegan para detonar el encuentro entre Derecho e Inteligencia Artificial?

### Philippe:

A eso me refería principalmente para el desarrollo de la protección de datos personales. Actualmente hay casos de IA que están en curso, es decir, va a haber cierto desarrollo jurisprudencial.

Hay muchas incógnitas y casos en curso para cuestiones de responsabilidad sobre la IA en Estados Unidos, relacionados con temas de Propiedad Intelectual. Hay una disputa que existe ahora, por ejemplo, el problema jurídico de saber qué derechos tenemos si la IA ha copiado nuestro trabajo o algo que hemos publicado. Se involucra a todos los sistemas de Inteligencia Artificial Generativa, como DALL-E, o ChatGPT, que van a crear texto, imágenes y videos contestando un *prompt* –es decir, una solicitud del usuario–. Van a requerir bases de datos muy grandes e identificar en ellas patrones, para aprender cómo puede contestar la pregunta del usuario. Entonces, si hemos publicado en línea ciertos datos, hay una buena probabilidad de que la IA generativa haya sido entrenada con los datos publicados en línea.

Por tanto, aquí hay problemas que surgen con los artistas. Por ejemplo, que digan: «yo no quiero que la IA pueda reproducir mi trabajo». Hay varias formas de proteger el trabajo: podríamos publicar en línea con Copyright si hablamos de Estados Unidos, con derechos de autor. También hablamos de marcas de agua y metadatos, es decir, datos sobre los datos para identificar los trabajos publicados. En noviembre de 2022 se presentó una demanda, una *class action*, para tratar de ir a contra esta práctica. Es uno de los casos que hay que seguir de 2022 y hay varios casos por el estilo.

Más allá de la jurisprudencia, también hay textos que ya están en preparación: textos jurídicos, algunos de los cuales están ya vigentes. Por ejemplo, una recomendación de la UNESCO

## La ética de la IA, si la analizamos, muchas veces es utilitarista porque en buena medida vamos a tener aplicaciones que no tendrían que vulnerar los derechos humanos, pero sí lo hacen.

sobre la ética de la IA, que ya está adoptada y firmada por México. Hay otros textos en preparación, como el reglamento de la Unión Europea.

### Alejandro:

Pensando en estos documentos, ¿cómo estamos en México? ¿Qué hay sobre este tipo de regulación? Es un momento en el cual el tema de Protección de Datos Personales ofrece una guía para este tipo de desarrollo, pero más allá de la Ley Federal de Protección de datos personales en posesión de particulares y su reglamento, ¿en México qué hay en torno a este tema? ¿Se está cocinando alguna ley? ¿Hay algún caso icónico? ¿El IMPI ha dicho algo? ¿Qué hay dentro de la regulación mexicana en torno a la IA?

### Philippe:

Buena pregunta: vamos a decir que sí hay algunas cosas, principalmente una estrategia digital nacional. Es un documento no enfocado directamente a la IA. Podemos ver que hay ciertas iniciativas, hubo una de un senador del PAN hace algunos meses, pero en concreto no hay algo específico para la IA. Muchas veces vamos a aplicar mecanismos de Derecho Común, Derecho Civil, Derecho Administrativo y en realidad es un debate interesante, porque muchos se preguntan si necesitamos una regulación o adaptar todo nuestro Derecho.

La IA se aplica en tantos ámbitos, que en realidad lo que vamos a tratar de adaptar es toda la regulación, administrativa, todo el Derecho Administrativo, para incluir temas de toma de decisión automáticas, etcétera. O de responsabilidad por el uso de IA, o de limitar los sesgos de la IA: por el momento eso no se ha hecho.

Pero tampoco se ha hecho la alternativa, es decir, una regulación en específico para la IA. Eso es también muy difícil de hacer. Es muy difícil incluir todos los temas que vamos a necesitar para regular la IA en una sola ley. Es sumamente difícil hacer una ley suficientemente precisa, pero no tan precisa, porque podría llegar ser una ley de 500 páginas, para poder regularla.

¿Entonces qué hacer? Generalmente, y lo que seguramente México va a hacer, es inspirarse en algunas formas de regular. Existen dos formas de regulación que pueden ser usadas. La primera es una basada en principios. Aquí, por

ejemplo, se hablaría de la IA centrada en el Ser Humano, el respeto a la protección de datos personales, etcétera. Esto es, principios generales lo suficientemente flexibles para aplicarse a muchos casos, como conocemos en Derecho Civil.

El otro enfoque es una regulación basada en el riesgo: el acercamiento que ha tenido la Unión Europea. Pero la razón por la que me refiero mucho al Derecho Internacional es porque no va a bastar con hacer nuestra pequeña regulación nacional e ignorar al mundo. La IA es un tema trascendental, porque estamos hablando de compañías multinacionales. Estamos tratando datos que vienen de muchos países.

Obviamente, en cierta medida necesitamos una cooperación internacional. Además, México particularmente es un país con muchos acuerdos de libre comercio y que tiene vínculos con muchos países y con la Unión Europea en particular. México es un aliado estratégico.

Entonces, regresando al tema del reglamento, el AI Act adopta un enfoque basado en riesgo -incluye algunos principios-, que clasifica los sistemas de IA según el riesgo. En ese sentido llegan a existir los riesgos inaceptables, en los cuales de plano esos sistemas de IA tienen que prohibirse. Están, por ejemplo, los algoritmos que van a tratar de hacer daño en calificación social, como se puede conocer en China, y luego están los de alto riesgo, mediano y bajo.

Según el riesgo, vamos a tener ciertas obligaciones en materia de seguridad, de pruebas de uso de los datos, conservación e información al público sobre el uso de esos sistemas. Aquí no se legisla solamente para los que van a producir los sistemas de IA; también van a existir obligaciones para los usuarios: las empresas que van a adoptar sistemas de IA. Entonces tiene que haber una colaboración, porque ellos también van a ser responsables.

Todos los que tratamos de datos personales en México sabemos el impacto que ha tenido el GDPR. Es muy probable que la regulación sobre IA de la Unión Europea tenga una estricta extraterritorialidad, que va a influenciar a regulaciones nacionales, pero que además en buena medida va a tener que aplicarse por las empresas a nivel nacional. Para concluir sobre este tema, regulación específica no tenemos, generalmente van a ser mecanismos de Derechos Civiles. Estamos esperando controversias que se vayan a poder resolver. Ya podemos estudiar el Derecho Internacional e incluso el Derecho Regional, como en el caso de la Unión Europea.

### Alejandro:

Es impresionante cómo se detonó este tema; particularmente desde noviembre del año pasado, la velocidad con la que han lanzado todo tipo de herramientas, que están al alcance de prácticamente cualquiera. Parecen tremendamente atractivas para un modelo de negocio, para aumentar la eficiencia, para ahorrar. Los empresarios mexicanos están queriendo implementar o



desarrollar esta tecnología, los distintos tipos de ella. ¿Qué debería tener en la mente la mujer y el hombre de empresa en México –parámetros, guías– para saber cómo interactuar, justo en este momento en el que no hay una regulación estable?

#### **Philippe:**

Se pueden inspirar en los principios que ya existen. La recomendación de la UNESCO sobre la ética de la IA es un marco internacional con principios generales que México va a tener que aplicar porque es firmante.

Dentro de esos principios, hay varios con los que ya nos tenemos que empezar a familiarizar, con un nuevo vocabulario jurídico y con nuevas prácticas que vamos a tener que implementar. Si es que realmente se aplica esta recomendación, y si seguimos la tendencia internacional, el desarrollo y uso de sistemas de IA ya no va a ser el «salvaje Oeste». Vamos a tener que aplicar ciertas reglas y los sistemas se van a tener que pensar muy bien para desarrollarse e implementarse.

También van a ser necesarios mecanismos de control interno, es decir, tal como tenemos un oficial para la protección de datos personales, vamos a tener un oficial de cumplimiento para la ética de la IA. Mejor empezar ahora a interesarse en esos temas, porque podría ser costoso en términos de tiempo y dinero, formarse y aplicar la reglamentación, si realmente hay un mecanismo para el *enforcement*.

¿Cuáles son esos principios? Voy a hablar de algunos: transparencia y explicabilidad, los sistemas de IA y sus procesos de toma de decisión deberán ser claros y comprensibles. Los empresarios deben ser capaces de determinar cómo un sistema de IA llegó a una decisión en particular, para que puedan interpretar si tomó los datos, por ejemplo del sector salud, al analizar un electrocardiograma. ¿Cómo llegó a esa conclusión? De la misma manera, si el sistema preconiza un tratamiento, debería mencionar sus fundamentos y la fecha de corte de su entrenamiento.

Eso no lo hacen las herramientas más conocidas actualmente, pues además la transparencia y explicabilidad significan que los sistemas deben ir acompañados de instrucciones de uso claras, comprensibles. Este principio está asociado

al Derecho a la Información de los que van a utilizar el sistema.

Aquí obviamente están los temas de complejidad de los sistemas de IA, que pueden dificultar la explicabilidad. No necesariamente tendríamos tiempo de entrar en esto, pero es el tema de los *black box*, la complejidad de los sistemas cuando usamos redes neuronales profundas para justificar una decisión y explicar cómo llegamos ésta es sumamente difícil. Se necesitan ingenieros para llevar a cabo esas tareas.

Segundo, sería el principio de responsabilidad. Las entidades deben ser responsables del funcionamiento de sus sistemas, asegurarse de comprender los riesgos conocidos y previsibles relacionados con el uso de IA, y eso implica estar al tanto de los posibles errores o incoherencias que puedan surgir en el sistema y cómo pueden afectar a la seguridad de los usuarios. Esto está vinculado al derecho a la reparación.

Otro principio es el de equidad. Los sistemas de IA no deben ser discriminatorios o reforzar prejuicios, y deben promover la igualdad y la no discriminación. Aquí hay toda una lista de deseos que plantea la UNESCO para los sistemas de IA, que son objetivos. Es decir, como siempre en materia de Derechos Humanos, hablamos de derechos adquiridos, pero en realidad muchas veces son derechos que queremos lograr. Para lograr la equidad, se trata del entrenamiento de los sistemas de IA, evitar que los datos estén sesgados y que den lugar a resultados discriminatorios.

Luego están los otros principios, como la protección de la privacidad y la gobernanza de los datos. Hablamos sobre todo de sistemas críticos, pero no solamente está el principio de robustez, seguridad y fiabilidad, los sistemas de IA durante todo su ciclo de vida deben ser resistentes a los ataques y fallos y funcionar de manera confiable. Aquí la participación del empresario en este proceso es indispensable.

Un principio importante es la vigilancia humana, que los sistemas de IA deben estar sujetos a intervención humana y a la supervisión adecuada para garantizar que funcionen como se espera, respetando los derechos humanos, los valores que están en la recomendación. Esto implica controlar el funcionamiento del sistema, detectar posibles errores o comportamientos inesperados

e intervenir si es necesario, para garantizar la seguridad de la precisión de los datos.

Es muy importante esto, y en la industria de la IA lo que se vislumbra son normas técnicas. El tema aquí es que la participación humana es fundamental; son procesos que tenemos que implementar, gente que va a constar que todo ha sido desarrollado de forma ética y que respete esos principios. De eso trata la norma ISO.

#### **Alejandro:**

No deja de ser interesante, seguramente ya habrás reparado en ello, que los marcos de referencia que tiene hoy el Derecho para enfrentar no sólo los retos, sino las oportunidades del desarrollo de la inteligencia artificial, sean éticos. No deja de ser curioso, tanto en la ISO como en el documento de la UNESCO. ¿Entonces la Ética y el Derecho se volverán a encontrar so pretexto de la IA? Otra pregunta que pudiera ayudar a aterrizar la anterior: Google, Twitch, Microsoft, por mencionar algunas, se han unido a la fila de las Big Tech que han despedido a sus equipos de IA. Entonces se necesita supervisión humana, los marcos son éticos, pero a su vez las empresas parecen ser que están actuando –al menos las Big Tech, que son las que van estableciendo el ritmo de desarrollo de muchas otras– precisamente en sentido contrario. Quitando la supervisión humana, la referencia ética, ¿qué piensas de esto?

#### **Philippe:**

Aquí estás tocando un tema sensible, que es completamente cierto. Estuve participando en un foro donde estaba el director jurídico global de Meta. Teníamos ciertos temas sobre los cuales definitivamente no estábamos de acuerdo. La «filosofía de las virtudes», como bien lo mencionas, parece decir: «Congreso, ustedes no entienden nada, déjenos innovar, porque si regulan, no vamos a poder innovar correctamente». Están tomando el problema al revés: no es que los derechos humanos impidan la innovación, sino que es la innovación, la IA, la que amenaza a los derechos humanos. Muchas veces se voltea esta lógica de una forma falaz –en mi opinión– sobre la ética. En el uso del término ética también soy muy crítico porque la ética se ha vuelto un adjetivo: esto es ético, aquello no es ético y yo no sé

## En la parte jurídica habrá que prepararse, leer sobre temas de ética, de epistemología, de IA, para saber lo que es posible y lo que no.

qué significa. ¿Qué significa que algo es ético? Eso no tiene sentido.

La ética no es un adjetivo. La ética consiste en todo un conjunto de reglas morales que pueden ser muy diferentes entre una persona y otra, para llegar a un resultado aceptable para una persona o un grupo de personas. Esa definición no es oficial, pero el tema es que vamos a tener un resultado muy diferente si usamos una ética que si utilizamos una ética utilitarista. Decir que algo es ético no tiene mucho sentido.

La ética de la IA, si la analizamos, muchas veces es utilitarista porque en buena medida vamos a tener aplicaciones que no tendrían que vulnerar los derechos humanos, pero si lo hacen. Eso significa que nunca podríamos desarrollar nada de IA, porque el riesgo cero no existe.

Pero por otro lado, nos podemos referir a la recomendación de la UNESCO. Si admite que haya *sandbox* para hacer pruebas, incluso a veces pruebas en vivo a gran escala de sistemas de IA, aquí es cuando hacemos un *trade off*: un cálculo sobre si el riesgo es aceptable o no.

A veces estamos más inclinados hacia una ética kantiana, porque para ciertos casos ya de plano no se puede usar la IA. Por ejemplo, para la calificación social. Si hablamos de ética kantiana, nunca debería usarse la IA en el desarrollo de armas letales autónomas. ¿Existen las armas letales autónomas? Sí. ¿Están prohibidas? No. Tenemos aquí cierta controversia sobre el tema, con intereses que sin lugar a duda son políticos, económicos y a veces un doble discurso.

Efectivamente, entre nosotros vamos a hacer todo ético, lo que sea que eso significa, pero eso no es una regulación que realmente nos vaya a imponer algo. ¿Por qué se usó la ética? Porque ética inicialmente es flexibilidad, es autorregulación y sabemos que la autorregulación no funciona, pero era para mostrar la buena cara. para decir, «pues a nosotros nos importan los derechos humanos». Vamos a ver qué ocurre cuando realmente ya nos traslademos del campo de la ética al campo de la regulación, que yo creo que sí sería necesario.

### Alejandro:

En medio de toda esta visión HD, digamos, tú tienes la perspectiva técnica, la jurídica, una



perspectiva también global de cómo lo están viendo desde las Big Tech, y a la par estás dentro de un laboratorio en una universidad mexicana. ¿Qué le pedirías a la mujer, al hombre de empresa mexicano de cara a este desarrollo e implementación de la IA? Es decir, ¿qué esperarías de ellos? ¿A qué los invitarías para ir creando innovación, desarrollo, una aplicación que sea más adecuada?

### Philippe:

Depende aquí del objetivo. Desde un punto de vista técnico, diría que es importante interesarse no solamente en la superficie de las cosas. Si hablamos de una empresa que no se dedica a la tecnología, la invitaría a que no sea solamente usuaria de tecnología, sino que trate de innovar con un equipo de científicos de datos. Eso puede permitir a México no depender de reglas de otros países, ya sea sociales, jurídicas o técnicas, sino adoptar una técnica e innovación nacionales. La IA es tan diversa que se debe adaptar a contextos y a realidades que difieren según los países. No hay un sistema que se vaya a adaptar a todo. En vez de usar nada más sistemas existentes de otros países, por buenos que sean, interesarse en la investigación y el desarrollo.

Eso para la parte técnica. Para la parte jurídica, prepararse, leer sobre esos temas justamente de ética, de epistemología, de IA, para saber lo que es posible y lo que no (eso es otro tema: evitar las falacias). Sí, hay muchas ventajas, pero también hay desventajas en la IA. En términos de pruebas de implementación, verificar si el sistema conviene a la tarea que quiero resolver. Pero sobre todo, ese estudio previo me va a servir para mi *compliance*.

Al implementar un sistema de IA, se debe dedicar cierto tiempo a estudiar que el sistema sea robusto, que lo pueda entender, que tenga la documentación específica, etcétera. Sé que va a requerir tiempo, que va a representar dinero, pero es necesario si quiero que mi práctica tenga cumplimiento de los derechos humanos. Respetar estos derechos no va a ser algo necesariamente fácil con la IA. Los empresarios aquí que quieran implementar la IA deben tomar en cuenta este paso, y mejor temprano que tarde. Una vez que ya esté funcionando, formalizar todos los procesos, hacer un mapeo, una gobernanza de los datos y los algoritmos que estamos usando, será más caro. Deben prever en materia de recursos y acercarse a profesionales que puedan dar una asesoría sobre esto, justo desde el principio. </>