

INTELIGENCIA ARTIFICIAL Y LA ERA DE LAS SOLUCIONES MÉDICAS INEXPLICABLES: NAVEGANDO LA OPACIDAD ALGORÍTMICA EN LA MEDICINA ACTUAL

CARLOS J. REGAZZONI^{1, 2}

¹Comité de Salud Global y Seguridad Humana, Consejo Argentino para las Relaciones Internacionales (CARI),

²Instituto de Salud Global, Universidad J. F. Kennedy, Buenos Aires, Argentina

Dirección postal: Carlos J. Regazzoni, Comité de Salud Global y Seguridad Humana, Consejo Argentino para las Relaciones Internacionales (CARI), Uruguay 1037, piso 1º, 1016 Buenos Aires, Argentina

E-mail: cregazzoni@gmail.com

Recibido: 30-IV-2025

Aceptado: 4-VII-2025

Resumen

Las tecnologías basadas en inteligencia artificial (IA) están transformando profundamente el cuidado de la salud mediante técnicas de *machine learning* (o aprendizaje automático) y redes neuronales. Estas herramientas estadísticas revolucionan la programación determinista clásica al aprender directamente patrones probabilísticos a partir de los datos. Estos sistemas permiten abordar problemas clínicos antes inaccesibles por complejidad o ambigüedad, optimizando así diagnóstico, tratamiento, gestión sanitaria e investigación biomédica. Sin embargo, esta revolución tecnológica plantea un desafío epistemológico sin precedentes: mientras la estadística tradicional busca explicaciones causales de los fenómenos bajo observación, los modelos predictivos de IA priorizan el rendimiento predictivo, independientemente de la comprensión teórica subyacente. Esta "opacidad algorítmica", derivada de modelos con millones de parámetros ajustados automáticamente, contrasta con el razonamiento clínico humano basado en métodos analíticos o heurísticos. La medicina, que históricamente fundamenta sus decisiones en causalidad y evidencia explicativa, se enfrenta ahora a herramientas predictivas cuya lógica interna frecuentemente resulta inaccesible incluso para sus desarrolladores. Esta divergencia genera retos significativos a nivel educativo, profesional y ético, requiriendo que los médicos adquieran nuevos equipamientos conceptuales en estadística avanzada e informática para na-

vegar efectivamente esta transición. Este artículo explora las bases bioestadísticas y conceptuales de esta tensión entre predicción y explicación en IA, contrastando ambas aproximaciones desde el proceso inferencial hasta la evaluación causal, destacando la necesidad imperiosa de que la medicina comprenda profundamente estos nuevos paradigmas para integrarlos críticamente en la práctica clínica y la investigación científica.

Palabras clave: inteligencia artificial, inferencia, estadística, razonamiento clínico, inferencia causal, algoritmo

Abstract

Artificial intelligence and the age of unexplainable medical solutions: navigating algorithmic opacity in today's medicine

Artificial intelligence (AI)-based technologies are profoundly transforming healthcare through machine learning techniques and neural networks. These statistical tools revolutionize classical deterministic programming by directly learning probabilistic patterns from data. Such systems enable the resolution of clinical problems previously inaccessible due to their complexity or ambiguity, thus optimizing diagnostics, treatment, healthcare management, and biomedical research. However, this

technological revolution presents an unprecedented epistemological challenge: whereas traditional statistics seek causal explanations for observed phenomena, predictive AI models prioritize predictive performance regardless of the underlying theoretical understanding. This “algorithmic opacity,” resulting from models with millions of autonomously adjusted parameters, contrasts with human clinical reasoning based on analytical or heuristic methods. Medicine, historically grounded in causality and explanatory evidence, now confronts predictive tools whose internal logic is often inaccessible even to their developers. This divergence poses significant educational, professional, and ethical challenges, requiring physicians to acquire new conceptual competencies in advanced statistics and informatics to effectively navigate this transition. This article examines the biostatistical and conceptual foundations of this tension between prediction and explanation in AI, contrasting both approaches from the inferential process to causal evaluation. It emphasizes the urgent necessity for medicine to deeply comprehend these emerging paradigms in order to critically integrate them into clinical practice and scientific research.

Key words: artificial intelligence, inference, statistics, clinical reasoning, causal inference, algorithm.

PUNTOS CLAVE

Conocimiento actual

- Existe una amplia difusión de los conceptos básicos de las bioestadísticas en medicina, pero ahora enfrentamos el desafío de asimilar los principios de funcionamiento de las tecnologías basadas en inteligencia artificial. La opacidad algorítmica se constituye así en un riesgo cierto.

Contribución del artículo

- El artículo desarrolla los conceptos de la bioestadística, la inferencia estadística e inferencia causal, y explica el funcionamiento de los modelos de inteligencia artificial. Este conocimiento permitirá que los médicos tengan un rol protagónico en la incorporación de estas tecnologías, basadas en la eficiencia y no en la inferencia causal.

Las tecnologías basadas en inteligencia artificial (IA) permean todas las dimensiones del

cuidado de la salud; precisión diagnóstica, tratamientos personalizados, investigación acelerada, y gestión de salud eficiente^{1,2}. Todo mediante aparatos que simulan nuestra inteligencia.

En esencia la IA representa una revolución de la programación informática³. Los programas, clásicamente convertían los valores de ciertas variables de entrada, en valores de salida de las variables de resultado, utilizando reglas explícitas codificadas por humanos, con una lógica determinista⁴. Vale decir, si la computadora recibe una información u orden específica, devuelve tal información y/o tarea definida. A una entrada le correspondía siempre una salida “programada”, y no otra. La IA en cambio, mediante redes neuronales (RN) y técnicas de *machine learning* (ML) o aprendizaje automático, cambia el paradigma⁵. El programa no sigue ya reglas predefinidas sino que aprende patrones probabilísticos directamente de los datos de entrada⁶, utilizando funciones estadísticas complejas. Durante el entrenamiento estos modelos ajustan sus parámetros de funcionamiento internos de manera autónoma para mejorar la precisión con que predicen los resultados requeridos por el operador. Detectan patrones complejos en los datos de entrada (por ejemplo, imágenes, lenguaje natural, o múltiples señales del ambiente) que resultaban impracticables de programar manualmente. Gracias a ello, no solo pueden procesar volúmenes masivos de datos, sino que también pueden abordar problemas cuyas reglas son ambiguas o desconocidas (luego no eran programables). Así, las cuatro funciones esenciales del quehacer médico: diagnóstico, toma de decisiones, prestación de servicios de salud, e investigación biomédica; se verán profundamente reconfiguradas y eficientizadas.

Sin embargo, en paralelo a esta revolución de la eficacia, la salud entra en una etapa de insurrección epistemológica⁷. La IA representa un giro copernicano respecto de las concepciones clásicas de la estadística, que podríamos resumir como un desplazamiento del foco epistémico desde la verdad explicable hacia el eficaz desempeño predictivo⁸. El principal objetivo de la inferencia estadística tradicional era identificar, en una fenomenología variable, las regularidades de la naturaleza que producían los datos, y a partir de allí generar predicciones sobre el comportamiento de los eventos en estudio, ba-

sándose en modelos explicativos. Por contraposición, los modelos de IA son evaluados por su capacidad de predecir eventos independientemente de la explicación mecanicista que pudiera subyacer, o no, en el proceso de inferencia. La medicina se enfrenta entonces a dispositivos dotados de capacidades computacionales avanzadas que ofrecen soluciones cuya fundamentación, en muchos casos, resulta opaca o difícil de interpretar. Esta disociación entre eficacia técnica y comprensión teórica plantea nuevos desafíos educativos, profesionales, y éticos, en salud.

Básicamente las herramientas de IA médica más potentes funcionan como una “caja negra” que, alimentada con datos, luego del análisis devuelve resultados tales como diagnósticos o recomendaciones terapéuticas, con pocas o ninguna explicación interpretable sobre el razonamiento utilizado para llegar a ellos⁹. Las funciones estadísticas utilizadas en IA, no son estructuras diseñadas para ser interpretables, sino para maximizar la precisión predictiva. Esta tensión entre rendimiento y transparencia¹⁰ se denomina “opacidad algorítmica”¹¹. Los modelos de IA, con millones de parámetros para modelar funciones matemáticas complejísticas, toman decisiones mediante procesos que incluso sus creadores no pueden comprender. De este modo nos alejamos de la tradición médica históricamente basada en la causalidad, la fundamentación científica, y la justificación de las intervenciones. Mientras que la ciencia se interesó siempre por entender, la IA propone resolver sin explicar¹², y los médicos nos encontramos carentes del equipamiento conceptual para navegar estas aguas.

Este artículo explora las características del razonamiento utilizado por los modelos predictivos y las técnicas de aprendizaje automático profundo que motorizan los sistemas de IA en salud. Utilizaremos como punto de partida los conceptos fundacionales de la bioestadística. Además, se contraponen el modelo de inferencia de la IA al razonamiento clínico humano y la teoría del diagnóstico¹³.

Del razonamiento clínico a la bioestadística

Los sistemas de IA funcionan de manera distinta a como lo hace el cerebro del médico. La

medicina evalúa situaciones clínicas como diagnóstico o tratamiento, inmersa en una constante incertidumbre¹⁴. Esta realidad médica caracterizada por la variabilidad y la imprecisión, gira en torno del llamado razonamiento clínico¹⁵; que es la problematización de una situación dada en los términos propios de la ciencia médica.

Actualmente se reconocen dos grandes aproximaciones al razonamiento clínico: la analítica, y la intuitiva experta o heurística¹⁶. La aproximación analítica aplica el método hipotético-deductivo y la deliberación. Avanza desde la obtención de información, a la formulación de algunas hipótesis, y la selección de un curso de acción. El método intuitivo experto o heurística¹⁷, por el contrario, es propio del profesional entrenado que, con menor cantidad de datos entiende rápidamente el problema y se mueve hacia la solución. En una aproximación heurística el decisor evita sobrecargas de información y se basa en algunos predictores para concluir, y actuar. Esta “racionalidad aproximada” advertida por Herbert Simon¹⁸ en 1956 (uno de los padres de la teoría de la decisión), en realidad es el modo de acción más común, cuando información y capacidad de cómputo son limitados. Es la forma como el ser humano maneja situaciones problemáticas complejas donde por definición una solución exacta es imposible¹⁹, y debe entonces encontrar un camino “satisfactorio” para decidir en incertidumbre. En el caso de la medicina, la experiencia profesional proporciona las capacidades necesarias para este tipo de razonamiento intuitivo²⁰.

Sin embargo, la cuestión es problemática, y la toma de decisiones en salud tendrá su contrapartida en el denominado error médico²¹, causa de 795 000 pacientes con secuelas serias o muerte cada año solo en los EE. UU.²². Y “error” es una palabra a retener, porque resulta indispensable para entender el funcionamiento de los modelos de IA y la revolución epistemológica en ciernes. La realidad del error médico llevó a intentos de enriquecer el razonamiento clínico y la resolución de problemas de salud, utilizando evidencia científica y mediante diferentes formalizaciones del proceso de toma de decisiones²³. Entra en juego entonces la bioestadística, a la cual ya nos hemos acostumbrado, excelente equipamiento conceptual, aunque frecuente-

mente soslayado²⁴, para asimilar lo que serán los modelos de IA y la nueva forma de generar soluciones en salud.

Bioestadística

Necesidad de las estadísticas en medicina

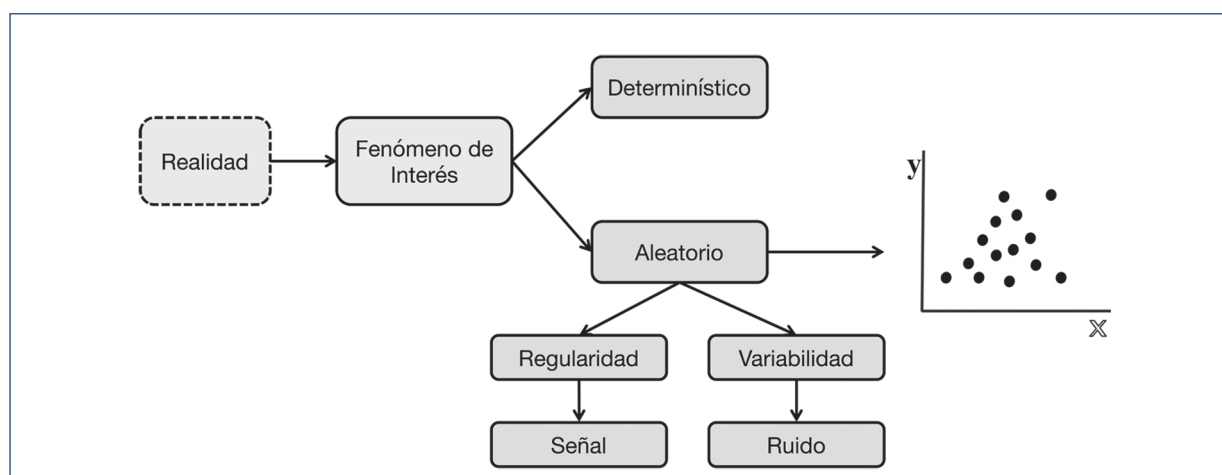
Como decía David R. Cox²⁵, solemos tomar al análisis estadístico por un mero ritual elaborado luego del trabajo arduo y verdadero, para satisfacer a árbitros en las publicaciones científicas y regulaciones gubernamentales. Sin embargo, enfrentados a la incertidumbre y a lo desconocido, tanto las exigencias del método científico como la demanda de racionalidad en la toma de decisiones invocan a la ciencia de la cuantificación de lo probable²⁶. En concreto, los médicos aprendemos de una experiencia que llega fragmentada y cambiante, y debemos introducir el análisis estadístico para extraer algún tipo de enseñanza durable respecto de los procesos que generan los datos que observamos. Sin la variabilidad en los atributos de los fenómenos que estudiamos²⁷, los eventos serían absolutamente predecibles, viviríamos en un mundo determinístico, y no necesitaríamos de la estadística (Fig. 1). Pero, como decía Nicolás de Cusa²⁸, la precisión no es de aquí. Luego debemos abordar la incertidumbre. La estadística²⁹ utiliza probabilidades para abstraer patrones o regularidades subyacentes en el cambio, y capturar en cantidades matemáticamente operables el nivel de incertidumbre de nuestras observaciones y con-

clusiones. La bioestadística, por su parte, aplica estos métodos estadísticos a fenómenos de interés específicos de las ciencias de la vida³⁰.

Fenómeno de interés y proceso aleatorio, y proceso generador de datos

Denominamos fenómeno de interés (ej.: enfermedad de base, efecto de un tratamiento, etc.) al aspecto de la realidad biológica que produce las observaciones que examinamos (Ej.: curva térmica, lesiones cutáneas). Estas observaciones, que varían en intensidad y frecuencia de aparición debido a factores biológicos, ambientales, o técnicos, son la manifestación empírica del fenómeno de interés. Dado que dicha variabilidad no puede explicarse completamente por mecanismos deterministas conocidos, es decir que no podemos establecer siempre el mecanismo que la causa, debemos modelar³¹ entonces el fenómeno como un proceso aleatorio (estocástico). Un proceso aleatorio es un modelo matemático que describe, mediante una familia de variables aleatorias, la evolución probabilística de un sistema. Este tipo de modelo permite cuantificar la incertidumbre, y realizar inferencias dinámicas sobre el comportamiento del sistema bajo análisis. Por ejemplo, caracterizar mediante un proceso aleatorio la distribución probabilística del sobrepeso en una población permite estimar el riesgo individual de desarrollar obesidad³². Se utiliza también el término “proceso estocástico” para describir

Figura 1 | Descripción de la clasificación del fenómeno de interés, en aleatorio y determinístico



estos fenómenos aleatorios cuyos valores evolucionan según leyes probabilísticas³³. Otra denominación, y con menores exigencias matemáticas, es la de “proceso generador de datos” (PGD). Llamamos PGD al mecanismo subyacente y real, generalmente desconocido, que produce las observaciones. El PGD se modela como una distribución de probabilidades (o varias, asociadas o no) a lo largo de los valores posibles de las variables en estudio³⁴ (espacio muestral). Por ejemplo, podríamos hipotetizar que algún modelo de regresión explique el PGD que produce las mediciones de estatura asociada a la edad³⁵, y concluir que la altura de los individuos de una población aumenta con la edad hasta cierto momento y luego desciende. Comprender el PGD permitirá realizar inferencias estadísticas sobre el fenómeno de interés³⁶, y, en nuestro ejemplo, predecir si una determinada estatura es adecuada para una edad dada. La ciencia médica utiliza múltiples soluciones estadísticas para modelar procesos aleatorios, como el modelo de Markov para estudiar cohortes de pacientes³⁷, modelos autorregresivos para anticipar la probabilidad de un fenómeno futuro³⁸, o diversos modelos de correlación para estimar la dinámica de contagios en una epidemia³⁹. La enfermedad es en sí misma de naturaleza estocástica⁴⁰, y modelarla como PGD permite entender cómo produce observaciones con variabilidad, interpretables con herramientas estadísticas (Fig. 2). La inteligencia artificial concibe la realidad como un vasto

proceso generador de datos, modelándola a partir de la inferencia estadística sobre datos observados.

Modelos conceptual y observacional

Pregunta a investigar y conocimiento previo fundamentan el modelo conceptual; una explicación biológica plausible de las observaciones que nos problematizan, y que debe preceder a las abstracciones matemáticas. Por ejemplo, a la observación de la incidencia de diarrea asociada a la ubicación geográfica, le podría convenir un modelo conceptual para explicar su potencial asociación con la calidad del agua (esto hizo John Snow en 1851)⁴¹. El modelo conceptual tiene que ver con nuestra teoría acerca del fenómeno de interés.

Al modelo conceptual le sigue la aproximación al fenómeno⁴² mediante un modelo observacional; una sonda exploratoria dispuesta a captar información sensible al fenómeno de interés (Ej.: la función cardíaca podría aproximarse por su actividad eléctrica) (Fig. 3). El modelo observacional puede ser experimental, retrospectivo, o simplemente experiencial. El modelo observacional aísla el fenómeno objetivo y permite realizar inferencias sobre la realidad latente o proceso generador de datos, que produce las observaciones⁴³.

En contraste, los modelos usados en IA son más bien exploratorios⁴⁴; procesan datos masivos (estructurados o no), no siempre apoyados

Figura 2 | Descripción del proceso generador de datos y su transformación en observación

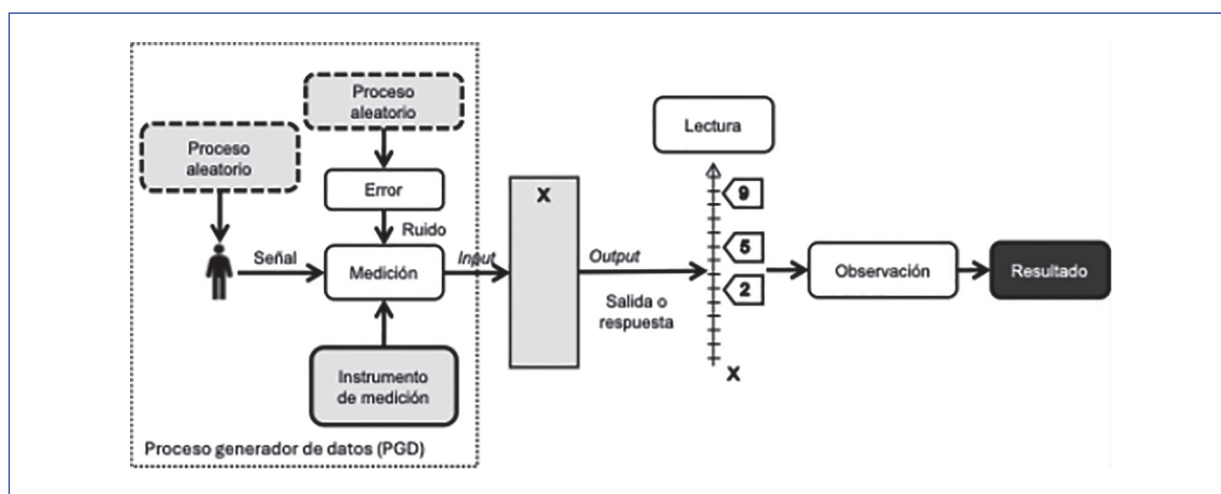
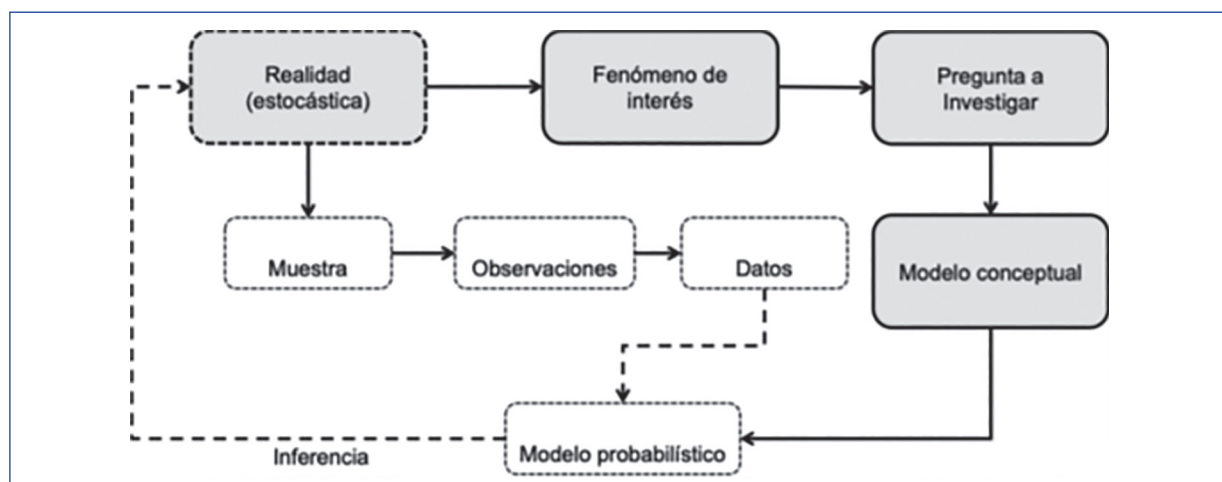


Figura 3 | Relaciones entre modelos conceptual, observacional y probabilístico en el proceso de inferencia

en un modelo conceptual, y detectan patrones emergentes mediante ajuste automático de sus parámetros⁴⁵. Obviamente que esta distinción no es absoluta: existen métodos de IA que incorporan conocimiento previo, hipótesis conceptuales (como los modelos bayesianos), y causalidad⁴⁶. Pero finalmente solo buscan patrones en la nube de datos. Luego la IA prescinde de las teorías médicas para formular sus respuestas.

Modelo estadístico

Las observaciones emergentes del modelo observacional serán encuadradas en un marco probabilístico para representar y manejar la incertidumbre y para hacer predicciones propias del análisis de datos científicos, ML (*machine learning*) y la IA⁴⁷. Se asume que el PGD, previamente modelado en términos conceptuales, es el mecanismo real que genera las observaciones, es decir los valores medidos para las variables de interés en una muestra determinada, y ahora el mismo deberá modelarse en términos estadísticos. Es decir, probabilísticos. Por ejemplo, se modelará el rango más probable para la tensión arterial para cada grupo etario.

El modelo estadístico será una representación simplificada y necesariamente imperfecta del PGD, expresada en términos probabilísticos⁴⁸. El modelo observacional contiene variables aleatorias de interés que irán adquiriendo valores a medida que el PGD produce observaciones; a esos valores se les asignarán probabilidades de

ocurrencia, para determinar así la estructura probabilística real del proceso en estudio. ¿Cuál será la estatura más probable de este niño en Alemania? ¿Cuál es el riesgo de morir por neumonía? Son preguntas con un componente probabilístico⁴⁹ y forman parte de un modelo estadístico.

La construcción de modelos estadísticos comienza por su formulación o especificación, es decir la selección de una estructura probabilística teórica para los datos (distribución normal, modelo lineal, regresión de Cox, etc.). Este paso depende también del tipo de variable de respuesta; continua o categórica, tiempo hasta el evento, probabilidad de éxito en una serie de intentos, etc. Tras las observaciones sigue el ajuste o modificación de los parámetros del modelo (coeficientes de regresión, parámetros de forma en distribuciones, etc.) para minimizar la discrepancia entre predicciones del modelo y observaciones. Para minimizar el error de predicción, se buscará minimizar el MSE –*Minimal Squared Error*– en modelos lineales, o maximizar la verosimilitud –*likelihood*– en modelos paramétricos⁵⁰. Todo esto para validar (o ajustar) el modelo, es decir confirmar su capacidad predictiva y de generalización.

En esencia, los modelos estadísticos⁵¹ son idealizaciones matemáticas del PGD que proponen relaciones probabilísticas entre los valores obtenidos para las variables de interés y las variables aleatorias seleccionadas como resultado

u output, apelando a funciones estadísticas para su representación. Los modelos suelen combinar una parte determinista, ya que descansan en una función matemática que relaciona variables (por ejemplo, $Y = \beta_0 + \beta_1 X$), y una parte estocástica (término de error aleatorio, ϵ , que captura la incertidumbre, de manera que, en el ejemplo, el modelo sería: $Y = \beta_0 + \beta_1 X + \epsilon$). Las variables seleccionadas como predictoras suelen obedecer a teorías de causalidad; por ejemplo, se busca relacionar el peso corporal con la tensión arterial media, debido a teorías fisiopatológicas que vinculan ambos fenómenos. Esto no ocurre en la mayoría de los modelos de IA. Y definitivamente no ocurre cuando la medicina utiliza grandes modelos de lenguaje.

Variables aleatorias

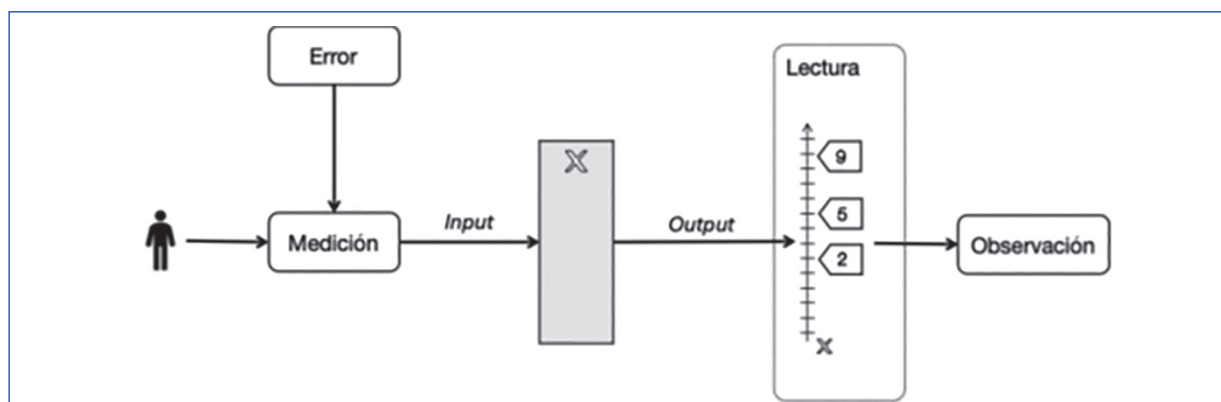
Los datos con que se alimenta el modelo estadístico son valores de ingreso (input) correspondientes a variables aleatorias, y que caracterizan atributos del PGD en estudio. La intensidad o frecuencia con que los atributos bajo análisis se expresan en el PGD que nos interesan (antes incluidos en nuestro modelo conceptual), deben convertirse en cantidades para poder hacer cálculos; se transforman así en valores de variables aleatorias (Fig. 4).

Una variable aleatoria⁵² es una función, similar a las funciones matemáticas. Una función es una regla matemática o lógica que asigna a cada elemento de un conjunto X exactamente un elemento de otro conjunto Y, y así los valores de

Y son función de los valores de X: $f(x) = y$. La variable aleatoria es una función que mapea cada elemento del espacio donde se mide el fenómeno de interés (técnicamente espacio muestral), a un valor numérico, una cantidad oposición en la escala de números reales. Ella puede mapear desde la existencia o no del atributo de interés, a una variable dicotómica (0,1), o desde la intensidad con que este se expresa el fenómeno de interés, a una variable continua (valor de $-\infty$ a $+\infty$); pero siempre resultará una cantidad de algo. El PGD se mide en términos de cantidades de alguna o muchas variables aleatorias cuyo valor depende de una medición, y cuyos valores se distribuyen según probabilidades de ocurrencia.

En la medición⁵³, un dispositivo traslada una señal física producida por el PGD a una escala numérica calibrada. Esto genera un valor que modelamos como la realización de una variable aleatoria de interés, y que incluye la señal emitida por el PGD, ruido del entorno, y error de medición. Así el dato observado es un valor particular que toma la variable aleatoria cuando se realiza la observación, más los errores que lo distorsionan. Dada la naturaleza estocástica del PGD, los datos fluctúan dentro de un rango posible, y su variabilidad se cuantifica como “varianza”, es decir como la dispersión de los valores alrededor de su media o proporción reales. Esta varianza manifiesta la estocasticidad inherente al PGD⁵⁴. La frecuencia con que se observa cada valor particular de la variable aleatoria dependerá de la distribución de probabilidades

Figura 4 | Esquemización del concepto de variable aleatoria



Una variable aleatoria es una función que asigna (o “mapea”) cada resultado del espacio muestral (el espacio de medición) a un valor en una escala numérica

asociada a ella, determinada por el PGD. Por ejemplo, en una distribución normal los valores cercanos a su media son más probables que los extremos. En los modelos de IA que incorporan incertidumbre o aprendizaje a partir de datos (la mayoría de los utilizados en salud), las variables de entrada y salida son representados como variables aleatorias, con su valor esperado y varianza. Esto es distinto al razonamiento clínico clásico, basado en mecanismos de acción para explicar la observación.

Muestreo

El modelo estadístico⁵⁵ específico condiciona a su vez la forma en que se obtienen los datos del proceso aleatorio de referencia⁵⁶, es decir que influye en el modelo observacional y viceversa, ya que ciertas limitaciones del proceso observacional podrían obligarnos a cambiar el modelo estadístico (dificultades de medición, rareza del fenómeno, etc.). En este contexto⁵⁷ resulta fundamental la muestra aleatoria.

La teoría del muestreo⁵⁸ establece que bajo ciertas condiciones probabilísticas un subconjunto finito de observaciones (muestra) seleccionado al azar de un universo definido, permite inferir propiedades de la población o del PGD, con un margen de error controlado. Hablamos de condiciones probabilísticas cuando las observaciones (siempre de una muestra) son a su vez el resultado de otro mecanismo no determinista, el muestreo aleatorio; el valor de una observación no influye en otra (son independientes); y las observaciones pertenecen a la misma distribución (los diferentes valores posibles de la variable de interés tienen la misma probabilidad de ser observados con cada muestra). Si la muestra es representativa (lo que depende también del tamaño) y aleatoria, los estadísticos derivados de ella (media, varianza, etc.) se aproximan a los parámetros poblacionales, siguiendo leyes probabilísticas como el teorema central del límite y el error estándar, etc. La idea de base es que estos conceptos de error estándar y varianza permiten modelar lo que ocurriría en muestras repetidas⁵⁹, que serán, obviamente, casi siempre hipotéticas.

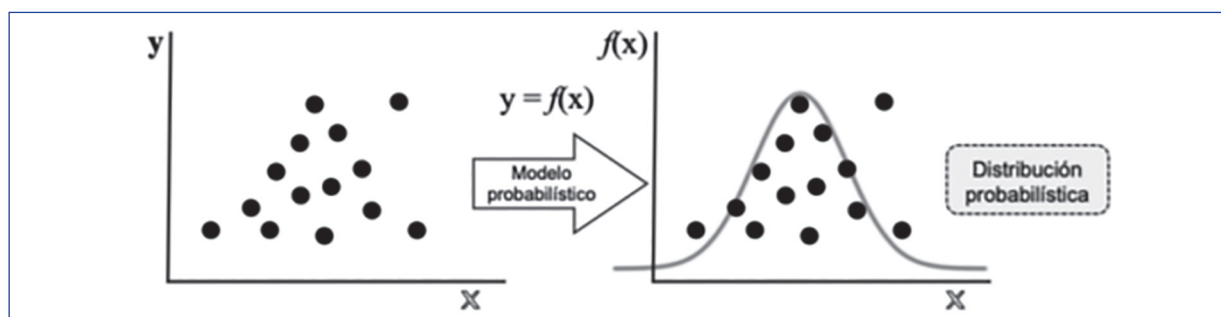
Respecto de los modelos de IA, a pesar de ser entrenados con millones de datos, igual se apoyan en los fundamentos de la teoría del muestreo⁶⁰ para garantizar que los datos del entrena-

miento reflejen la distribución real que produce el PGD de interés, y que sean generalizables. Ignorar los principios del muestreo en IA conduce a modelos sesgados, o poco generalizables, incluso ante datos masivos.

Distribuciones probabilísticas

El proceso observacional produce una varianza (cambios en los valores de una misma variable de interés entre medición y medición) que sigue patrones impuestos por el PGD subyacente. Dichos patrones o regularidades pueden modelarse mediante una distribución probabilística⁶¹ sobre el espacio de observaciones, la cual atribuye probabilidades de ocurrencia a los valores posibles de las variables aleatorias incluidas en el modelo. Por ejemplo, las cifras de tensión arterial cambian de un momento a otro en un mismo individuo y entre individuos, pero lo hacen siguiendo una regla probabilística. Una distribución probabilística es una función matemática que toma cada valor o rango de valores de una variable aleatoria, y le asigna una probabilidad de ocurrencia, creando así modelos probabilísticos (Normal, Poisson, Binomial, Exponencial, o funciones especiales ajustadas por algoritmos de ML, etc.) para cada variable, y también distribuciones conjuntas para más de una variable (Fig. 5). Se considera que las observaciones son extraídas de alguna de las distribuciones modeladas.

En el caso de analizar la asociación entre variables, se modelan las dependencias entre ellas, como ocurre con los modelos lineales o las clasificaciones. Por ejemplo, la frecuencia con la que se registran temperaturas corporales en una unidad de terapia intensiva tiende a seguir un patrón característico, modelable con una distribución probabilística, y que podría asociarse a otros valores de otras variables aleatorias, como ser recuento de neutrófilos en sangre, siguiendo algún modelo estadístico (Ej.: regresión lineal). Igualmente, los datos que ingresan en un modelo de IA se tratan como realizaciones de variables aleatorias provenientes de un proceso generador de datos (PGD) desconocido. A su vez, las salidas –ya sean valores puntuales o distribuciones explícitas– constituyen estimaciones sujetas a incertidumbre y poseen su propia distribución probabilística. Por ello, las conclusiones que ofrece la IA deben entenderse

Figura 5 | Esquematación de la creación de una distribución probabilística en un conjunto de datos

como aproximaciones probabilísticas a la verdad, nunca como certezas absolutas. También en un gran modelo de lenguaje, cada palabra que el sistema genera corresponde al valor esperado de una distribución probabilística de palabras posibles. Sin embargo, su interpretación se modifica y se complejiza cuando se le solicita, por ejemplo, un diagnóstico médico a partir de un conjunto de síntomas provisto como *prompt* (entrada o solicitud que se le da a un sistema de IA para que realice una tarea o genere una respuesta específica).

Estadística descriptiva

El modelo estadístico sugerirá también el modo como se describen las observaciones; en otras palabras, el resumen de los datos⁶². Mediante esta descripción el investigador, a través del modelo, infiere la estructura de la muestra y del PGD del cual ésta es extraída⁶³. La estructura de que hablamos es una abstracción matemática de la frecuencia con que se observan los valores de las variables de interés, condensada en una distribución probabilística, como hemos dicho. Describir una muestra mediante estadísticos descriptivos consiste en contar individuos, presentar cantidades de agregados (proporciones, medias, medidas de dispersión), y mapear las observaciones con su frecuencia de aparición (definir distribuciones probabilísticas).

Por ejemplo, al modelar los valores observados se obtiene una medida de tendencia central que describe el “centro” de la distribución probabilística de los valores observados para la variable de interés. Además, están las medidas de dispersión, como la varianza, que determinan cuánto se alejan los valores de ese centro, infor-

mando la variabilidad de la muestra. Estos son los primeros estadísticos.

Un estadístico⁶⁴ es una aproximación a un parámetro poblacional desconocido, calculado a partir de en una muestra aleatoria tomada de dicha población. Los estadísticos generalmente difieren de los verdaderos valores poblacionales.

Quedará luego analizar ahora los comportamientos de los valores de las diferentes variables de interés entre sí (Ej.: relación entre peso corporal y esperanza de vida), y en relación con espacios de oportunidad como tiempo (Ej.: pacientes ingresados por hora), espacio (Ej.: vacunados por hogar), u otro contexto (Ej.: fallecidos por familia). Esto se hace mediante la estadística analítica. Los modelos de IA pueden incluir, como variables de entrada, estadísticas descriptivas que resumen la muestra (por ejemplo, promedios o proporciones), y con ellas construir predictores complejos. No obstante, en la IA predictiva la estructura estadística del proceso generador de datos se aprende de forma automática durante la optimización de la función que minimiza el error de predicción, mientras que el experto define únicamente el espacio funcional –por ejemplo, la arquitectura de la red– sin especificar explícitamente la distribución de probabilidad de cada variable de entrada.

Estadística analítica

El modelo estadístico permite inferir, además de la descripción de la estructura probabilística del PGD, relaciones entre variables dentro del mismo (correlación, regresión, o modelos multivariados), así como comparar diferentes muestras para deducir el proceso aleatorio que las origina (pruebas de hipótesis, t-test, ANO-

VA, análisis de *clusters*). Aquí es donde la significación (*p*-valores, intervalos de confianza) y la asociación estadística (coeficientes de regresión, etc.) entran en juego. Imaginemos que un conjunto de variables se observa de manera simultánea y se detecta un patrón en la variación de sus valores dentro del “espacio de oportunidad” (tiempo, espacio experimental, geografía, o contexto) de la medición. Esto sugiere que existe algún tipo de coordinación, y que varían con cierta dependencia o asociación. Cuando los valores de un grupo de variables aleatorias muestran patrones que permiten predecir su comportamiento, esto se capta mediante modelos matemáticos, ya sean deterministas o estocásticos. De no hallarse tal coordinación, o de no detectarse, entonces las variaciones se atribuyen al azar, también definido como ruido. Las inferencias estadísticas de todo tipo emplean modelos estadísticos que incorporan supuestos teóricos. Las variables aleatorias, los intervalos de confianza y las probabilidades *a posteriori*⁶⁵, por ejemplo, existen en este mundo teórico. Y cuando usamos un modelo estadístico para realizar inferencias estadísticas, implícitamente afirmamos que la variabilidad mostrada por los datos es adecuadamente capturada por dicho modelo, por lo que el mundo teórico corresponde razonablemente bien con el mundo real⁶⁶. Estos datos exhiben tanto regularidades (a menudo descritas teóricamente como una “señal”, que en ocasiones responde a descripciones matemáticas simples o “leyes”), como variabilidad no explicada, que habitualmente se considera “ruido”. Los modelos estadísticos permiten, dependiendo del modelo, separar señal de ruido.

Para la estadística analítica resulta crítica la noción de probabilidad en su doble acepción epistemológica⁶⁷, como representación de la variabilidad de la realidad bajo evaluación (ej.: existe una “*p*” probabilidad de obtener A como resultado de un PGD cualquiera, $p(A)=x$, que depende de su frecuencia de aparición), y como medida de la incertidumbre de quien pretende conocer (adjudico subjetivamente una probabilidad *X* a que el resultado sea A). Y es en este momento que la dimensión epistemológica de las probabilidades nos lleva a la gran cuestión de la inferencia.

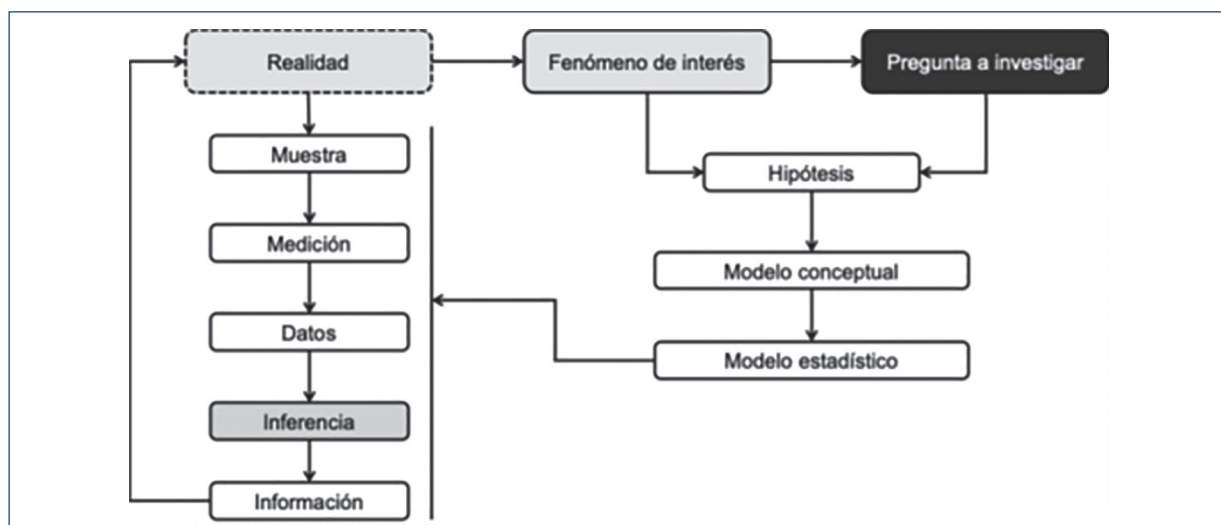
El proceso de inferencia en la ciencia

Inferencia estadística

Proceso observacional, modelo conceptual, modelo estadístico y muestreo, tienen como objetivo realizar inferencias acerca de la estructura probabilística subyacente al PGD⁴⁸. En sí misma la inferencia, en sus variantes de inductiva o deductiva, es un proceso cognitivo por el cual concluimos un resultado reflexionando sobre otra información⁶⁸. Podemos inferir los efectos a partir de ciertos principios (deducción) o, a la inversa, los principios causales a partir de ciertos efectos (inducción)⁶⁹. La inferencia en general puede ser lógica, estadística, causal, o la inferencia en el contexto del ML y la IA, que veremos en seguida.

En el caso de la inferencia estadística se concluye una estructura probabilística para el PGD a partir de los resultados obtenidos en un proceso observacional. Más precisamente, en la inferencia estadística⁷⁰, o bien desde las observaciones se infiere la estructura probabilística del PGD y las potenciales asociaciones entre variables, o bien se contrasta la distribución probabilística que genera las observaciones con otra distribución llamada hipótesis alternativa⁷¹, y se infiere el PGD que produjo nuestras observaciones (Fig. 6). La inferencia también identifica el grado de consistencia o verosimilitud (*likelihood*) entre las configuraciones del modelo estadístico seleccionado y los datos observados de acuerdo con la calidad de las predicciones que este produce⁷². Por ejemplo, si seleccionamos un modelo de distribución normal para representar los valores de estatura de los niños en la población, pero las observaciones no se agrupan simétricamente en torno a su media, diremos que el modelo no se ajusta a los datos, y requerimos otra distribución.

La inferencia estadística asume que los datos, como se ha dicho, resultan de un PGD aleatorio latente, cuya estructura probabilística resulta desconocida pero que es inferida a partir de observaciones (siempre muestrales)⁷³. Concretamente los datos nos permiten obtener una función que se ajusta a la distribución probabilística de los valores de la variable de interés, para realizar predicciones ante nuevos valores de entrada. Por ejemplo, dada la edad de un individuo podría inferir los probables valores de presión

Figura 6 | Esquematación del proceso de inferencia estadística

arterial media frente a una nueva medición. Si la muestra nos brindó cien valores de X que se comportan como una distribución normal, cada uno asociado a una probabilidad (no entraremos en el tecnicismo de la densidad de probabilidad), luego podremos conocer la probabilidad de ocurrencia de cualquier otro de los valores de X no observados, usando dicha función. En síntesis, el modelo sirve de base para la predicción empírica o inferencia estadística⁷⁴.

Inferencia en los modelos de inteligencia artificial

Ahora bien, las variables incluidas en el modelo estadístico clásico (los predictores en una regresión logística, por ejemplo) se consideran explícitamente relacionadas (explicativa o predictivamente) con la variable de resultado. Es decir, son modelos con un componente explicativo y un fundamento teórico que los respalda. Por contrapartida, muchos modelos de IA se centran en la predicción misma, y no pretenden establecer relaciones causales o explicativas entre variables (entre los valores observados para cada variable)⁷⁵. Los modelos de IA construyen funciones para ajustar sus estimaciones a los datos de entrenamiento y, posteriormente, realizar inferencias sobre nuevos datos sin asumir relaciones explicativas entre las variables de entrada⁴⁵. Durante el entrenamiento se expone el modelo a un conjunto seleccionado de datos

para que aprenda a reconocer patrones. Luego el modelo construye una función estadística (Ej.: RN) que prediga esos patrones. Posteriormente, ese modelo entrenado identifica dichos patrones en datos nuevos y realiza predicciones basadas en ellos. Así, la inferencia refleja que el algoritmo ha aprendido patrones a partir de datos conocidos y puede reconocerlos en datos no vistos previamente. Pero no requiere de un sustento teórico de tipo causal.

Podemos asumir aquí que el PGD es una caja negra donde, para ciertos valores de las variables de entrada X , se generan ciertos valores de la o las variables de salida Y . La estadística clásica asume un modelo (Ej.: regresión lineal) para X tal que resultará Y .

La IA, en cambio, asume el PGD complejo, luego busca una función que ajuste las relaciones entre X e Y simplemente buscando minimizar el error de predicción⁷⁶. El modelo de IA predecirá un valor de Y para cada nuevo valor de X y medirá su desvío (o error de predicción) basándose en los datos de entrenamiento, e irá cambiando sus parámetros durante el aprendizaje para ajustarse a la realidad. En su forma habitual de aprendizaje el modelo de ML recibe, durante su entrenamiento, pares (x, y) formados por los valores de entrada x (por ejemplo, un patrón radiográfico) y su salida o etiqueta correspondiente y (el diagnóstico). A partir de esos datos aprende una función predictora $f_0(x) = \hat{y}$ -parametrizada

por θ — que, para un x dado, produce una predicción $\hat{y}$. Tras calcular la diferencia entre la predicción y el valor real ($\hat{y}-y$), expresada mediante una función denominada de pérdida, y el algoritmo, ajusta los parámetros θ de la función predictora para minimizar ese error (Fig. 7). Una vez que la función de pérdida alcanza un mínimo aceptable, el modelo queda listo para estimar nuevos valores $\hat{y}$ cuando se le presenten entradas x no vistas; es decir, para hacer diagnóstico de neumonía cuando vuelva a “ver” el patrón radiográfico aprendido.

Parámetros

Conviene detenernos entonces en el concepto de parámetro. Los sistemas estudiados como PGD se caracterizan o especifican mediante parámetros⁵⁵. Un parámetro es una cantidad que determina o especifica algún aspecto fundamental de la distribución probabilística de los valores de las variables de interés en la población, o PGD bajo análisis. El parámetro es un valor (Ej.: media (μ), varianza (σ^2), proporción (p), pendiente de regresión (β), etc.) que determina cuantitativamente alguna característica esencial de la función que describe la distribución subyacente de los datos (Ej.: μ y σ^2 determinan la forma de una distribución normal), o sus asociaciones (Ej.: coeficiente de correlación ρ). El parámetro es siempre un objeto matemático que no existe en la realidad sino en el modelo que abstrae el

PGD. Se podría decir que el modelo existe en dos formas, una poblacional desconocida (por ejemplo, caracterizada por algún valor de μ y σ^2) y otra muestral conocida (por ejemplo, caracterizada por media y varianza muestrales, que son calculables). El propósito de la inferencia estadística es estimar o formular conjeturas del tipo y valor específico de dicho parámetro, a partir de la información contenida en una muestra extraída de la población o proceso en cuestión. Para dicha estimación se usarán los estadísticos (media muestral, desvío estándar, coeficiente de regresión, etc.).

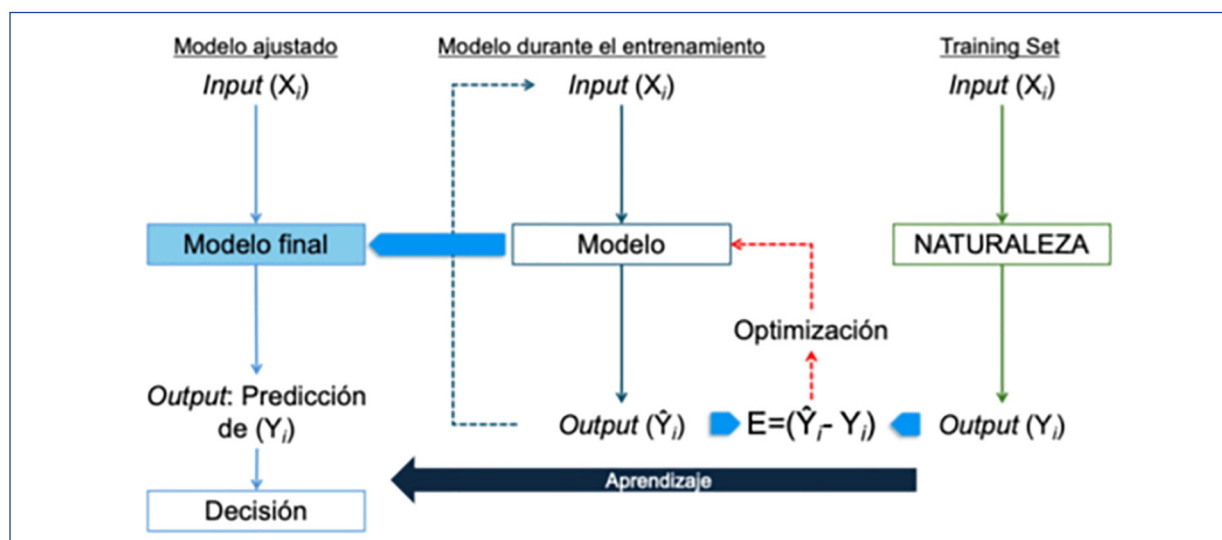
Dicho sintéticamente, el objetivo del análisis estadístico tradicional es inferir la distribución probabilística subyacente a las observaciones, estimar los parámetros que la definen (con sus respectivos valores) y asumir que la muestra procede de ese proceso generador de datos (PGD) modelado estadísticamente⁷⁷.

En la concepción de un modelo estadístico cualquiera, un paso esencial consiste en su parametrización⁷⁸.

Parametrización del modelo estadístico

Para recapitular lo discutido hasta aquí diremos que hablamos de evidencia experimental cuando una situación experimental, que es un PGD, se representa mediante un modelo matemático estadístico adecuado⁷⁹ que denominamos “E”. Denominamos “evidencia experimen-

Figura 7 | Esquema de entrenamiento de un modelo de aprendizaje automático



tal” al par modelo estadístico y resultados (E, x), donde E es el modelo estadístico que formaliza el PGD o experimento, incluyendo tanto la parametrización, definida por el espacio de parámetros Θ que caracteriza las distribuciones probabilísticas del PGD, como el espacio muestral (todos los posibles resultados de la variable X , tal que $X=x$), usando la función probabilística $\phi(x|\theta) - x$ (minúscula) es el valor de X (mayúscula) observado tras el muestreo. Bajo este planteamiento, las técnicas de inferencia estadística, como las pruebas de hipótesis y la estimación por intervalos, exploran la función probabilística $\phi(x|\theta)$ para distintos valores θ del espacio de parámetros Θ , lo que permite describir de forma cuantitativa y cualitativa la evidencia aportada por los datos.

Ahora bien, la parametrización no es igual en las estadísticas clásicas que en los modelos de IA⁸⁰. Primero, en estadística clásica, la parametrización está vinculada a supuestos teóricos sobre el PGD o la población en estudio. Por ejemplo, en estadística clásica podría asumirse la existencia de un parámetro llamado “pendiente”, para definir la relación lineal entre índice de masa corporal y mortalidad, con un respaldo teórico fisiopatológico en dicha relación. En IA⁵, en cambio, la parametrización es un mecanismo de aproximación funcional, donde los parámetros (denominados “pesos” en el caso de las redes neuronales) no tienen interpretación teórica directa (vale decir no son el valor más frecuente de una distribución normal o la pendiente de una recta de regresión), es decir que no corresponden a una distribución de los datos en el PGD, sino que emergen del entrenamiento del modelo con datos. En estadística, los parámetros son propiedades fijas (aunque desconocidas) de la población. Se asume que la población está caracterizada por el verdadero parámetro. En IA, los parámetros son variables de ajuste que no tienen significado fuera del modelo entrenado⁸¹. Esto vuelve al modelo de IA por naturaleza opaco.

Inferencia causal

Llegamos entonces a la cuestión de la inferencia causal, el punto central de las soluciones médicas explicables *versus* las inexplicables. Conceptualmente, una relación causal sería aquella donde el evento en la variable X (es decir

la ocurrencia de un valor determinado de X) es necesario o suficiente para que ocurra el evento $Y=y$, es decir para que ocurra un valor determinado en la variable de resultado Y ⁷⁷. Por ejemplo, dado que ocurre que un individuo alcanza una edad (en el ejemplo variable X) determinada, su riesgo de morir (en el ejemplo, variable Y), será de tanto. Se infiere entonces una relación causal entre X e Y . Como eventualmente para probar que un evento sea necesario o suficiente para la ocurrencia de otro evento se requiere evaluar el cambio en el sistema entre dos situaciones, una con el evento causal presente y otra con el evento causal ausente, transición en la cual se observará la ocurrencia de y , es que se dice que la asociación caracteriza condiciones estáticas (x e y se observan juntas), y la causalidad implica situaciones dinámicas (con y sin x , observar la ocurrencia o no, o la intensidad de presentación de y).

La estadística se encuentra profundamente involucrada en el análisis de asociación y en el planteo de modelos causales⁸², con la advertencia de que el análisis causal exige información teórica adicional al modelo estadístico del PGD, sin la cual la inferencia causal es imposible. Detrás de cualquier conclusión de causalidad debe existir algún supuesto teórico de causalidad, es decir alguna hipótesis de mecanismo productor del evento observado. Todas las ideas de correlación, dependencia estadística, regresión, verosimilitud, o cualquier idea de riesgo, todas ellas habitan exclusivamente en el universo de la asociación estadística. Para hablar de causalidad, en cambio, se requiere comparar, explicar influencias, medir efectos, despejar confundidores, evaluar intervenciones, y contar con explicaciones biológicas o mecánicas plausibles. De ningún grado de asociación estadística *per se*, podemos deducir nivel de causalidad alguno. Para inferir causalidad de una asociación entre variables, técnicamente de la distribución probabilística conjunta de ciertas variables, deben cumplirse otros postulados⁸³. Clásicamente se inferirá asociación causal cuando la asociación estadística entre variables es fuerte; es consistentemente advertida por múltiples experimentos; es específica para un efecto particular en una población determinada; la potencial causa precede temporalmente al efecto; existe una re-

lación dosis-respuesta entre exposición y efecto; se advierte un mecanismo biológico plausible de causalidad; la interpretación causal es coherente con el conocimiento existente; y existen relaciones causales análogas o similares establecidas previamente.

De todos modos, descubrir el vínculo de causalidad, hasta hoy es una facultad exclusiva de la mente humana. Los algoritmos llamados causales, por mayor sofisticación que tengan, encuentran límites que hasta ahora únicamente el ser humano puede traspasar⁸⁴. Todos los criterios de evaluación causal carecen de umbrales objetivos universalmente aplicables, lo que hace que la determinación de causalidad dependa en última instancia del juicio interpretativo humano que dice a partir de qué nivel, la asociación entre variables es suficientemente fuerte, consistente, y específica como para inferir causalidad. Es el juicio humano el que percibe causalidad en algunas asociaciones dosis-respuesta y no en otras.

Finalmente, la evidencia debe persuadir a un evaluador que integra múltiples dimensiones de información y que debe tomar una decisión en base a una relación de causalidad advertida⁸⁵. Estos criterios se han simplificado en el caso de la IA y el ML, y se podría hablar de causalidad cuando se observa⁸⁶ asociación entre variables, cuando advertimos los efectos de la intervención con una variable en el comportamiento de otra, y basándonos en la contrafactualidad de pensar en el comportamiento del sistema en estudio con y sin la variable potencialmente causal. Así y todo, se deberán establecer valores umbral en los algoritmos, y finalmente eso queda reservado al ser humano, que se persuade o no de la existencia de vínculo causal. De hecho, el método heurístico de razonamiento clínico discutido antes atestigua a favor de la peculiar capacidad de nuestra mente para advertir causalidad.

Por consiguiente, los modelos de IA, aunque implementen sofisticados marcos formales de causalidad⁸⁷ (gráficos acíclicos dirigidos, análisis contrafactual, etc.), operan inevitablemente dentro de limitaciones establecidas por umbrales, es decir puntos de corte a partir de los cuales se acepta causalidad, definidos por sus diseñadores humanos, patrones de inferencia causal presentes en sus datos de entrenamiento, y suposiciones causales incorporadas en su arquitectura.

Conclusión

La cuestión de la causalidad nos remite entonces al punto inicial de la interpretabilidad de las respuestas que un sistema de IA pueda llegar a brindar a los problemas de salud que se le planteen, y la necesidad de trabajar en la transparencia de los algoritmos que se integren al quehacer médico y científico⁸⁸. Estos sistemas pueden identificar asociaciones estadísticas e implementar reglas de inferencia causal, pero la validación definitiva de relaciones causales sigue requiriendo interpretación humana, es decir teorización humana⁷⁴, especialmente en dominios complejos como la medicina, las ciencias sociales o la política pública.

La capacidad predictiva de los algoritmos que motorizan los sistemas de IA es extraordinaria. Pero en definitiva consisten en abstracciones matemáticas de una realidad caracterizada por la incertidumbre y cuyos nexos de causalidad nunca podrán comprender, aunque sí representar. El desafío para la comunidad médica es incorporar a su conjunto de herramientas un nuevo tipo de maquinarias que simulan la inteligencia humana, no solo sin perder el control del proceso a través del cual se concluyen soluciones a problemas médicos, sino también, siendo parte en el enorme salto que significa tener una ayuda para pensar, sin precedentes. La realidad nos obligará a aceptar que habrá cuestiones sobre las cuales hoy se ocupa el razonamiento clínico que serán enteramente suplantadas por una máquina que lo hará mejor. Entramos en la era de la automatización de la inferencia. Algo parecido ocurrió con la calculadora y los calculistas. Pero las matemáticas teóricas no desaparecieron. Y esto fue porque los matemáticos teóricos son quienes mejor conocen los principios de automatización de la inferencia que ocurren en la calculadora, y ahora en la IA. La medicina debe revolucionar la forma en que piensa, incorporar matemáticas y computación avanzada en su formación, aceptar el desafío de ceder irremediablemente parte de sus dominios, a ciertos aparatos que infieren automáticamente⁸⁹, y ser parte activa de la interminable lista de mejoras que estas máquinas requerirán en el futuro. Todo esto será imposible si los médicos dejamos de entender lo que la computadora realmente hace.

Conflicto de intereses: Ninguno para declarar

Bibliografía

1. Olawade DB, Wada OJ, David-Olawade AC, Kunonga E, Abaire O, Ling J. Using artificial intelligence to improve public health: a narrative review. *Front Public Health* 2023; 11:1196397.
2. Stokes JM, Yang K, Swanson K, et al. A deep learning approach to antibiotic discovery. *Cell* 2020; 180: 688-702.e13.
3. Jordan MI, Mitchell TM. Machine learning: Trends, perspectives, and prospects. *Science* 2015; 349:255-60.
4. Krizhevsky A, Sutskever I, Hinton G. ImageNet classification with deep convolutional neural networks. *Neural Information Processing Systems* 2012; 25.
5. LeCun, Y., Bengio, Y. & Hinton, G. Deep learning. *Nature* 2015; 521: 436-44.
6. Brynjolfsson E, McAfee A. The business of artificial intelligence. *Harvard Business Review* 2017; 7: 3-11.
7. Wang H, Fu T, Du Y, et al. Scientific discovery in the age of artificial intelligence [published correction appears in *Nature*. 2023 Sep;621(7978):E33]. *Nature* 2023; 620: 47-60.
8. Otsuka J. Philosophical implications of deep learning. En: Otsuka J (ed). *Thinking about statistics*. New York: Routledge, 2023, p 129.
9. Topol EJ. High-performance medicine: the convergence of human and artificial intelligence. *Nat Med* 2019; 25: 44-56.
10. Lipton ZC. The mythos of model interpretability: In machine learning, the concept of interpretability is both important and slippery. *Queue* 2018; 16: 31-57.
11. Grote T, Berens P. On the ethics of algorithmic decision-making in healthcare. *J Med Ethics* 2020; 46: 205-11.
12. Freiesleben, T., König, G., Molnar, C. et al. Scientific Inference with Interpretable Machine Learning: Analyzing Models to Learn About Real-World Phenomena. *Minds & Machines* 2024; 34:32.
13. Bhise V, Rajan SS, Sittig DF, Morgan RO, Chaudhary P, Singh H. Defining and Measuring Diagnostic Uncertainty in Medicine: A Systematic Review. *J Gen Intern Med* 2018; 33: 103-15.
14. Witt EE, Onorato SE, Schwartzstein RM. Medical Students and the Drive for a Single Right Answer: Teaching Complexity and Uncertainty. *ATS Sch* 2021; 3: 27-37.
15. Corazza GR, Lenti MV, Howdle PD. Diagnostic reasoning in internal medicine: a practical reappraisal. *Intern Emerg Med* 2021; 16: 273-9.
16. Croskerry P. A universal model of diagnostic reasoning. *Acad Med* 2009; 84: 1022-8.
17. Marewski JN, Gigerenzer G. Heuristic decision making in medicine. *Dialogues Clin Neurosci* 2012; 14: 77-89.
18. Simon HA. Rational choice and the structure of the environment. *Psychol Rev* 1956; 63: 129-38.
19. Simon HA, Dantzig GB, Hogarth R, et al. Decision Making and Problem Solving. *Interfaces* 1987; 17: 11-31
20. Elstein AS, Schwartz A. Clinical problem solving and diagnostic decision making: selective review of the cognitive literature [published correction appears in *BMJ* 2006 ;333: 944. Schwarz, Alan [corrected to Schwartz, Alan]]. *BMJ* 2002; 324: 729-32.
21. Norman GR, Monteiro SD, Sherbino J, Ilgen JS, Schmidt HG, Mamede S. The causes of errors in clinical reasoning: cognitive biases, knowledge deficits, and dual process thinking. *Acad Med* 2017; 92: 23-30.
22. Newman-Toker DE, Nassery N, Schaffer AC, et al. Burden of serious harms from diagnostic error in the USA. *BMJ Qual Saf* 2024; 33: 109-20.
23. Letho MR, Nah FF, YI JS. Decision-making models, decision support and problem solving. En: *Handbook of Human Factors and Ergonomics*, 4th ed., 2012, p 192-242.
24. Windish DM, Huot SJ, Green ML. Medicine residents' understanding of the biostatistics and results in the medical literature. *JAMA* 2007; 298: 1010-22.
25. Cox DR. Statistical science: a grammar for research. *Eur J Epidemiol* 2017; 32: 465-71.
26. Schall JD. Decision making. *Curr Biol* 2005; 15: R9-R11.
27. Altman DG. Theoretical distributions. En: *Practical statistics for medical research*. London: Chapman & Hall, 1992, p. 48.
28. Nicolás de Cusa. *Un ignorante discurre acerca de la sabiduría*. No. 38, 1995. Buenos Aires: EUDEBA.
29. Cox DR, Efron B. Statistical thinking for 21st century scientists. *Sci Adv* 2017; 3: e1700768.
30. Chiang CL, Zelen M. What is biostatistics? *Biometrics* 1985; 41: 771-5.
31. Altman DG, Bland JM. Uncertainty and sampling error. *BMJ* 2014; 349: g7064
32. Heymsfield SB, Wadden TA. Mechanisms, Pathophysiology, and Management of Obesity. *N Engl J Med* 2017; 376: 254-66.

33. Puri PS. Biomedical Applications of Stochastic Processes. Mimeograph Series No. 370. Department of Statistics, Division of Mathematical Sciences, Purdue University, June 1974.
34. Aoki M. Data Generating Processes. En: State Space Modeling of Time Series. Serie Universitext. Berlín, Heidelberg: Springer-Verlag, 1990, p 8-20.
35. Gasser T, Muller HG, Kohler W, Molinari L, Prader A. Nonparametric Regression Analysis of Growth Curves. *Ann Statist* 1984; 12: 210-29.
36. Butler RJ, Butler MJ, Wilson BL. The Data-Generating Process and Scientific Inference. En: Advanced Statistics for Health Research. Capítulo 4, 2023. World Scientific Publishing Co. Pte. Ltd., p 65-82.
37. Iskandar R, Berns C. Markov cohort state-transition model: a multinomial distribution representation. *Med Decis Making* 2023; 43: 139-42.
38. Rasmussen SM, Molgaard J, Haahr-Raunkjaer C, Meyhoff CS, Aasvang E, Sorensen HBD. Forecasting of Continuous Vital Sign Using Multivariate Auto-Regressive Models. *Annu Int Conf IEEE Eng Med Biol Soc* 2022; 2022: 385-8.
39. Iannelli F, Koher A, Brockmann D, Hövel P, Sokolov IM. Effective distances for epidemics spreading on complex networks. *Phys Rev E* 2017; 95: 012313.
40. Coggon DI, Martyn CN. Time and chance: the stochastic nature of disease causation. *Lancet* 2005; 365: 1434-7.
41. Snow J. On the mode of propagation of cholera. *Medical Times* 1851, Nov 29:559-562. Presented at the London Epidemiological Society; May 5, 1851.
42. Betancourt M. Probabilistic Modeling and Statistical Inference. En: https://betanalpha.github.io/assets/case_studies/modeling_and_inference.html; consultado abril 2025.
43. Sogunuru GP, Kario K, Shin J, et al. Morning surge in blood pressure and blood pressure variability in Asia: Evidence and statement from the HOPE Asia Network. *J Clin Hypertens (Greenwich)* 2019; 21: 324-34.
44. Beam AL, Kohane IS. Big Data and Machine Learning in Health Care. *JAMA* 2018; 319: 1317-8.
45. Bzdok D, Altman N, Krzywinski M. Statistics versus machine learning. *Nat Methods* 2018; 15: 233-4.
46. Jiao L, Wang Y, Liu X, et al. Causal inference meets deep learning: a comprehensive survey. *Research* 2024; 7: 0467.
47. Ghahramani Z. Probabilistic machine learning and artificial intelligence. *Nature* 2015; 521: 452-9.
48. Box GEP. Science and statistics. *J Am Stat Assoc* 1976; 71: 791-9.
49. Cox DR. Role of formal theory of inference. In: Principles of Statistical Inference. Cambridge, UK: Cambridge University Press, 2006:3.
50. Wang S, Qian L, Carroll RJ. Generalized empirical likelihood methods for analyzing longitudinal data. *Biometrika* 2010; 97: 79-93.
51. Cox DR. Role of models in statistical analysis. *Stat Sci* 1990; 5: 169-74.
52. Evans MJ, Rosenthal JS. Chapter 2: Random variables and distributions. In: Probability and Statistics: The Science of Uncertainty. Toronto: University of Toronto, p 33.
53. Tukey JW. The future of data analysis. *Ann Math Stat* 1962; 33: 1-67.
54. Neyman J, Pearson ES. On the problem of the most efficient tests of statistical hypotheses. *Philos Trans R Soc A* 1933; 231: 289-337.
55. Fisher RA. On the mathematical foundations of theoretical statistics. *Philos Trans R Soc Lond A* 1922; 222: 309-68.
56. Gelman A, Vehtari A. What are the most important statistical ideas of the past 50 years? *J Am Stat Assoc* 2021; 116: 2087-97.
57. Tillé Y, Wilhelm M. Probability sampling designs: principles for choice of design and balancing. *Stat Sci* 2017; 32: 176-89.
58. Tillé Y, Wilhelm M. Probability sampling designs: principles for choice of design and balancing. *Stat Sci* 2017; 32: 176-89.
59. Cox DR. Foundations of statistical inference: the case for eclecticism. *J Stat Soc Aust* 1978; 20: 43-59.
60. Chetverikov SF, Arzamasov KM, Andreichenko AE, Novik VP, Bobrovskaya TM, Vladzimirsky AV. Approaches to sampling for quality control of artificial intelligence in biomedical research. *Sovrem Tekhnologii Med* 2023; 15: 19-25.
61. Coskun A, Oosterhuis WP. Statistical distributions commonly used in measurement uncertainty in laboratory medicine. *Biochem Med (Zagreb)* 2020; 30: 010101.
62. Barde MP, Barde PJ. What to use to express the variability: Standard deviation or standard error of mean? *Perspect Clin Res* 2012; 3: 113-6.
63. Spanos A. Where do statistical models come from? Revisiting the problem of specification. *Lect Notes Monogr Ser* 2006; 49: 98-119.
64. Anderson AA. Assessing statistical results: magnitude, precision, and model uncertainty. *Am Stat* 2019; 73(suppl 1): 118-21.
65. Van der Linden LR, Hias J, Walgraeve K, Flamaing J, Spriet I, Tournoy J. Introduction to Bayesian statis-

- tics: a practical framework for clinical pharmacists. *Eur J Hosp Pharm* 2019; 0: 1-5.
66. Kass RE. Statistical inference: the big picture. *Stat Sci* 2011; 26: 1-9.
 67. Reid N, Cox DR. On some principles of statistical inference. *Int Stat Rev* 2015; 83: 293-308.
 68. White AR. Inference. *Philosophical Quarterly* 1971; 21: 289-302.
 69. Neta R. What is an inference? *Philosophical Issues* 2013; 23: 388-407.
 70. Barnard GA. Statistical Inference. *J R Stat Soc Ser B* 1949; 11: 115-49.
 71. Lindley DV. Statistical inference. *J R Stat Soc Ser B Methodol* 1953; 15: 30-76.
 72. Savage LJ, Barnard G, Cornfield J, et al. On the foundations of statistical inference: discussion. *J Am Stat Assoc* 1962; 57: 307-26.
 73. Un modelo de error de media. En: Freedman D, Pisani R, Purves R, Adhikari A (eds). *Estadística*. Barcelona: Editorial Antoni Bosch, 1993, p 495.
 74. Cox DR. Causality: some statistical aspects. *J R Stat Soc A* 1992; 155: 291-301.
 75. Shmueli G. To explain or to predict? *Stat Sci* 2010; 25: 289-310.
 76. Breiman L. Statistical modeling: The two cultures. *Stat Sci* 2001; 16: 199-231.
 77. Pearl J. Causal inference in the health sciences: a conceptual introduction. *Health Serv Outcomes Res Method* 2001; 2: 189-220.
 78. Kass RE, Wasserman L. The selection of prior distributions by formal rules. *J Am Stat Assoc* 1996; 91: 1343-70.
 79. Birnbaum A. On the foundations of statistical inference. *J Am Stat Assoc* 1962; 57: 269-306.
 80. Ghahramani Z. Probabilistic machine learning and artificial intelligence. *Nature* 2015; 521: 452-9.
 81. Pearl J. The seven tools of causal inference, with reflections on machine learning. *Commun AC*. 2019; 62: 54-60.
 82. Holland PW. Statistics and causal inference. *J Am Stat Assoc* 1986; 81: 945-60.
 83. Hill AB. The environment and disease: association or causation? *Proc R Soc Med* 1965; 58: 295-300.
 84. Bareinboim E, Pearl J. Causal inference and the data-fusion problem. *Proc Natl Acad Sci USA* 2016; 113: 7345-52.
 85. Schenone P. Causality: a decision theoretic approach [preprint]. *arXiv*. Published December 18, 2018. arXiv:1812.07414 [econ.TH]. doi:10.48550/arXiv.1812.07414.
 86. Jiao L, Wang Y, Liu X, Li L, Liu F, Ma W, Guo Y, Chen P, Yang S, Hou B. Causal inference meets deep learning: a comprehensive survey. *Research* 2024; 7: 0467.
 87. Cavique L. Implications of causality in artificial intelligence. *Front Artif Intell* 2024; 7: 1439702.
 88. Chen IY, Joshi S, Ghassemi M, Ranganath R. Probabilistic Machine Learning for Healthcare. *Annu Rev Biomed Data Sci* 2021; 4: 393-415.
 89. Doan R, Levy A. La inteligencia en la era de su reproductibilidad técnica. *Le Grand Continent*. June 16, 2023. En: <https://legrandcontinent.eu/es/2023/06/16/la-inteligencia-en-la-era-de-su-reproductibilidad-tecnica/>; consultado abril 2025.