

Wilson A. Pérez-Oviedo

wperez@flacso.edu.ec

Facultad Latinoamericana de
Ciencias Sociales, FLACSO
(Quito – Ecuador)

ORCID: 0000-0002-9285-1501

**DE LA ECONOMETRÍA AL
RAZONAMIENTO CAUSAL:
INTELIGENCIA ARTIFICIAL Y
FORMACIÓN ECONÓMICA EN
TIEMPOS DE BIG DATA**

*FROM ECONOMETRICS
TO CAUSAL REASONING:
ARTIFICIAL INTELLIGENCE AND
ECONOMIC EDUCATION IN THE
AGE OF BIG DATA*

DOI:

<https://doi.org/10.37135/kai.03.16.10>

Recibido: 01/12/2025

Aceptado: 31/12/2025

Resumen

La rápida expansión de la inteligencia artificial (IA) y del Big Data transforma el mercado laboral y la investigación en economía, profundizando los cuestionamientos sobre la credibilidad de los resultados empíricos, el énfasis en la causalidad y la formación universitaria vigente. El artículo analiza estos cambios como parte de la credibility revolution en economía empírica, destacando sus aportes en diseño de investigación y sus límites frente a sistemas complejos y dinámicos. A partir de una revisión crítica de la literatura reciente, se argumenta que más datos y mayor capacidad computacional no garantizan mejores inferencias sin marcos conceptuales adecuados. Se propone una integración entre enfoques basados en diseño, modelos causales estructurales y aprendizaje automático, y se discuten sus implicaciones para la enseñanza de la economía y la empleabilidad en la era de la IA.

Palabras clave: Inteligencia artificial, Big Data, Inferencia causal, Credibilidad empírica, Econometría aplicada, Formación en economía

Abstract

The rapid expansion of artificial intelligence (AI) and Big Data is transforming both the labor market and economic research, intensifying concerns about the credibility of empirical results, the emphasis on causality, and current university training. This article examines these changes in the context of the credibility revolution in empirical economics, highlighting its contributions to research design as well as its limitations when addressing complex and dynamic systems. Based on a critical review of recent literature, it argues that larger datasets and greater computational power do not ensure better inference without appropriate conceptual frameworks. The article proposes a methodological integration of design-based approaches, structural causal models, and machine learning techniques, and discusses the implications of this convergence for economics education and future economists' employability in the AI era.

Keywords: Artificial intelligence, Big Data, Causal inference, Empirical credibility, Applied econometrics, Economics education.

DE LA ECONOMETRÍA AL RAZONAMIENTO CAUSAL: INTELIGENCIA ARTIFICIAL Y FORMACIÓN ECONÓMICA EN TIEMPOS DE BIG DATA

*FROM ECONOMETRICS
TO CAUSAL REASONING:
ARTIFICIAL INTELLIGENCE
AND ECONOMIC EDUCATION
IN THE AGE OF BIG DATA*

DOI:

<https://doi.org/10.37135/kai.03.16.10>

Introducción

La veloz difusión de la inteligencia artificial (IA) y de las tecnologías asociadas al Big Data está transformando de manera profunda el mercado de trabajo, los procesos productivos y, por lo tanto, las competencias requeridas para la inserción profesional de los jóvenes graduados universitarios. Aunque existe una considerable incertidumbre sobre la magnitud exacta de sus efectos agregados en la economía y la sociedad en su conjunto, la evidencia disponible parece mostrar que estos cambios ya están alterando de forma heterogénea las trayectorias laborales, afectando con particular intensidad a los trabajadores jóvenes y a los recién graduados universitarios.

En este contexto, las brechas entre expectativas sobre la educación universitaria, demandas del mercado laboral y transformaciones tecnológicas aparecen como un problema central para las universidades y, en particular, para disciplinas cuya identidad se ha construido históricamente sobre el dominio de herramientas cuantitativas, tal como la economía.

Este artículo quiere aportar a ese debate desde una perspectiva dual. Por un lado, analizamos cómo la expansión de la IA y del Big Data interactúa con problemas ya de vieja data en la economía, y en especial en la econometría, como son la credibilidad empírica en la investigación y la necesidad de mayores desarrollos metodológicos orientados al abordaje directo de la inferencia causal. Por otro, examinamos las implicaciones de estos cambios para la enseñanza de la economía, argumentando que el futuro laboral de los nuevos economistas depende cada vez más de su capacidad para articular herramientas cuantitativas, pensamiento causal y comprensión estructural de sistemas complejos. Partiendo de esta base, el trabajo propone que la adaptación de la formación económica a la nueva realidad tecnológica no es únicamente una cuestión técnica, sino un desafío epistemológico y pedagógico de primer orden.

Desempleo de graduados

Según muchos pensadores y empresarios involucrados en su creación y desarrollo, la Inteligencia Artificial (IA) podría ser una innovación tecnológica similar -al menos- al internet en cuanto a su impacto social. Como toda revolución tecnológica en proceso, la presente también tiene un alto contenido de incertidumbre: mientras Acemoglu (2024) estima su impacto en el crecimiento del PIB de Estados Unidos en un magro 0.09% a 0.12% anual (p. 35), Briggs y Kodnani (2023) calculan este impacto en un 1.5% anual.

Las proyecciones sobre el efecto en empleo tienen aún mayor incertidumbre, mejor ejemplificadas por las dramáticas declaraciones de Sam Altman, director de OpenAI, afirmó que “clases enteras de empleos desaparecerán” (Fortune, 2025); o de Dario Amodei, director

ejecutivo de Anthropic, “la IA podría eliminar la mitad de todos los empleos de oficina de nivel inicial y llevar el desempleo a cifras de dos dígitos” (CNN Business, 2025).

Por supuesto, es imposible predecir incluso el futuro cercano, pero ya en 2025 hay señales alarmantes. Los indicadores más recientes muestran que el desempleo de los recién graduados en Estados Unidos se mantiene por encima de los niveles generales del mercado laboral: en 2025, el Federal Reserve Bank of St. Louis (2025) reportó que los graduados jóvenes de 23 a 27 años tienen una tasa promedio de desempleo de 4.59 %, frente a alrededor de un 4% para la población económicamente activa en general. Este panorama lo confirma la Federal Reserve Bank of New York (2025) que estima que la tasa de desempleo para los recién graduados de la universidad asciende a 5.3 % en el segundo trimestre de 2025.

Por su parte Hosseini y Lichtinger (2025) muestran que, en las firmas que adoptan activamente tecnologías de IA generativa, el empleo de trabajadores junior (niveles de entrada) comienza a reducirse en forma relevante a partir del primer trimestre de 2023, rompiendo la trayectoria paralela que mantenían con firmas no adoptantes hasta el año 2022. Simultáneamente, el empleo de trabajadores senior continúa su crecimiento normal, sin mostrar la misma caída.

Es importante notar que este cambio se debe fundamentalmente a una reducción en la contratación de personal junior, más que a un aumento en las desvinculaciones o promociones. Esto parece señalar que la adopción de IA genera un patrón de cambio sesgado según la antigüedad y experiencia laboral (“seniority-biased”), erosionando los escalones inferiores de la carrera profesional mientras preserva o favorece los niveles más altos.

Hosseini y Lichtinger (2025), adicionalmente, muestran que la adopción de IA generativa por parte de las empresas se acelera muchísimo tras el lanzamiento de ChatGPT en noviembre de 2022: el número mensual de firmas que publican su primera vacante “GenAI integrator” (empleados a cargo del desarrollo y aplicación de IA en el flujo productivo de la empresa) prácticamente se multiplica por diez entre diciembre de 2022 y mediados de 2023, mientras que el acumulado crece de forma casi exponencial hasta superar las diez mil empresas en 2025.

Esta rápida difusión tecnológica es coherente con los patrones observados en el empleo: la expansión de la adopción presiona a las firmas a reorganizar tareas, reduciendo contrataciones *junior* y preservando los puestos *senior*, produciendo el sesgo por antigüedad documentado en el estudio. Si bien para el caso latinoamericano, a la fecha, no existen estudios tan detallados, se podría prever entornos profesionales similares para los jóvenes graduados de nuestra región.

En este sentido, diversos estudios recientes presentan evidencia de que una parte importante de los problemas de los jóvenes graduados en el mercado laboral se explica por un desajuste

entre las habilidades que adquieren en la universidad y las que requieren los empleadores —especialmente en competencias de IA, habilidades digitales avanzadas y “soft skills” complementarias (Segbenya *et al.*, 2023; Wut *et al.*, 2024; Portocarrero Ramos *et al.*, 2025; Kamaruddin *et al.*, 2023; OECD, 2023).

Por lo tanto, este deterioro relativo en la empleabilidad de los jóvenes graduados no puede interpretarse como un fenómeno pasajero o temporal del mercado laboral, sino fundamentalmente como un reflejo de desajustes en los procesos de formación universitaria. En el caso de la economía, cuya identidad como disciplina se ha construido históricamente sobre el dominio de herramientas y métodos cuantitativos, el desajuste entre las competencias que se enseñan y las que demandan los empleadores resulta especialmente significativo. La incorporación tan acelerada de tecnologías de IA y análisis de datos, no solo que redefine las tareas productivas, sino que redefine el conjunto de habilidades cuantitativas que valora el mercado laboral. En este contexto, la revisión crítica de la formación económica tradicional se vuelve un componente central para comprender —y por supuesto, corregir— los problemas de empleabilidad que ya se observan y pueden profundizarse en el futuro.

La enseñanza de la economía en la era de la IA y el Big Data

En este documento, buscamos aportar a la reflexión sobre la necesidad de cambios en el proceso de enseñanza aprendizaje universitario, enfocándonos en el caso de la formación en Economía y, específicamente, en los métodos cuantitativos de la economía. Consideramos que este es un buen punto de inicio porque la aplicación de métodos cuantitativos, así como del lenguaje y las herramientas matemáticas desarrolladas por la física en general, le ha permitido a la disciplina económica presentarse como una “verdadera” ciencia social e, incluso, ampliar la aplicación de sus métodos y herramientas a temas del ámbito de otras ciencias sociales (Pérez, 2024).

Si los métodos cuantitativos son una ventaja comparativa de la economía, entonces su adecuación a la nueva realidad tecnológica se vuelve central para la empleabilidad de los jóvenes economistas. También es muy importante señalar que, como mostramos más adelante en este documento, las herramientas cuantitativas creadas y desarrolladas en el ámbito del Big Data y la Inteligencia Artificial podrían originar un cambio profundo en el enfoque cuantitativo de la Economía y de las Ciencias Sociales en general.

El estado de la econometría a la llegada de la IA

El predominio de los métodos cuantitativos en la disciplina ha sido una fuente de sano orgullo para los economistas. Sin embargo, varios autores han levantado críticas sólidas sobre la validez de los estudios econométricos que se publican en las revistas académicas especializadas.

Podemos tomar como punto de referencia de esta discusión la publicación del artículo "The Power of Bias in Economics Research" por Ioannidis, Stanley y Doucouliagos (2017), que tuvo un impacto considerable en la disciplina económica, siendo citado más de 800 veces desde su publicación, convirtiéndose rápidamente en una referencia central en debates sobre lo que se ha calificado como una crisis de replicabilidad en economía (Ferraro y Shukla, 2020).

El estudio de Ioannidis, Stanley y Doucouliagos (2017) usa meta-análisis para evaluar la credibilidad de la investigación empírica en economía. Analizando 159 temas económicos, más de 6 700 estudios publicados en reputadas revistas académicas y más de 64 000 estimaciones, los autores examinan dos dimensiones fundamentales: el poder estadístico y la magnitud del sesgo en los resultados publicados.

Los hallazgos revelan un panorama preocupante: la potencia mediana de los estudios económicos oscila entre 10% y 18%, y solo alrededor del 10% de las estimaciones alcanza un poder adecuado. En numerosas áreas no existe ningún estudio suficientemente potente, y cerca del 80% de los efectos reportados están inflados o exagerados, en ocasiones por factores de dos y hasta de cuatro. Esta combinación de bajo poder e inflación sistemática de resultados implica que muchas estimaciones "significativas" pueden ser muy engañosas, cuestionando la solidez del conocimiento empírico en economía.

A raíz de la crisis de credibilidad empírica mencionada, las prácticas de investigación econométrica sostenidamente se han movido hacia mejores estándares de transparencia y rigor. Este cambio es visible, especialmente, en economía experimental, donde la adopción del pre-registro y de planes de análisis previo (PAP) se han hecho más frecuentes, gracias también a las plataformas disponibles como el AEA RCT Registry y el Open Science Framework, así como a las discusiones metodológicas sobre reproducibilidad (Page, Noussair & Slonim 2021). El pre-registro tiende a mejorar la credibilidad empírica cuando se acompaña de un PAP detallado, reduciendo la flexibilidad analítica y los riesgos de p-hacking, según muestra la evidencia reciente (Strømmland 2019).

Por otro lado, el uso del meta-análisis ha crecido como herramienta estándar para sintetizar evidencia y, también, reevaluar efectos empíricos en cada vez áreas más numerosas de la disciplina económica (Gechert *et al.* 2023). Sin embargo, este proceso de adopción continúa siendo heterogéneo: mientras la economía experimental ha avanzado de manera más rápida en implementar estas prácticas, otros subcampos —especialmente aquellos basados en datos observacionales— han incorporado estas herramientas de forma más gradual, tal como se reconoce en la literatura sobre transparencia y replicabilidad en economía (Christensen & Miguel 2018).

Sin embargo, el análisis de Bartoš *et al.* (2024), basado en más de 68 000 meta-análisis que abarcan medicina, psicología, ciencias ambientales y economía, presenta evidencia convincente de que esta última se encuentra entre las áreas más afectadas por sesgo de selección de publicación. En economía, si ajustamos por dicho sesgo, la probabilidad mediana de un efecto positivo cae de 99,9 % a 29,7 % y también disminuye el tamaño de efecto mediano de $d = 0,20$ a $d = 0,07$. Esto indica claramente una sobrestimación sistemática de los resultados que se publican académicamente, incluso en revistas de mucho prestigio.

Esto aporta aún mayor evidencia de la fragilidad estadística en ciencias empíricas y, adicionalmente, muestra que las editoriales continúan filtrando resultados nulos o modestos que, por ser de tales características, nunca llegan a ser impresos, a pesar de que pueden ser estimaciones válidas. El panorama resultante es dual: la disciplina avanza hacia mayor transparencia mediante pre-registro, meta-análisis y requisitos de disponibilidad de datos, pero enfrenta aún estructuras de publicación que limitan la credibilidad y replicabilidad de parte importante de su evidencia empírica.

Buscando establecer las posibles causas de esta crisis de credibilidad, podemos referirnos a un artículo incluso anterior, del mismo Ioannidis (2005). En este artículo, Ioannidis argumenta que el problema de falsedad posible de los resultados publicados se presenta sobre todo en campos académicos que tienen: tamaños de muestra pequeños; efectos reales pequeños; y/o, alta flexibilidad analítica. Aunque el artículo se centra en investigación biomédica, su diagnóstico general aplica a ciencias sociales y economía, en donde también se observan las características que el autor menciona como posibles causas del fenómeno que analizamos.

Llega el BIG DATA

El volumen global de datos digitales continúa expandiéndose a un ritmo impresionante: las proyecciones más recientes de la IDC (International Data Corporation) dicen que la Global DataSphere alcanzará aproximadamente 291 zettabytes ($2.9E23$ bytes) en 2027, gracias a la aceleración del cómputo en la nube, la expansión del Internet de las Cosas, y el crecimiento de sistemas basados en inteligencia artificial (IDC, 2024).

Esta cifra es el resultado, no solo el aumento en la generación de datos, sino también de la captura, replicación y procesamiento de estos datos a través de infraestructuras digitales distribuidas (IDC, 2024). Estas tendencias muestran claramente que el ecosistema global de datos se encuentra en una fase de expansión exponencial lo cual tiene implicaciones profundas para la investigación científica y la producción académica, también en ciencias sociales.

Para empezar, si una de las causas del bajo poder estadístico de las estimaciones en las publicaciones académicas es la escases de datos: ¿La solución estaría llegando con el Big

Data? En parte, sí: por supuesto que más datos disponibles significa mayores posibilidades de análisis y la inclusión de nuevas variables relevantes. Sin embargo, como lo mostramos más adelante, sin un marco metodológico adecuado, el Big Data puede amplificar —no corregir— los problemas diagnosticados por la crisis de replicabilidad.

Más aún, la expansión del Big Data sin duda cambia cualitativamente la forma en que podemos estudiar los fenómenos sociales. La posibilidad de trabajar con variables cada vez más desagregadas —por individuos, hogares, firmas, territorios, vínculos en redes, trayectorias temporales o patrones de interacción— posibilitará una comprensión más fina de la relación entre micro-comportamientos y los resultados agregados. Esto permitirá abordar con mayor profundidad lo que D. S. Wilson y E. O. Wilson (2007) identifican como «el problema fundamental de la vida social», esto es: la tensión entre comportamientos orientados al beneficio propio dentro de los grupos y la superioridad adaptativa de aquellos grupos formados por cooperadores (p. 328).

Desde la perspectiva económica, disponer de datos altamente desagregados permitirá examinar mucho mejor cómo las decisiones individuales generan retroalimentaciones sistémicas, cómo emergen o no patrones cooperativos, y cómo estas dinámicas se agregan en fenómenos macroeconómicos, muchas veces relativamente beneficiosos para los individuos, pero perjudiciales para los grupos. En suma, el Big Data no solo amplía la escala y granularidad de las observaciones, sino que habilita una investigación más fina de la vida social, cuya naturaleza es profundamente interrelacional

Por ejemplo, en el campo de la macroeconomía, la literatura reciente sobre heterogeneidad y redes productivas, especialmente el trabajo de Baqaee y Farhi (2019a, 2019b, 2019c, 2021), muestra que los agregados macroeconómicos pueden responder de manera no lineal a shocks microeconómicos debido a heterogeneidades, cuellos de botella, complementariedades locales y estructuras de red. Insistamos en el hecho de que la observación y medida de este tipo de características es cada vez más factible, gracias al Big Data.

Estos hallazgos empíricos y teóricos refuerzan la idea de que una economía no puede comprenderse plenamente sin entender la interacción multinivel entre agentes, sectores y niveles de organización. Desde esta perspectiva, la expansión del Big Data hace posible el estudio, de manera más directa y rigurosa, de los patrones individuales con los fenómenos agregados.

Sin embargo, el creciente uso del Big Data en la investigación académica también puede acarrear riesgos institucionales, que pueden comprometer la independencia y la autonomía de los investigadores sociales. En primer lugar, la asimetría de capacidades: se posicionan

para dominar la generación de conocimiento aquellas instituciones con más recursos —infraestructura tecnológica, personal capacitado, fondos—, lo que puede amplificar las brechas entre instituciones desarrolladas y otras con menos recursos, pudiendo ser perjudicadas particularmente las instituciones académicas del Sur Global (Couldry & Mejías, 2019; Kitchin, 2021).

Adicionalmente, la privacidad, gobernanza y control de datos se vuelven un problema crítico: la recolección masiva de datos individuales requiere marcos éticos y normativos robustos. Por supuesto, cuando estos marcos faltan, hay riesgo de vulneraciones, inequidad de acceso o usos indebidos de la información (Taylor *et al.*, 2020). Un riesgo adicional es la dependencia de tecnologías y plataformas privadas: gran parte de la infraestructura para almacenar y procesar Big Data está en las manos de grandes corporaciones, lo que podría comprometer la transparencia, la soberanía de datos y la reproducibilidad de estudios independientes (Couldry & Mejías, 2019).

Finalmente, tomemos en cuenta que la presión institucional por resultados “impactantes” puede incentivar un uso indiscriminado de Big Data con fines de productividad científica, priorizando cantidad sobre calidad y robustez, lo que puede erosionar estándares éticos y de reproducibilidad (Cukier & Mayer-Schönberger, 2017; Kitchin, 2021).

Por otro lado, desde una perspectiva metodológica, el uso de Big Data no elimina —y en algunos casos podría agravar— los problemas clásicos de validez, inferencia causal y sesgo en investigación social. Uno de los riesgos más importantes es la aparición de correlaciones espurias derivadas simplemente del volumen de datos, es decir, encontrar relaciones estadísticamente significativas que carecen de sustancia causal real (Ioannidis, 2020).

También, la facilidad para realizar múltiples análisis con grandes bases de datos incrementa el riesgo de “data dredging” o “p-hacking” (van der Zee & Reichardt, 2021). Adicionalmente, los datos masivos muchas veces carecen de información sobre factores sociales, históricos o culturales que son esenciales para interpretar resultados de forma adecuada, lo que puede llevar a malas interpretaciones de los efectos estimados (Ribeiro *et al.*, 2022). Un ejemplo claro de lo dicho es la conocida paradoja de Simpson, que se puede volver más frecuente con el Big Data, donde relaciones que parecen claras a nivel agregado pueden invertirse cuando se examinan subgrupos (creando una gran oportunidad para que el “efecto estadístico” sea seleccionado a conveniencia política). Es decir, “más datos” no significa “mejor inferencia” sin un buen diseño (Pearl, 2019; Wagner & Cooper, 2023).

La Escalera de la Causalidad de Judea Pearl y sus implicaciones para la Econometría y la IA

Judea Pearl (2019) ha propuesto que toda inferencia científica se organiza en tres niveles jerárquicos, la denominada Escalera de la Causalidad, que ordena los tipos de preguntas que pueden responderse con datos y modelos. El primer nivel, correlacional, se ocupa de patrones en los datos y relaciones del tipo probabilidad condicional $P(Y|X)$; en este nivel operan la mayoría de las técnicas estadísticas tradicionales. El segundo nivel, de intervención, introduce el operador $do(\cdot)$, permite responder preguntas del tipo ¿qué pasa si intervenimos en X? ¿qué pasa si *hacemos* X, como una intervención exógena, no originada dentro del mismo sistema?, es decir, se trata de calcular $P(Y|do(X))$.

El tercer nivel, contrafactual, responde preguntas de carácter subjuntivo o contrafactual —¿habría ocurrido Y si X no hubiera ocurrido? — para lo que se requiere comparar mundos posibles estructuralmente anclados en un modelo causal explícito (Pearl, 2019). En *The Book of Why*, Pearl y Mackenzie (2018) explican que estos niveles no son intercambiables: cada escalón permite responder un conjunto de preguntas estrictamente más amplio que el escalón anterior, y ningún volumen de datos correlacionales, por grande que sea, permite escalar sin un modelo causal estructural.

La mayor parte del aprendizaje de máquina contemporáneo, incluidos los modelos de aprendizaje profundo (*Deep Learning*) y los grandes modelos de lenguaje (*Large Language Models*, LLM), opera fundamentalmente en el primer escalón de la escalera: aprende patrones correlacionales en los datos de la forma más eficiente y exhaustiva posible (Schölkopf *et al.*, 2021), tal como lo explicamos más abajo. Pearl (2019) es explícito en señalar que el éxito predictivo de estos modelos no implica capacidad de explicación o análisis causal, pues no permiten razonar sobre intervenciones de política (por ejemplo), cambios de entorno que afecten la dinámica del sistema o escenarios contrafactuales.

Esta limitación es crítica para disciplinas como economía o ciencias sociales, donde las preguntas relevantes corresponden al segundo y tercer escalón: ¿qué ocurre si se modifica una política?, ¿cómo cambia un mecanismo bajo un nuevo régimen de incentivos?, ¿qué habría pasado si no se aplicaba un subsidio? Peters, Janzing y Schölkopf (2017) muestran que incluso algoritmos muy sofisticados fallan fuertemente en cuanto a su invariancia bajo cambios de distribución, generalización fuera de dominio y traslado entre entornos, precisamente porque carecen de representación causal.

La conclusión, compartida también por Schölkopf *et al.* (2021), es que, sin un marco causal explícito, Big Data no puede responder las preguntas que importan para políticas públicas,

evaluación de intervenciones o tratamientos y teoría económica, porque esas preguntas viven, por definición, en los escalones superiores de la escalera causal. En resumen: la aplicación de la IA no puede restringirse al nivel 1 de Pearl, sino que debe ampliarse a los otros dos escalones, como efectivamente se está haciendo.

Para Pearl (2019), el instrumento válido para representar estas relaciones causales es el DAG o Gráficos Dirigidos Acíclicos (*Directed Acyclic Graph*) como el dibujado en la figura 1, en donde los nodos representan las variables aleatorias y los arcos son aristas dirigidas en la dirección de la causalidad. El carácter acíclico de estos gráficos significa que no hay ninguna variable en donde se pueda empezar una cadena de causalidades que termine en la misma variable. Esto, por supuesto, impone una restricción fuerte cuando se quiere modelar sistemas dinámicos complejos (macroeconomía con retroalimentaciones, coevolución institucional, difusión en redes, mercados con expectativas, etc.) ya que por construcción, un DAG no puede representar bucles de realimentación contemporáneos (*feedback loops*), de modo que: (i) obliga a “romper” la simultaneidad mediante supuestos de orden causal, (ii) empuja a “desplegar el tiempo” (modelar $Y_t \rightarrow X_t \rightarrow Y_{t+1}$) para capturar retroalimentación solamente vía rezagos, y (iii) dificulta formalizar mecanismos como equilibrios (fijos) o dinámicas endógenas donde variables se determinan mutuamente en el mismo horizonte.

Esta limitación es reconocida en la teoría causal: al permitir ciclos, varias propiedades convenientes de los modelos acíclicos dejan de estar garantizadas (existencia/unicidad de soluciones, correspondencia clara entre grafo y distribuciones observacionales/intervencionales/ contrafactuales), lo cual ha obligado a introducir condiciones adicionales para sostener semánticas causales bien definidas (Bongers *et al.*, 2021). En los últimos años, varios autores (Bongers *et al.*, 2021), (Forré *et al.*, 2022), (Mooij *et al.*, 2024), (Ghassami *et al.*, 2024) han reconocido un punto metodológico central: la aciclicidad facilita inferencia e identificación en muchos problemas, pero para sistemas con presencia fuerte de retroalimentación la investigación contemporánea requiere desarrollar herramientas que unifiquen los DAG acíclicos y la dinámica con ciclos— para no perder precisamente lo que hace “complejos” a esos sistemas.

La revolución de la credibilidad

La expansión masiva de datos, las tensiones metodológicas y los cuestionamientos a la credibilidad empírica, constituyen el contexto donde se propone la llamada *credibility revolution* en economía empírica, formulada de manera influyente por Angrist y Pischke (2010). Esta revolución no surge como respuesta o aprovechamiento del Big Data, sino como una reacción a la fragilidad de las identificaciones de gran parte de la econometría aplicada tradicional, señaladas más arriba. Probablemente, su rasgo más distintivo es el énfasis en el diseño de la investigación, siempre privilegiando estrategias que permitan aislar variación causal creíble mediante experimentos aleatorizados o cuasi-experimentales.

En este sentido, la *credibility revolution* puede entenderse como un intento de la disciplina económica por fortalecer la inferencia empírica, restringiendo el conjunto de supuestos necesarios para identificar efectos causales y haciendo explícita y debatible la fuente de variación utilizada. De esta manera, la econometría profundiza su aproximación a la causalidad.

El impacto de este giro metodológico en la profesión ha sido muy importante. Según Athey e Imbens (2017), la economía aplicada contemporánea ha incorporado sistemáticamente diseños tales como: regresiones discontinuas, diferencias-en-diferencias, variables instrumentales y métodos afines, además de más transparentes diagnósticos y pruebas de robustez. De esta forma, el enfoque ha contribuido al incremento de los estándares de evaluación de políticas públicas y ha redefinido con mayor claridad qué se considera evidencia empírica persuasiva dentro de la disciplina económica, y en general en las ciencias sociales. Sin embargo, el éxito de la *credibility revolution* se concentra principalmente en la mejora de la validez interna, es decir, en la capacidad de establecer relaciones causales locales bajo supuestos explícitos, más que en la construcción de explicaciones generales o mecanismos teóricos subyacentes.

Precisamente allí están varias de las críticas más importantes a esta revolución. Deaton (2010) y Deaton y Cartwright (2018) advierten que la identificación causal, aun cuando sea difícilmente cuestionable desde el punto de vista del diseño, no garantiza una real comprensión de los mecanismos que generan los efectos observados; más importante aún, desde el punto de vista de posible aplicación de resultados causales en políticas públicas, tampoco garantiza su extrapolación a otros contextos institucionales o temporales.

Mientras, Heckman y Urzúa (2010) destacan que muchos estimadores privilegiados por la *credibility revolution* identifican adecuadamente parámetros causales locales, en general relevantes, pero que no siempre coinciden con los objetos de interés económico o de política pública, recordando que la elección del método debe estar subordinada a la pregunta sustantiva que se desea responder. Por supuesto, estas críticas no niegan ni minimizan los avances logrados, pero sí señalan límites estructurales del enfoque cuando se enfrenta a problemas complejos, dinámicos o multiescalares.

La irrupción del Big Data y de la IA contribuye nuevos e importantes elementos a este debate. Por un lado, la disponibilidad de grandes volúmenes de datos amplifica los riesgos ya identificados de correlaciones espurias, p-hacking y sobreinterpretación estadística; por otro, introduce herramientas computacionales capaces de manejar alta dimensionalidad, heterogeneidad y no linealidades. En este escenario, la *credibility revolution* enfrenta el desafío de adaptarse a entornos donde el diseño experimental estricto no siempre es factible y donde la inferencia depende en forma creciente de la combinación de múltiples fuentes de datos imperfectos (como suelen ser los datos que provienen del registro digital de la actividad humana diaria).

Este contexto abre espacio para enfoques causales alternativos y complementarios que permitan razonar explícitamente sobre supuestos, mecanismos e identificación en sistemas complejos, cuestión que conduce naturalmente a la discusión del marco de causalidad estructural desarrollado por Judea Pearl y colaboradores, que abordamos en la sección siguiente.

De la revolución de la credibilidad a la integración causal en la era de la IA y el Big Data

La discusión que precede muestra que la *credibility revolution* representó un avance importante y sustantivo en la disciplina económica al elevar los estándares de identificación causal y reducir la dependencia de supuestos estadísticos poco claros. Sin embargo, también revela que este giro metodológico, aun siendo necesario, resulta insuficiente para abordar los problemas empíricos de nuestros días que están caracterizados por: sistemas complejos, datos heterogéneos y cambios estructurales rápidos.

En particular, cuando la evidencia disponible proviene de múltiples fuentes imperfectas, poblaciones distintas o regímenes institucionales no comparables, la credibilidad de los estudios empíricos ya no depende únicamente del diseño de un estudio aislado, sino de la capacidad para integrar información causal dispersa bajo supuestos sólidos y explícitos. Esta limitación ha sido reconocida tanto por críticos de los enfoques puramente experimentales como por autores que buscan ampliar el alcance de la inferencia causal más allá de contextos locales (Deaton & Cartwright, 2018; Heckman & Urzúa, 2010).

En este punto, la convergencia entre la *credibility revolution* y los modelos causales estructurales adquiere gran relevancia metodológica. Mientras los diseños cuasi-experimentales restringen el conjunto de inferencias posibles mediante la selección de variación exógena, los modelos causales estructurales permiten representar explícitamente los mecanismos generadores de los datos y derivar, de manera formal, qué efectos son identificables dadas ciertas estructuras causales. Esta explicitación es particularmente relevante y válida en escenarios donde el ideal experimental es inalcanzable o donde múltiples fuentes de sesgo interactúan simultáneamente (es decir, en la vida real).

La literatura reciente -al parecer- avanza en esta dirección, precisamente, proponiendo marcos formales para la fusión de datos causales (*data fusion*) y la transportabilidad de resultados entre contextos distintos. Bareinboim y Pearl (2016) muestran que, bajo supuestos causales que estén especificados claramente, es posible combinar evidencia observacional y experimental proveniente de diferentes poblaciones para responder preguntas causales que ningún conjunto de datos aislado permitiría resolver.

Más recientemente, Hünermund y Bareinboim (2023) han traducido estas ideas al lenguaje de la econometría aplicada, mostrando cómo el *do-calculus* puede utilizarse para evaluar

identificabilidad y sesgos en diseños empíricos del mundo real, que se caracteriza por selección no aleatoria o heterogeneidad que no ha sido observada. En este sentido, el enfoque causal estructural no reemplaza la *credibility revolution*, sino que la extiende, proporcionando herramientas para razonar formalmente sobre supuestos, extrapolación y combinación de evidencia.

La irrupción tanto de la inteligencia artificial como del Big Data, sin duda, refuerzan la urgencia – y posibilidad- de esta integración metodológica. Si bien las técnicas de aprendizaje automático han demostrado una capacidad impresionante para modelar patrones complejos y manejar alta dimensionalidad, la literatura académica coincide en que su éxito predictivo no se traduce automáticamente en capacidad explicativa o causal (Athey & Imbens, 2019; Pearl, 2019).

En respuesta a esta dificultad han surgido varios enfoques híbridos que buscan preservar la inferencia causal en entornos de datos masivos, tales como el *double/debiased machine learning*, que combina estimadores flexibles con garantías asintóticas para parámetros causales de interés (Chernozhukov *et al.*, 2018). Adicionalmente, trabajos en *causal representation learning* buscan la robustez frente a cambios de distribución, así como la generalización fuera de dominio y la evaluación de intervenciones, objetivos que requieren incorporar estructura causal explícita en los modelos de aprendizaje (Peters, Janzing & Schölkopf, 2017; Schölkopf *et al.*, 2021).

En conjunto, estos desarrollos parecen evolucionar hacia una segunda etapa en la evolución metodológica de la economía empírica, en la que la credibilidad ya no se define únicamente por el diseño de estudios individuales, sino por la coherencia causal entre múltiples fuentes de datos, modelos y contextos. En la era de la IA y el Big Data, la pregunta central deja de ser únicamente *si* un efecto es identificable bajo un diseño específico, y pasa a ser *bajo qué supuestos estructurales* dicho efecto puede generalizarse, combinarse con otra evidencia o utilizarse para evaluar contrafactuales relevantes para la política pública. Esta transición tiene implicaciones directas no solo para la práctica empírica y la formulación de políticas, sino también para la forma en que se enseñan los métodos cuantitativos en economía, cuestión que abordamos en la sección final.

Conclusiones

El análisis desarrollado en el presente documento sugiere que la irrupción de la IA y del Big Data no constituye un fenómeno aislado ni meramente tecnológico, sino un shock transversal que amplifica tensiones y problemas ya existentes en la investigación empírica y en la formación universitaria en economía. La evidencia revisada muestra que, aunque el aumento exponencial

en la disponibilidad de datos y en la capacidad computacional abre oportunidades inéditas para el análisis social, también exacerba riesgos asociados a correlaciones espurias, sesgos de publicación y fragilidad inferencial. En este contexto, la llamada *credibility revolution* representó sin duda un avance decisivo al elevar los estándares de identificación causal mediante un énfasis en el diseño de investigación; sin embargo, sus propios límites se vuelven más visibles en entornos caracterizados por datos heterogéneos, sistemas complejos y ausencia de condiciones experimentales ideales.

Frente a estos desafíos, el artículo sostiene que el futuro de la econometría —y, por extensión, de la enseñanza de la economía— pasa por una integración más profunda entre enfoques basados en diseño, marcos causales estructurales (*a la* Judea Pearl) y herramientas de aprendizaje automático. La incorporación explícita del razonamiento causal, la formalización de supuestos y la capacidad de evaluar contrafactuales y extrapolaciones se vuelven competencias centrales en un mundo donde la predicción por sí sola resulta insuficiente frente a las necesidades de formulación de políticas, por ejemplo.

Desde esta perspectiva, la reforma de los métodos cuantitativos en la formación económica no debe entenderse únicamente como una actualización curricular, sino como una profunda reorientación conceptual: formar economistas capaces de interpretar, diseñar e intervenir en sistemas sociales complejos, utilizando la IA y el Big Data no como sustitutos del análisis causal, sino como instrumentos al servicio de una inferencia más rigurosa y socialmente relevante.

Declaración de contribución de autoría CRediT

Wilson A. Pérez-Oviedo: Conceptualización, Redacción – borrador original, Redacción (revisión y edición)

Agradecimientos

Los autores desean agradecer a FLACSO Ecuador por su apoyo en la realización de este artículo.

Declaración de conflictos de interés

Los autores declaran no tener ningún conflicto de intereses.

Referencias

1. Acemoglu, D. (2024). The simple macroeconomics of AI. *Economic Policy*, 40(121), 13–58. <https://doi.org/10.1093/epolic/eiae042>
2. Angrist, J. D., & Pischke, J.-S. (2010). The credibility revolution in empirical economics: How better research design is taking the con out of econometrics. *Journal of Economic Perspectives*, 24(2), 3–30. <https://doi.org/10.1257/jep.24.2.3>

3. Athey, S., & Imbens, G. W. (2017). The state of applied econometrics: Causality and policy evaluation. *Journal of Economic Perspectives*, 31(2), 3–32. <https://doi.org/10.1257/jep.31.2.3>
4. Athey, S., & Imbens, G. W. (2019). Machine learning methods that economists should know about. *Annual Review of Economics*, 11, 685–725. <https://doi.org/10.1146/annurev-economics-080217-053433>
5. Baqaee, D. R., & Farhi, E. (2019a). *Macroeconomics with heterogeneous agents and input–output networks* (NBER Working Paper No. 24684). National Bureau of Economic Research.
6. Baqaee, D. R., & Farhi, E. (2019b). The macroeconomic impact of microeconomic shocks: Beyond Hulten’s theorem. *Econometrica*, 87(4), 1155–1203. <https://doi.org/10.3982/ECTA15202>
7. Baqaee, D. R., & Farhi, E. (2019c). *The microeconomic foundations of aggregate production functions* (NBER Working Paper No. 25293). National Bureau of Economic Research.
8. Baqaee, D. R., & Farhi, E. (2021). *Networks, barriers, and trade* (NBER Working Paper No. 26108). National Bureau of Economic Research.
9. Bareinboim, E., & Pearl, J. (2016). Causal inference and the data-fusion problem. *Proceedings of the National Academy of Sciences*, 113(27), 7345–7352. <https://doi.org/10.1073/pnas.1510507113>
10. Bartoš, F., Maier, M., Wagenmakers, E.-J., Nippold, F., Doucouliagos, H., Ioannidis, J. P. A., Otte, W. M., Sladekova, M., Deressa, T. K., Bruns, S. B., Fanelli, D., & Stanley, T. (2024). Footprint of publication selection bias on meta-analyses in medicine, environmental sciences, psychology, and economics. *Research Synthesis Methods*. <https://doi.org/10.1002/jrsm.1703>
11. Bongers, S., Peters, J., Mooij, J. M., & Schölkopf, B. (2021). Foundations of structural causal models with cycles and latent variables. *Annals of Statistics*, 49(5), 2885–2915. <https://doi.org/10.1214/20-AOS2017>
12. Briggs, J., & Kodnani, D. (2023). Generative AI could raise global GDP by 7%. *Goldman Sachs Research*. <https://www.goldmansachs.com/intelligence/pages/generative-ai-could-raise-global-gdp-by-7-percent.html>

13. Chernozhukov, V., Chetverikov, D., Demirer, M., Duflo, E., Hansen, C., Newey, W., & Robins, J. (2018). Double/debiased machine learning for treatment and structural parameters. *The Econometrics Journal*, 21(1), C1–C68. <https://doi.org/10.1111/ectj.12097>
14. Christensen, G., & Miguel, E. (2018). Transparency, reproducibility, and the credibility of economics research. *Journal of Economic Literature*, 56(3), 920–980. <https://doi.org/10.1257/jel.20171350>
15. CNN Business. (2025, 29 mayo). *Why this leading AI CEO is warning the tech could cause mass unemployment*. <https://edition.cnn.com/2025/05/29/tech/ai-anthropic-ceo-dario-amodei-unemployment>
16. Couldry, N., & Mejías, U. A. (2019). *The costs of connection: How data colonizes human life and appropriates it for capitalism*. Stanford University Press.
17. Cukier, K., & Mayer-Schönberger, V. (2017). *Big data: A revolution that will transform how we live, work, and think* (Rev. ed.). Eamon Dolan/Mariner Books.
18. Deaton, A. (2010). Instruments, randomization, and learning about development. *Journal of Economic Literature*, 48(2), 424–455. <https://doi.org/10.1257/jel.48.2.424>
19. Deaton, A., & Cartwright, N. (2018). Understanding and misunderstanding randomized controlled trials. *Social Science & Medicine*, 210, 2–21. <https://doi.org/10.1016/j.socscimed.2017.12.005>
20. Federal Reserve Bank of New York. (2025). *The labor market for recent college graduates*. <https://www.newyorkfed.org/research/college-labor-market>
21. Federal Reserve Bank of St. Louis. (2025, August 25). *Recent college graduates bear brunt of labor market shifts*. <https://www.stlouisfed.org/on-the-economy/2025/aug/recent-college-grads-bear-brunt-labor-market-shifts>
22. Ferraro, P. J., & Shukla, P. (2020). Is a replicability crisis on the horizon for environmental and resource economics? *Review of Environmental Economics and Policy*, 14(2), 339–351. <https://doi.org/10.1093/reep/reaa011>
23. Forré, P., Mooij, J. M., & Peters, J. (2022). Causal discovery in the

- presence of cycles. *Journal of Machine Learning Research*, 23(146), 1–53. <https://www.jmlr.org/papers/v23/21-031.html>
24. Fortune. (2025, 22 julio). “*Intelligence too cheap to meter*” is AI’s next frontier, Sam Altman says. <https://fortune.com/2025/07/23/sam-altman-artificial-intelligence-too-cheap-jobs/>
 25. Gechert, S., Rannenberg, A., & Watzka, S. (2023). Fiscal multipliers, macroeconomic conditions, and publication bias. *Journal of Economic Surveys*, 37(2), 307–339. <https://doi.org/10.1111/joes.12503>
 26. Ghassami, A., Zhang, K., & Schölkopf, B. (2024). *Causal discovery in feedback systems: Methods and challenges*. *Foundations and Trends® in Machine Learning*.
 27. Heckman, J. J., & Urzúa, S. (2010). Comparing IV with structural models: What simple IV can and cannot identify. *Journal of Econometrics*, 156(1), 27–37. <https://doi.org/10.1016/j.jeconom.2009.09.007>
 28. Hosseini, S. M., & Lichtinger, G. (2025). *Generative AI as seniority-biased technological change: Evidence from U.S. résumé and job posting data* (SSRN Working Paper No. 5425555). <https://doi.org/10.2139/ssrn.5425555>
 29. Hünermund, P., & Bareinboim, E. (2023). Causal inference and data fusion in econometrics. *The Econometrics Journal*, 26(1), S1–S28. <https://doi.org/10.1093/ectj/utac031>
 30. IDC. (2024). *The digitization of the world: From edge to core (Updated forecast 2023–2027)*. International Data Corporation.
 31. Ioannidis, J. P. A. (2005). Why most published research findings are false. *PLoS Medicine*, 2(8), e124. <https://doi.org/10.1371/journal.pmed.0020124>
 32. Ioannidis, J. P. A. (2020). The challenge of reforming nutritional epidemiologic research. *JAMA*, 324(6), 559–560. <https://doi.org/10.1001/jama.2020.11028>
 33. Kamaruddin, R., Ahmad, A., & Mansor, N. (2023). Employability skills and labor market mismatch: Evidence from higher education graduates. *Journal of Education and Work*, 36(3), 247–263. <https://doi.org/10.1080/13639080.2023.2183579>
 34. Kitchin, R. (2021). *Data lives: How data are made and shape our world*. Bristol University Press.

35. Mooij, J. M., Janzing, D., & Schölkopf, B. (2024). Causal discovery for time series and dynamical systems. *Annual Review of Statistics and Its Application*, 11, 1–26. <https://doi.org/10.1146/annurev-statistics-033121-113547>
36. OECD. (2023). *OECD skills outlook 2023: Skills for a resilient green and digital transition*. OECD Publishing. <https://doi.org/10.1787/27452b1c-en>
37. Page, L., Noussair, C. N., & Slonim, R. (2021). The credibility crisis in economics: The importance of replication and pre-analysis plans. *Oxford Economic Papers*, 73(4), 1057–1084. <https://doi.org/10.1093/oep/gpab022>
38. Pearl, J. (2019). The seven tools of causal inference, with reflections on machine learning. *Communications of the ACM*, 62(3), 54–60. <https://doi.org/10.1145/3241036>
39. Pearl, J., & Mackenzie, D. (2018). *The book of why: The new science of cause and effect*. Basic Books.
40. Peters, J., Janzing, D., & Schölkopf, B. (2017). *Elements of causal inference: Foundations and learning algorithms*. MIT Press.
41. Portocarrero Ramos, J., Ruiz, K., & Alvarado, C. (2025). University training, skills mismatch, and youth employability in Latin America. *Journal of Labor and Society*. Advance online publication.
42. Ribeiro, M. T., Saxena, N. A., & Lundberg, S. (2022). Challenges of explaining AI models when critical contextual variables are missing. *Nature Machine Intelligence*, 4, 673–685. <https://doi.org/10.1038/s42256-022-00532-0>
43. Schölkopf, B., Locatello, F., Bauer, S., Ke, N. R., Kalchbrenner, N., Goyal, A., & Bengio, Y. (2021). Toward causal representation learning. *Proceedings of the IEEE*, 109(5), 612–634. <https://doi.org/10.1109/JPROC.2021.3058954>
44. Segbenya, M., Kpessa-Whyte, M., & Teye, J. K. (2023). Graduate employability and skills mismatch in emerging economies. *International Journal of Educational Development*, 96, 102702. <https://doi.org/10.1016/j.ijedudev.2022.102702>
45. Strömberg, E. A. (2019). Publication bias and selective reporting in experimental economics. *Journal of Economic Behavior & Organization*, 163, 252–268. <https://doi.org/10.1016/j.jebo.2019.05.006>