

Ciencia de datos, *big data* e inteligencia artificial en las ciencias sociales computacionales. Crítica epistemológica y aportes para la investigación *

Data Science, Big Data, and Artificial Intelligence in Computational Social Sciences: Epistemological Critique and Contributions for Research

Ciência de dados, *big data* e inteligência artificial nas ciências sociais computacionais: crítica epistemológica e contribuições para a pesquisa

Aminahuel Aimé^a

DOI: <https://doi.org/10.11144/Javeriana.syp44.cdbd>

Universidad Nacional de Río Cuarto, Argentina

aime.aminahuel@hum.unrc.edu.ar

ORCID: <https://orcid.org/0000-0001-5504-5175>

Recibido: 06 octubre 2025

Aceptado: 20 noviembre 2025

Publicado: 30 diciembre 2025

Resumen:

El artículo examina el impacto del *big data* y la inteligencia artificial en el campo de las ciencias sociales computacionales. Se destaca la incidencia de este campo interdisciplinar en los estudios sociales en el marco de tres problemas específicos: a) el reduccionismo ontológico, b) el mito de la neutralidad y c) las complejidades en la distinción entre causalidad y correlación de variables en los fenómenos sociales. Desde una perspectiva crítica, se problematizan los supuestos epistemológicos que sostienen el uso acrítico de algoritmos y modelos de aprendizaje automático en la investigación social aislados del contexto de producción teórica. A partir de una revisión del estado del arte y de la presentación del análisis de un *dataset* mediante la aplicación de métodos de agrupamiento no supervisado, específicamente el algoritmo de K-means y el clúster jerárquico, se proponen estrategias que integran enfoques computacionales con tradiciones cualitativas, con el fin de potenciar investigaciones aplicadas desde la triangulación metodológica y atentas a los contextos sociopolíticos donde los datos son producidos y utilizados.

Palabras clave: ciencias sociales computacionales, *big data*, inteligencia artificial (IA).

Abstract:

The article examines the impact of big data and artificial intelligence in the field of computational social sciences. It highlights the influence of this interdisciplinary field on social studies within the framework of three specific identified problems: a) ontological reductionism, b) the myth of neutrality, and c) complexities in the distinction between causality and correlation of variables in social phenomena. From a critical perspective, the paper problematizes the epistemological assumptions that sustain the uncritical use of algorithms and machine learning models in social research when isolated from the context of theoretical production. Based on a state-of-the-art review and the presentation of a case study involving dataset analysis through the application of unsupervised clustering methods—specifically the K-means algorithm and hierarchical clustering—strategies are proposed to integrate computational approaches with qualitative traditions. The goal is to enhance applied research through methodological triangulation, remaining attentive to the sociopolitical contexts where data are produced and utilized.

Keywords: Computational Social Sciences, Big Data, Artificial Intelligence (AI).

Resumo:

O artigo examina o impacto da *big data* e da inteligência artificial no campo das ciências sociais computacionais. Destaca-se a incidência deste campo interdisciplinar nos estudos sociais no âmbito de três problemas específicos identificados: a) reducionismo ontológico, b) mito da neutralidade e c) complexidades na distinção entre causalidade e correlação de variáveis nos fenômenos sociais. A partir de uma perspectiva crítica, problematizam-se os pressupostos epistemológicos que sustentam o uso acrítico de algoritmos e modelos de aprendizado de máquina na pesquisa social isolados do contexto de produção teórica. Com base em uma revisão do estado da arte e na apresentação de um caso de análise de um conjunto de dados (*dataset*) mediante a aplicação de métodos de agrupamento não supervisionado —especificamente o algoritmo K-means e o cluster hierárquico— propõem-se estratégias que integram abordagens computacionais com tradições qualitativas. O objetivo é potencializar pesquisas aplicadas a partir da triangulação metodológica e atentas aos contextos sociopolíticos onde os dados são produzidos e utilizados.

Palavras-chave: ciências sociais computacionais, *big data*, inteligência artificial (IA).

Notas de autor

^a Autora de correspondencia. Correo electrónico: aime.aminahuel@hum.unrc.edu.ar

Introducción

Las ciencias sociales computacionales (CSC) representan un campo interdisciplinario en el que convergen tradiciones de las ciencias sociales con desarrollos computacionales contemporáneos. Esta convergencia contempla tanto la incorporación instrumental de técnicas provenientes de la informática, la ciencia de datos y la inteligencia artificial como también una reconfiguración de los modos de producción de conocimiento social hacia enfoques orientados a la modelización, la simulación y el análisis algorítmico que dialogan con problemas clásicos de explicación sociológica. Si bien la aplicación de métodos cuantitativos no es un fenómeno reciente en las ciencias sociales, pues se podrían mencionar vastos desarrollos disciplinares en el campo de la econometría o la sociología cuantitativa, el surgimiento de las CSC a finales del siglo XX y principios del XXI supuso un quiebre paradigmático que tensionó los debates epistémicos entre el comprensivismo —predominante en los estudios sociales latinoamericanos tal y como sostiene Dieter Nohlen (2007)¹— y el conductismo digital (Rouvroy y Berns, 2016), como un reemergente positivista contemporáneo, cuya fuente empírica es el comportamiento de los internautas en redes sociales digitales.

El campo emergente de las CSC se desarrolla en diálogo —y en ocasiones en tensión— con corrientes sociológicas que problematizan el alcance explicativo de los métodos computacionales. Desde la sociología analítica, por ejemplo, ya se advierte que la identificación algorítmica de patrones no equivale a explicación social. En esa línea, autores como Keuschnigg *et al.* (2018) enfatizan la necesidad de reconstruir mecanismos causales que articulen agencia, estructura y procesos micro-macro para la integración de los análisis computacionales a teorías generales en ciencias sociales. En paralelo, otros enfoques, como la sociología digital (Marres, 2017), proponen comprender las tecnologías de datos como ensamblajes sociotécnicos que configuran activamente las formas de observación de lo social, lo que implica un cuestionamiento a la idea de neutralidad metodológica en la recolección de *big data* o en su análisis computacional con técnicas informáticas. En suma, ambas perspectivas no rechazan las ciencias sociales computacionales, pero sí advierten la ingenuidad de las perspectivas clásicas, a la vez que enfatizan que su potencial depende de una integración crítica con marcos teóricos que eviten reduccionismos correlacionales o tecnodeterministas.

En concordancia, desde el posicionamiento de este escrito, la irrupción epistémica de las CSC es consecuencia de factores técnicos, políticos y globales que se relacionan con tres fenómenos clave cuya problematización comprende la teoría general del abordaje propuesto. Por un lado, la emergencia de la sociedad red (Castells, 1996; Van Dijck, 2016), que promovió la predominancia de las interconexiones sociales; por otro lado, el desarrollo del capitalismo financiero con grandes volúmenes de información como base y su correlato simultáneo, el capitalismo de plataformas, cognitivo o informacional (Fumagalli, 2010; Srnicek, 2019), en el que la mercancía central es la información y en torno a ella se configuran nuevas infraestructuras de datos que posibilitaron la explosión del *big data*. Y, finalmente, la consolidación de la ciencia de datos —junto con los algoritmos de aprendizaje automático e inteligencia artificial— como parte de la infraestructura técnica para el procesamiento de grandes volúmenes de información. Este desarrollo configuró un ecosistema técnico que posteriormente sería apropiado y resignificado por las ciencias sociales computacionales (Crawford, 2022; Sosa Escudero, 2019; Sadin, 2020).

De acuerdo con lo expuesto, el presente artículo problematiza la utilización del *big data*, desde la ciencia de datos y la inteligencia artificial en las ciencias sociales computacionales, a partir de tres dimensiones epistemológicas fundamentales. Primero, se denuncia el reduccionismo ontológico inherente al *dataísmo* que simplifica la complejidad intrínseca de los fenómenos sociales (intención, significado, contexto), al reducirlos a meras variables cuantificables y relaciones estadísticas. Segundo, se cuestiona el denominado *mito de la neutralidad*, al sostener que los *datasets* y los algoritmos no son instrumentos objetivos, sino artefactos sociales que inevitablemente reflejan y perpetúan sesgos históricos, culturales y estructurales (como los de raza o género, así como los límites de ciertos algoritmos o estructura de datos en su representación de lo

social). Tercero, se examina la confusión entre causalidad y correlación, señalando la dificultad de establecer inferencias causales a partir de modelos predictivos de caja negra que operan sobre la correlación. Este punto subraya la necesidad de incorporar la teoría social como pilar fundamental para la interpretación y validación de los resultados. A partir de esta crítica, mediante la problematización de un caso de utilización de algoritmos para un análisis de datos sociales, se proponen estrategias que integran enfoques computacionales con tradiciones cualitativas de las ciencias sociales, con el fin de potenciar investigaciones aplicadas desde la triangulación metodológica y atentas a los contextos sociopolíticos donde los datos y las tecnologías son producidos y utilizados.

Desde el punto de vista metodológico, el presente trabajo se inscribe en una investigación de carácter exploratorio-reflexivo con sustento en un caso empírico de análisis de un *dataset* mediante la aplicación de métodos de agrupamiento no supervisado, específicamente el algoritmo de K-medias (K-means) y el clúster jerárquico. La estrategia combina una revisión crítica del estado del arte sobre las ciencias sociales computacionales con el análisis de un caso aplicado de modelización algorítmica, utilizado como dispositivo heurístico para problematizar los alcances y límites epistemológicos de la ciencia de datos en contextos sociales. Este diseño no persigue validaciones predictivas generalizables, sino la explicitación de decisiones analíticas, supuestos técnicos y sus implicancias interpretativas. En este sentido, la metodología se concibe como una triangulación entre análisis bibliográfico, experimentación computacional y reflexión teórica, con el objeto de reconocer los límites inherentes al carácter situado de los datos y a la mediación algorítmica en su procesamiento.

Genealogía de las ciencias sociales computacionales

El origen de las ciencias sociales computacionales (CSC) puede rastrearse en la incorporación temprana de herramientas computacionales —incluida la ciencia de datos— a problemas sociales, por ejemplo, a los análisis espaciales y el uso de algoritmos de clasificación. Por ello, la geografía europea —específicamente la geomática— fue pionera en estos abordajes de acuerdo con Day Montarcé (2021), puesto que, en su quehacer disciplinar, incluía acciones tales como el procesamiento, el análisis, la representación y la distribución de información geográficamente referenciada. Sin embargo, esta adopción técnica se acompañó de debates epistemológicos en la disciplina, respecto a cómo traducir fenómenos geográficos y socioambientales en modelos formales, y el debate académico desarrollado con mayor profundidad por la sociología analítica por autores como Keuschnigg *et al.* (2018). Con el tiempo, la incorporación progresiva de herramientas informáticas para el análisis de grandes volúmenes de datos y el desarrollo de técnicas de aprendizaje automático e inteligencia artificial de las últimas décadas ampliaron el repertorio metodológico de las ciencias sociales, lo que implicó un proceso de apertura a nuevas formas de modelización de lo social y exploración empírica que fueron adoptadas —con ritmos y alcances diversos— por múltiples disciplinas. Siguiendo los aportes de Gualda Caballero *et al.* (2023), en la actualidad, los nuevos desarrollos disciplinares proponen investigaciones aplicadas al campo de lo social, que van desde el análisis de dinámicas electorales hasta la identificación de patrones de consumo en contextos organizacionales, estudios de clústeres en temáticas de salud, pasando por abordajes de la opinión pública en redes sociales —entre los estudios más referenciados—, entre otros. Esto evidencia el modo de integración de los métodos computacionales a problemas y proyectos de investigación diversos en el campo de las ciencias sociales.

En este marco, resulta necesario precisar qué se entiende por ciencia de datos, dado que este campo constituye una de las infraestructuras técnicas que nutren —pero no definen por completo— a las ciencias sociales computacionales. Hastie *et al.* (2015) sostiene que se pueden delimitar antecedentes disciplinares y desarrollos tecnológicos:

- El aprendizaje automático (AA: *machine learning*) que construye algoritmos que pueden aprender de los datos.
- El aprendizaje profundo (*deep learning*) como una subdisciplina del *machine learning* que se caracteriza por el uso de redes neuronales artificiales con múltiples capas ocultas (*deep neural networks*) para extraer automáticamente representaciones, características (*features*) y abstracciones de alto nivel a partir de datos brutos (Goodfellow *et al.*, 2016).
- El aprendizaje estadístico (*statistical learning*) como una rama de la estadística aplicada que emergió en respuesta al AA, enfatizando en los modelos estadísticos y en la evaluación de la incertidumbre.
- La ciencia de datos (CD) es la extracción de conocimiento de los datos, utilizando ideas de matemática, estadística, AA e ingeniería.
- La inteligencia artificial comprendida como conjunto de técnicas, modelos y algoritmos, particularmente el *machine learning* y el *deep learning*, que permiten procesar, clasificar y, sobre todo, predecir patrones.

La emergencia de la ciencia de datos se relaciona a mutaciones históricas y confluencias disciplinares que sentaron sus bases, muchas de ellas gestadas antes de la irrupción del *big data*. La ciencia de datos es el resultado de la convergencia de líneas de investigación clásicas y emergentes, influenciadas por demandas tanto académicas como industriales y de defensa, entre las que se destacan las siguientes:

1) Mutaciones en la estadística y en la computación, así como la intersección de estas disciplinas. Esto incluye la evolución de la estadística inferencial hacia la estadística computacional y el aprendizaje automático (*machine learning*), en el cual se prioriza la optimización algorítmica y el manejo de grandes volúmenes de datos sobre las pruebas paramétricas tradicionales.

2) Desarrollos en lingüística y en procesamiento del lenguaje natural (PLN). La línea de la lingüística formal y de la semiótica ha sido crucial, con trabajos pioneros como los de Charles Peirce sobre la representación del significado (Everaert-Desmedt, 2004). Esta línea se sofisticó tecnológicamente con las investigaciones de instituciones como IBM y su histórico interés en la traducción automática y, posteriormente, en el procesamiento del lenguaje natural (PLN), fundamental para analizar textos y discursos sociales a escala masiva.

3) Ingeniería de reconocimiento de patrones (*pattern recognition*). Esta disciplina se orientaba inicialmente a problemas de visión artificial y de procesamiento de señales, pero sus algoritmos (como las redes neuronales iniciales) son precursores directos de las técnicas de IA utilizadas en la actualidad para detectar regularidades en datos sociales complejos.

4) Los problemas de la industria, la inteligencia y la defensa. Los desarrollos en ciencia de datos obedecieron, en sus orígenes, a necesidades pragmáticas y financieras. La industria (posteriormente las grandes corporaciones de *software*) demandó soluciones para la optimización logística, el *marketing* predictivo y la automatización. Paralelamente, las agencias de inteligencia y defensa, como la Agencia de Proyectos de Investigación Avanzada de Defensa de EE. UU. (Darpa), financiaron durante décadas la investigación en PLN, minería de datos y sistemas expertos, con el objetivo de mejorar la toma de decisiones estratégicas, la vigilancia y la seguridad nacional (Schab, 2020).

Estas trayectorias, que operaron de forma paralela, convergieron en el siglo XXI con la digitalización masiva de las conductas financieras y de consumo de los individuos, lo que consolidó los principales antecedentes del *big data*, además del ecosistema metodológico y tecnológico para su análisis. Luego, a partir de la irrupción de las redes sociodigitales (Aminahuel y Rodríguez, 2024), el abordaje de *big data* giró hacia un nuevo principio epistémico central: la presunción de que los comportamientos humanos, traducidos en conductas digitales, y las interacciones sociales —plataformizadas— generan hoy vastas trazas o huellas digitales (Morales, 2019). Estas incluyen los modos de interacción en redes sociales de los individuos sobre lo que

gusta o no gusta, perfiles sociodemográficos y económicos, intereses políticos, estados de ánimo, patrones de consumo, transacciones, estados financieros, geolocalización, etc., y son huellas porque pueden ser capturadas, almacenadas y analizadas a una escala y a una velocidad antes inimaginables. Sin embargo, esta disponibilidad masiva de datos presenta algunas dificultades para el investigador social. Dada la multicolinealidad de las bases de datos y la enorme cantidad de variables por observación, el análisis de *big data* en procesos sociales presupone una profundidad de conocimiento mucho mayor que otras modalidades de análisis cuantitativos, pero, a la vez, otorga un conocimiento profundo si se trabaja con diferentes técnicas de análisis.

Esto orientó a las ciencias sociales a repensar sus herramientas, migrando de los métodos de recolección tradicionales (encuestas por muestreo, observación participante, análisis de contenido cuantitativo, etc.) a aquellos que procesan grandes conjuntos de datos estructurados y no estructurados mediante algoritmos complejos y modelos de aprendizaje automático como Random Forest/árboles de decisión, K-Nearest Neighbors (K-NN), entre otros.

Autores como Gary King (2011) y David Lazer *et al.* (2009) definen las CSC como un espacio de investigación interdisciplinario que integra métodos computacionales con marcos conceptuales de las ciencias sociales, en el que la capacidad de análisis a gran escala debe articularse con teorías explicativas sobre los procesos sociales. La ventaja inicial de este tipo de análisis, para los autores clásicos, radicó en la capacidad de observar el comportamiento social casi en tiempo real y a nivel poblacional, lo que implicó una ganancia significativa en términos de escala así como de profundidad. Esta nueva aproximación, no obstante, trajo consigo una serie de desafíos teóricos que requieren un análisis crítico que observe tanto los sesgos —temática mayormente abordada en la literatura— como los alcances de los análisis, partiendo de la base de que el conocimiento de lo social se obtiene en torno a la sistematización cuantitativa de los comportamientos digitales, en palabras de la filósofa belga Antoinette Rouvroy (2013), el llamado *conductismo de datos*.

El desplazamiento metodológico desarrollado ha sido objeto de discusión también dentro de corrientes contemporáneas como la sociología analítica. Bajo este enfoque, se advierte que la creciente capacidad de modelización algorítmica no equivale automáticamente a una mejora en la explicación social; es decir, los patrones identificados mediante técnicas computacionales —agrupamientos, correlaciones masivas o predicciones— constituyen apenas un primer nivel descriptivo que requiere ser articulado con teorías de mecanismos sociales capaces de explicar cómo y por qué emergen determinadas regularidades. En este sentido, la sociología analítica combate la creciente sustitución de la comprensión causal por la predicción social —fenómeno recurrente en el uso acrítico de algoritmos en ciencias sociales—, al sostener que la potencia exploratoria de las ciencias sociales computacionales solo adquiere valor explicativo cuando se integra en marcos conceptuales que vinculan agencia individual, estructuras institucionales y dinámicas colectivas (Keuschnigg *et al.*, 2018).

En consecuencia, asumir las ciencias sociales computacionales exclusivamente como una extensión de la ciencia de datos aplicada a lo social resulta epistemológicamente restrictivo. El campo abarca no solo técnicas de clasificación, predicción, etc., sino también modelización de mecanismos sociales, simulaciones basadas en agentes, análisis de redes complejas y experimentación digital orientada a explorar procesos emergentes. Esta ampliación metodológica implica reconocer que los algoritmos son herramientas técnicas y, a la vez, dispositivos que median la construcción del objeto social investigado. Desde esta perspectiva, el desafío no radica en maximizar la eficiencia predictiva, sino en articular los métodos computacionales con tradiciones teóricas capaces de dotar de sentido sociológico a los patrones detectados (Keuschnigg *et al.*, 2018). En dicho marco, Keuschnigg *et al.* (2018) afirman que

las CSC se ubican en la intersección de las ciencias sociales y la computación, una intersección que incluye el análisis de datos observacionales a escala web, experimentos de estilo de laboratorio virtual y modelado computacional. Similar a AS [sociología analítica], las CSC perciben los fenómenos sociales como resultado de interacciones sociales entre agentes individuales que están integrados en sistemas sociales complejos y esto plantea desafíos explicativos considerables porque el

comportamiento de las entidades en una “escala” de la realidad no se puede rastrear fácilmente hasta las propiedades de las entidades en la escala inferior. (p. 4)²

En suma, aunque las ciencias sociales computacionales comparten herramientas con la ciencia de datos, su alcance metodológico excede ampliamente la optimización algorítmica o la modelización predictiva. Mientras la ciencia de datos privilegia la eficiencia técnica en la clasificación y predicción, las CSC incorporan prácticas como la simulación basada en agentes, el análisis de redes complejas, la experimentación social digital y la modelización de mecanismos colectivos orientados a la explicación de procesos sociales. En este sentido, el uso de algoritmos no constituye un fin en sí mismo, sino un medio para explorar dinámicas emergentes, relaciones causales y configuraciones sociotécnicas. Reconocer esta distinción resulta clave para evitar reducir las CSC a una aplicación instrumental de técnicas estadísticas, a partir de la comprensión de su carácter interdisciplinar en tanto espacio de conocimiento donde la teoría social, la modelización computacional y la reflexión epistemológica se articulan en la producción de nuevos saberes (Cioffi-Revilla, 2010; Marres, 2017; Keuschnigg *et al.*, 2018).

La irrupción del big data y el contexto capitalista

El desarrollo y la consolidación de las ciencias sociales computacionales (CSC) se consideran inherentes a la transformación epocal que supuso el surgimiento de la sociedad de la información (Castells, 1996) y, concretamente, la consolidación del *big data* como un insumo analítico esencial. Al conceptualizar una organización social cimentada en la infraestructura tecnológica y la digitalización de las interacciones, la sociedad de la información facilitó las condiciones para que los comportamientos y las dinámicas sociales empezaran a dibujar huellas digitales (Morales, 2019) que caracterizan el comportamiento de los individuos en las plataformas. En este contexto, el *big data* se erige como un nuevo paradigma de la observación social que contempla entre sus elementos esenciales la disponibilidad masiva de datos y de herramientas técnicas para su análisis. De acuerdo con Lazer *et al.* (2009), su caracterización principal —articulada en torno al volumen, la velocidad de producción y la variedad de sus formatos (las 3V)— implica un salto metodológico que permite a las ciencias sociales superar las restricciones inherentes al muestreo y analizar poblaciones íntegras con una granularidad y una inmediatez sin precedentes.

Este acceso a conjuntos de datos sociales, que se generan de forma automática mediante las interacciones en entornos digitales por parte de internautas (redes sociales, transacciones bancarias electrónicas, sensores, movimientos geolocalizados, intercambios de mensajes, etc.), introduce una dimensión de escalabilidad y profundidad que históricamente se había constituido como un desafío para las disciplinas sociales en los estudios cuantitativos. La capacidad de modelizar fenómenos sociopolíticos complejos como la polarización o los flujos informacionales, a una escala global y con procesamiento en tiempo real, constituye en la actualidad una ventaja clara frente a los métodos de recolección tradicionales. Sin embargo, es menester preguntarnos específicamente sobre el concepto de *big data* por fuera de las definiciones tradicionales. En dicho marco, para la autora Abarzúa Cutroni (2022) el *big data* es algo más que disponibilidad masiva de datos, pues se trata de un fenómeno social, político y cultural:

En principio, el *big data* puede ser pensado como un dispositivo técnico relativamente novedoso, asociado al almacenamiento y procesamiento algorítmico de una cantidad de datos ingente y diversa. Sin embargo, el *big data*, como parte de una serie de innovaciones tecnológicas, trasciende lo meramente técnico para transformarse en un fenómeno social, político y cultural que tiene como denominador común la datificación de las sociedades contemporáneas. En la actualidad, las relaciones sociales —al menos una parte de estas— no solo se han virtualizado o digitalizado, sino que están mediadas por algoritmos, en los que hemos delegado parte del trabajo de la sociedad y la cultura. Rodríguez (2018) plantea que la algoritmización de la sociedad es posible hoy “no es solo porque los usamos para cualquier cosa, sino también y sobre todo porque todos ellos se encuentran conectados a través de sistemas que son incesantemente alimentados por nuestros usos, de manera tal de poder procesar los registros de diferentes actividades (una solicitud de amistad, la visión de una serie televisiva, la frecuencia cardíaca de un

running en el parque, la búsqueda de un dato cualquiera en internet) en un suelo común que permita luego la ‘personalización’, la asignación de esa masa de datos a un individuo, la definición de un perfil”. (p. 158)

Para la autora, las 3V que caracterizaron el *mainstream* de la literatura mercadotécnica sobre el *big data* no son suficientes para la reflexión en las ciencias sociales “ya que para este ámbito se requiere también revisar la calidad de los datos, su representatividad y los dilemas éticos asociados” (Abarzúa Cutroni, 2022, p. 158). En el mismo sentido, Rouvroy (2013) advierte que la predominancia del *big data* produjo el advenimiento de un conductismo de datos en el que la difuminación de los límites entre datos, información y conocimiento social es indistinguible:

Llamaré “conductismo de datos” a esta nueva forma de producir conocimiento sobre actitudes, comportamientos o eventos de preferencias futuras sin tener en cuenta las motivaciones psicológicas, los discursos o las narrativas del sujeto, sino más bien basándome en los datos [...]. Son estas enormes cantidades de datos en bruto, estos enormes datos estadísticos impersonales, en constante evolución, lo que hoy constituye “el mundo” en el que los algoritmos “revelan” lo que la gubernamentalidad algorítmica toma para “la realidad”. “Realidad”: ese conocimiento que parece retener, no parece producido, pero siempre está ahí, inmanente a las bases de datos, esperando ser descubierto por procesos algorítmicos estadísticos. El conocimiento ya no se produce sobre el mundo, sino desde el mundo digital [las proyecciones que los algoritmos arrojan y que son base para la construcción de la realidad y que siempre se basan en la lógica de la ganancia económica, en el mejor de los casos]. (Rouvroy, 2013, pp. 2-4; traducción propia)

Esta marcada inclinación del concepto orientado hacia el volumen de información ha soslayado la discusión sobre la tensión epistemológica que emerge del debate académico en relación a los modos de extracción y, por ende, de estructuración que presentan estos datos, ya que su existencia se basa en intereses mercantiles privados de gigantes tecnológicos en el marco de lo que algunos autores denominan el *capitalismo cognitivo* (Fumagalli, 2010) o de *plataformas* (Srnicsek, 2019). En dicho marco, el italiano Andrea Fumagalli (2010) rastrea los orígenes del capitalismo cognitivo en relación con el creciente peso de las tecnologías lingüísticas y comunicativas. Mientras en el capitalismo fordista la función de la comunicación, principalmente publicitaria, estaba netamente separada de la fase de producción, en el capitalismo cognitivo la comunicación forma una unidad con la producción en el sentido que la determina y la dirige. Para Fumagalli (2010, p. 163), el nexo entre comunicación y producción está mediado por la actividad de consumo, una actividad que, con el devenir del capitalismo neoliberal, se ha vuelto cada vez más globalizada y central para el modelo especulativo.

Desde la perspectiva del autor, el acto de consumo en el capitalismo cognitivo no es reducible solamente a la adquisición final de un bien o servicio, ya que configura el proceso mercantil, en tanto el consumo implica una participación en la opinión pública, el *marketing* y el acto comunicativo, simultáneamente. El accionar de los internautas en las redes sociogiditales ejemplifica con claridad este fenómeno, asociado al *microtargeting* (Barbu, 2014): toda vez que los usuarios de redes sociales interactúan entre sí mediante las aplicaciones de celular —consumen contenido audiovisual, reaccionan a videos cortos o a publicaciones de distinta índole, buscan negocios mientras escuchan un tipo de música y conversan sobre la necesidad de algún consumo, etc. — generan información en escala. A partir de la correlación que establece el algoritmo entre los actos que desarrolla el usuario en el entorno virtual, se constituyen perfiles estandarizados de tipos de consumidores que configuran, luego, la estrategia comercial y productiva, así como el modelo de negocios.

A su vez, la aparición del *big data* está asociada a la aparición de herramientas para su análisis y de infraestructuras privadas de captación y almacenamiento. Abarzúa Cutroni (2022) argumenta que esta nueva forma de conocimiento de lo social representa un voluminoso cuerpo de información que difiere de los datos tradicionales de las ciencias sociales por su origen y contexto, ya que ofrece una capacidad fundamental de incrementar las escalas de observación. Este cambio ha impulsado la renovación metodológica, al exigir a las disciplinas repensar la articulación micro-macro y fomentar los métodos mixtos (Gualda, 2022; Parra Saiani y Piovani, 2021, citado en Abarzúa Cutroni, 2022, pp. 158-159). La autora enfatiza que la mayor parte de estos datos son extraídos y almacenados por corporaciones privadas como Amazon, Google o Meta, que

desarrollaron las plataformas sin priorizar la privacidad de los usuarios. Esta concentración del valor comercial y de la capacidad de procesamiento genera una barrera de acceso casi infranqueable para la investigación académica, especialmente en el sur global, dado que los conjuntos de datos no son, en la práctica, públicos ni de libre disposición, así como tampoco resultan transparentes los algoritmos utilizados para su recolección. A partir de esta crítica estructural, a continuación se presentan los tres dilemas epistemológicos mencionados anteriormente.

Debates estructurales en las CSC: los tres dilemas epistemológicos

La aplicación de las herramientas de la IA y la ciencia de datos para el abordaje del *big data* y su relación con los problemas sociales ha generado un intenso debate crítico que se articula en torno a tres problemas estructurales previamente identificados. Esta crítica busca evitar la subordinación de la teoría social a la lógica técnica del algoritmo (Mejías y Couldry, 2019). En este apartado, se presentan los tres dilemas epistemológicos que emergen con mayor presencia en los abordajes críticos sobre el uso de algoritmos para el análisis de procesos sociales, a saber: el dataísmo y el reduccionismo ontológico; el mito de la neutralidad y la replicación de sesgos, y la discusión sobre causalidad, correlación y el problema de la inferencia.

Dataísmo y reduccionismo ontológico

El primer dilema se refiere al reduccionismo ontológico conocido como *dataísmo* (Mejías y Couldry, 2019). Esta ideología, promovida por la lógica del *big data*, tiende a simplificar la complejidad de los fenómenos sociales al enfocarse únicamente en lo cuantificable. La acción social, tal como la concibieron pensadores clásicos como Weber (2014), está cargada de intención, significado, contexto y agencia. Estos elementos son difíciles de capturar completamente mediante variables estadísticas o modelos predictivos basados en la correlación. La primacía de la cantidad sobre la cualidad amenaza con una pérdida de densidad interpretativa, en la cual la capacidad del modelo para predecir se prioriza sobre su capacidad para explicar y comprender el porqué de un fenómeno. En este sentido, el determinismo tecnológico se cierne como un riesgo, en el que la explicación se subordina al poder algorítmico y se descuida la complejidad de los sistemas sociales (Turkle, 2011).

El foco central de la controversia reside en la tesis promovida hace tiempo por ciertos teóricos de la computación, como Chris Anderson (2008), sobre el presunto fin de la teoría. Esta postura argumenta que el *big data*, al posibilitar la identificación de patrones mediante una correlación exhaustiva en grandes volúmenes de datos, minimiza la necesidad de la construcción de hipótesis *a priori* y de modelos explicativos causales originados en la teoría social. De esta manera, la correlación se consideraría suficiente para orientar la acción y la capacidad predictiva del algoritmo eclipsaría la exigencia de comprender el porqué subyacente a los fenómenos observados. Esta tesis se centra en cómo la analítica masiva no busca comprender la complejidad del comportamiento individual o social (la intencionalidad), sino únicamente anticiparlo para ejercer un control normativo o comercial, lo cual se considera profundamente reductivo en su alcance político y epistemológico.

Frente a esta perspectiva, la crítica académica ha demarcado las fronteras metodológicas. Académicas/os como Antoinette Rouvroy y Thomas Berns (2013) problematizaron el concepto de *gubernamentalidad algorítmica* para evidenciar el riesgo de una dictadura de los datos que sustituye la deliberación y el conocimiento contextual por una gobernanza exclusivamente predictiva.

El economista argentino Walter Sosa Escudero (2019) planteó una tesis similar que sostiene que la fortaleza del *big data* radica en su capacidad para identificar patrones y correlaciones, lo que permite anticipar qué sucederá con alta precisión. Sin embargo, esta superioridad predictiva no implica automáticamente una

superioridad explicativa o causal. Desde su formación en econometría y estadística, Sosa Escudero (2019) destaca que, si bien el *big data* expande las oportunidades para la predicción, no suprime los problemas esenciales “clásicos” de investigación asociados a la inferencia causal y a la validez metodológica. El mero incremento en la cantidad de datos, aunque sea masivo, no asegura automáticamente la solidez de las inferencias si el diseño de la investigación presenta debilidades estructurales o si la fuente de datos introduce sesgos de selección sistemáticos. La determinación de la causalidad, crucial para las ciencias sociales, exige una estructura de supuestos y una sólida base teórica, *llámese coherencia epistemológica interna*, que los modelos basados exclusivamente en la correlación no logran aportar. En el mismo sentido, Cédric Durand (2022) advierte que las características intrínsecas al *big data* alimentan una epistemología empirista ingenua, tras la idea que este nuevo régimen de conocimiento procedería por pura inducción automática, ya que los datos entregarían la verdad sin pasar por el desvío de la teoría.

En consecuencia, la irrupción del *big data* obliga a una reorganización en el quehacer del investigador social, quien debe evitar la adhesión al dataísmo y la creencia ideológica de que los datos son inherentemente autoexplicativos. El desafío fundamental de las CSC es utilizar las herramientas algorítmicas para la identificación eficiente de patrones, pero siempre subordinando esos hallazgos a la teoría social con el propósito de otorgarles significado, contexto y poder explicativo. Esta crítica converge con los planteamientos de la sociología digital, que entiende los métodos computacionales como prácticas sociotécnicas que registran la realidad social a la vez que participan en su configuración. Desde este enfoque, las ciencias sociales computacionales no pueden reducirse a la aplicación instrumental de técnicas de ciencia de datos, ya que cada procedimiento de clasificación, modelización o visualización incorpora decisiones normativas y epistemológicas que delimitan qué dimensiones de lo social se vuelven visibles (Marres, 2017). La reflexividad epistemológica, por lo tanto, se convierte una condición necesaria para evitar que la infraestructura técnica naturalice interpretaciones parciales de los fenómenos sociales.³

El mito de la neutralidad y la replicación de sesgos

Una de las críticas más férreas a la IA y el *big data* en la investigación social se dirige al mito de la neutralidad algorítmica. Este mito asume que los datos son un reflejo puro y objetivo de la realidad; sin embargo, no resulta una novedad explicitar que todo dato, en tanto artefacto social, es el resultado de un conjunto de decisiones humanas de recolección, selección y clasificación (Lisa Gitelman, 2013), aun en tiempos automáticos, y, por lo tanto, son pasibles de contaminarse por sesgos en sus distintas fases.

Claudio Celis Bueno (2021) sostiene que los llamados *sesgos maquínicos* son la transmisión de prejuicios sociales a una máquina. Se podría trasladar esta concepción a toda una interfaz, como son las plataformas de redes sociodigitales, puesto que el autor afirma que los sesgos se presentan, en el caso de los algoritmos de aprendizaje automático, de manera estructural, ya que “su propio proceso de entrenamiento depende de las bases de datos que contienen en sí dichos sesgos” (p. 1). Citando a Joler y Pasquinelli (2021), el autor describe tres tipos de sesgos: los que dependen de la morfología de las bases de datos, los que dependen del programador y aquellos que dependen de factores históricos y sociales. Asimismo, en el artículo de Celis Bueno se presentan diferentes ejemplos empíricos de casos de sesgos algorítmicos para demostrar que el aprendizaje automático naturaliza un determinado orden del mundo que se manifiesta en los *datasets* y luego reproducen ese orden.

En este sentido, podría cuestionarse la arquitectura de las redes sociales que, sobre reglas de juego preestablecidas —con carácter mercantil— para la interacción de los internautas, limitan y disciplinan las formas de comportamientos sociales digitales que, luego, son sistematizados en escala y que resultan en datos masivos de conductas digitales (conductismo de datos), a la vez que son el mismo insumo sobre el que se construyen las tipologías de perfiles para realizar recomendaciones automáticas. Por lo tanto, en palabras del mismo autor, desde una perspectiva crítica, la pretensión de objetividad total de los algoritmos es imposible,

“ya que el sistema mismo requiere de los prejuicios heredados para su propio entrenamiento” (p. 2). En estudios recientes, por ejemplo, se ha identificado cómo, a través de los sesgos algorítmicos en la plataforma X, usuarios con el mismo perfil son expuestos a un contexto informacional diferente y esto responde a dinámicas de perfil socioeconómico y relacionamiento (popularidad, recencia y exposición). Un estudio de Andrés Abeliuk (2021) demostró diferentes sesgos presentes en la plataforma Twitter, ahora X. Ejemplos de algunos de estos en recomendaciones estaban relacionados con la popularidad de productos en la plataforma, al presentarles a los usuarios contenidos que ya son populares, lo que llevaba a una amplificación de productos y también de estereotipos y prejuicios en producciones audiovisuales —algo que deriva luego en sesgos informacionales para el caso de burbujas ideológicas algorítmicas—.

Por lo tanto, esta naturaleza social del dato implica que los *datasets* replican y amplifican sesgos históricos, culturales y estructurales preexistentes en la sociedad. Cathy O’Neil (2016) demostró cómo los modelos de *machine learning* pueden convertirse en armas de destrucción matemática al codificar la discriminación (racial, de género, económica) en sistemas automáticos de toma de decisiones (justicia penal, crédito, empleo). En consecuencia, el investigador de CSC tiene la responsabilidad ética de auditar las fuentes y los *features* utilizados para la captura de la información. Sin embargo, para el caso de las redes sociodigitales, resulta dificultoso acceder a los modos de recolección de la información dada la opacidad del modo en que se captura el movimiento digital de cada usuario, por lo que la crítica epistemológica aquí demanda pasar de una visión ingenua de la neutralidad a —al menos— un enfoque de sesgo consciente.

Causalidad, correlación y el problema de la inferencia

El tercer pilar fundamental de la crítica epistemológica a la aplicación de la ciencia de datos, IA y *big data* en las CSC se centra en la persistente dificultad de establecer inferencias causales a partir de métodos cuyo diseño estructural privilegia la correlación y la capacidad predictiva (Lipton, 2018). El *machine learning*, alimentado por el *big data*, sobresale en la identificación de patrones y en la detección de correlaciones inesperadas en volúmenes masivos de información. Sin embargo, en la investigación social aplicada y, fundamentalmente, en la formulación de políticas públicas, el conocimiento no se satisface con la mera anticipación del fenómeno (qué sucederá), sino que exige la explicación causal (por qué sucede). Sosa Escudero (2019) enfatiza esta distinción al afirmar que el *big data* es un recurso predictivo superior, pero es cauteloso al sostener que no anula ni resuelve los problemas esenciales de la inferencia causal. El investigador del Consejo Nacional de Investigaciones Científicas y Técnicas (Conicet) argumenta que la acumulación de datos (*n* grande) no subsana los desafíos metodológicos clásicos de la investigación social, como el sesgo de selección, la endogeneidad o la falta de validez interna y externa. En esencia, la cuantificación masiva de información no compensa un diseño de investigación defectuoso: si los datos están sesgados o si la relación causal se invierte, el volumen simplemente amplificará el error.

Para ilustrar de forma concreta por qué la correlación sin teoría es peligrosa, introduce el debate de la mala utilización de la correlación (Pearson: la medida del grado de relación entre una variable y otra) en series temporales con tendencia. El autor señala que la aplicación mecánica de la correlación de Pearson, una medida fundamental en la estadística, puede llevar a conclusiones absurdas o engañosas cuando las series de tiempo bajo análisis presentan una tendencia sistemática muy marcada. Si dos variables no relacionadas causalmente (como el número de premios de poesía en Chile y el precio del petróleo) simplemente crecen simultáneamente a lo largo del tiempo debido a factores exógenos no relacionados (como el crecimiento poblacional o la inflación global, respectivamente), el coeficiente de Pearson arrojará un valor alto, lo que sugeriría una fuerte correlación. En estos casos, el coeficiente de Pearson se vuelve metodológicamente inservible, ya que la asociación alta es una espuria generada por la tendencia compartida y no por una relación sustantiva. Con otros ejemplos, también sostiene que no se deben extrapolar causalidades a partir

de simples correlaciones. Este fenómeno obliga al investigador a recurrir a herramientas econométricas más complejas como la cointegración, los modelos de corrección de error o a la diferenciación de las series para eliminar el efecto de la tendencia. Este requerimiento metodológico demuestra que la vigilancia estadística y el conocimiento teórico del contexto de generación del dato son insustituibles por el volumen: el *big data* puede ofrecer de correlaciones, pero solo la teoría social puede señalar aquellas que resultan ilegítimas. Asimismo, la inferencia causal, como lo ha demostrado Judea Pearl (2018), exige la imposición de supuestos extralgorítmicos sobre la direccionalidad de las variables, supuestos que son provistos por la teoría social. El problema se agrava con el desarrollo de modelos de inteligencia artificial (IA) complejos, como el *deep learning*, que introducen el concepto de *caja negra* (Lipton, 2018), la incapacidad de comprender el proceso interno por el cual el algoritmo asigna pesos entre variables o clasifica una observación debilita su utilidad en la investigación social. La validación de una hipótesis social depende de la capacidad de justificar el mecanismo causal subyacente: si este mecanismo permanece opaco dentro del código, el hallazgo carece de valor para la acumulación de conocimiento teórico y resulta riesgoso para la implementación de políticas que requieren ser justificadas y auditadas socialmente. Los métodos de las CSC, por lo tanto, deben ser sometidos a un *test* de inferencia para asegurar que la capacidad predictiva no se convierta en el único estándar de verdad.

Finalmente, es dable destacar que la tensión entre correlación estadística y explicación sustantiva ha sido ampliamente discutida por la sociología analítica, que sostiene que la identificación de regularidades empíricas debe complementarse con la reconstrucción de mecanismos causales capaces de explicar cómo y por qué emergen determinados patrones sociales (Keuschnigg *et al.*, 2018). Desde esta perspectiva, los modelos algorítmicos funcionan como dispositivos heurísticos valiosos, pero su poder explicativo depende de su inserción en teorías que articulen niveles micro y macro, ya que, sin esta mediación conceptual, la predicción puede convertirse en una forma sofisticada de descripción que no alcanza a explicar los procesos sociales subyacentes.

Aportes metodológicos para la triangulación entre las CSC y los abordajes cualitativos

La síntesis de la crítica epistemológica desarrollada lleva a la propuesta de un marco metodológico híbrido que busca integrar enfoques computacionales con las tradiciones cualitativas y críticas de las ciencias sociales. Este marco se propone trascender la aplicación de algoritmos para enfocarse en la triangulación metodológica y en la validación contextual de los hallazgos, con el objetivo de fusionar la escalabilidad que ofrece el *big data* con la densidad interpretativa propia de las tradiciones cualitativas y críticas de las ciencias sociales (Denzin, 1970). Esto implica diseñar investigaciones que utilicen algoritmos para la identificación de patrones (por ejemplo, *topic modeling* masivo) y que luego contrasten y contextualicen esos hallazgos mediante técnicas tradicionales (como entrevistas en profundidad, etnografía digital).

El caso de análisis que se presenta en el siguiente apartado no pretende erigirse como una validación empírica cerrada, sino como un laboratorio metodológico que permite observar cómo las decisiones algorítmicas influyen en la construcción de categorías sociales. La selección de variables, métricas de distancia y criterios de agrupamiento constituye una forma de modelización que traduce realidades en estructuras computacionales. Esta operación, lejos de ser neutral, implicó en su quehacer supuestos teóricos sobre el desarrollo territorial y la desigualdad, lo que refuerza la necesidad de interpretar los resultados desde marcos sociológicos y no únicamente técnicos.

La ciencia de datos en la práctica: el caso del análisis censal de Ecuador

La solución metodológica reside en el diseño de investigaciones que utilicen los enfoques computacionales para la identificación de patrones y que luego contrasten y contextualicen esos hallazgos mediante técnicas cualitativas, un proceso que se puede ilustrar con el análisis del censo de Ecuador realizado en el marco de la Diplomatura de Posgrado en Ciencia de Datos (Aminahuel *et al.*, 2025). La fuente se trató de una base de datos convencional, pero las técnicas utilizadas para el abordaje de las variables que incluía cada observación fueron del campo de la ciencia de datos (con algoritmos del *machine learning* específicamente) con algoritmos que se detallarán en este apartado.

El trabajo analizó datos censales de Ecuador de 2010, con el objetivo de clasificar municipios, utilizando métodos de agrupamiento no supervisado para caracterizar los gobiernos locales según un conjunto de variables sociales y de acceso a servicios. Se aplicaron algoritmos de agrupamiento como K-medias y clúster jerárquico, con distintas métricas de distancia (Euclídea y Manhattan) para experimentar y determinar matemáticamente la cantidad de agrupamientos.

La determinación del número óptimo de clústeres para implementar un modelo predictivo para nuevos municipios se justificó mediante la combinación de los métodos del codo y de la silueta y se optó por una solución de tres agrupaciones o modelos —tipologías— de municipios. El análisis de los perfiles resultantes, validado a través de un análisis de varianza (Anova) para cada variable, permitió identificar tres clústeres distintivos que delimitaron los algoritmos asociados a diferentes niveles y tipos de desarrollo socioeconómico: municipios predominantemente urbanos, municipios rurales y municipios intermedios. Para demostrar la capacidad de predicción del modelo para nuevos municipios, se desarrolló y entrenó un modelo de regresión logística multinomial que demostró la aplicabilidad de los resultados para la clasificación y la toma de decisiones.

A fin de explicar con mayor claridad, se profundizará en las decisiones técnicas que se implementaron en el modelo y sus limitaciones. En primer término, la aplicación de métodos de agrupamiento no supervisado, específicamente el algoritmo de K-medias (*K-means*) y el clúster jerárquico, identificó una solución óptima de tres clústeres, a partir de los métodos del codo y de la silueta. El primer método, desde la interpretación de la figura 1, muestra que el número óptimo de clústeres es 3. El punto de inflexión es la curva, en la cual la disminución de la suma de cuadrados dentro de los grupos (SCDG) comienza a ser menos pronunciada en el codo y allí se puede apreciar que se encuentra alrededor de los 3 clústeres (mirando de arriba hacia abajo). A partir de este punto, la disminución del SCDG es menos marcada, lo que sugiere que agregar más clústeres no aportaría una mejora en la homogeneidad dentro de cada grupo o en la heterogeneidad entre grupos.

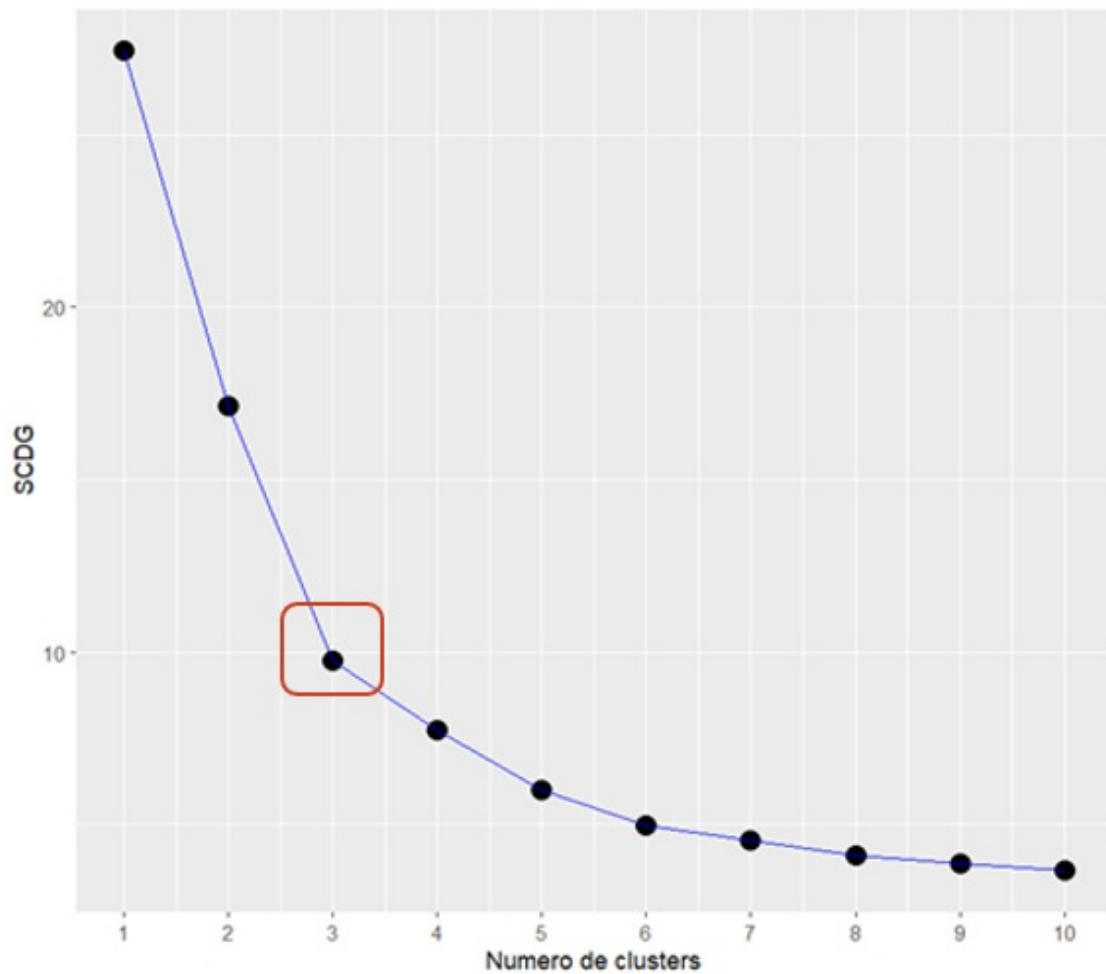


FIGURA 1.

“Método del codo”. Suma de cuadrados dentro de los grupos (SCDG)

Fuente: Aminahuel *et al.*, 2025.

Respecto a la separación de los clústeres, en la figura 2 se observa que los puntos (observaciones, o sea, municipios) que integran cada clúster están relativamente bien separados, lo que sugiere una buena agrupación. Aunque se advierte una buena separación entre el clúster 1 y 3, en el clúster 2 se observan solapamientos con los otros grupos. En este paso del análisis, también se probó separar por más clústeres (4 o 5) o por menos (2). En estas pruebas, se notó que los solapamientos se mantenían, por lo que se sostiene que, aunque hay algunos solapamientos, trabajar con 3 clústeres parecía ser un criterio adecuado para resumir la información del conjunto de datos. En cuanto a la forma, los clústeres parecen tener formas relativamente compactas.

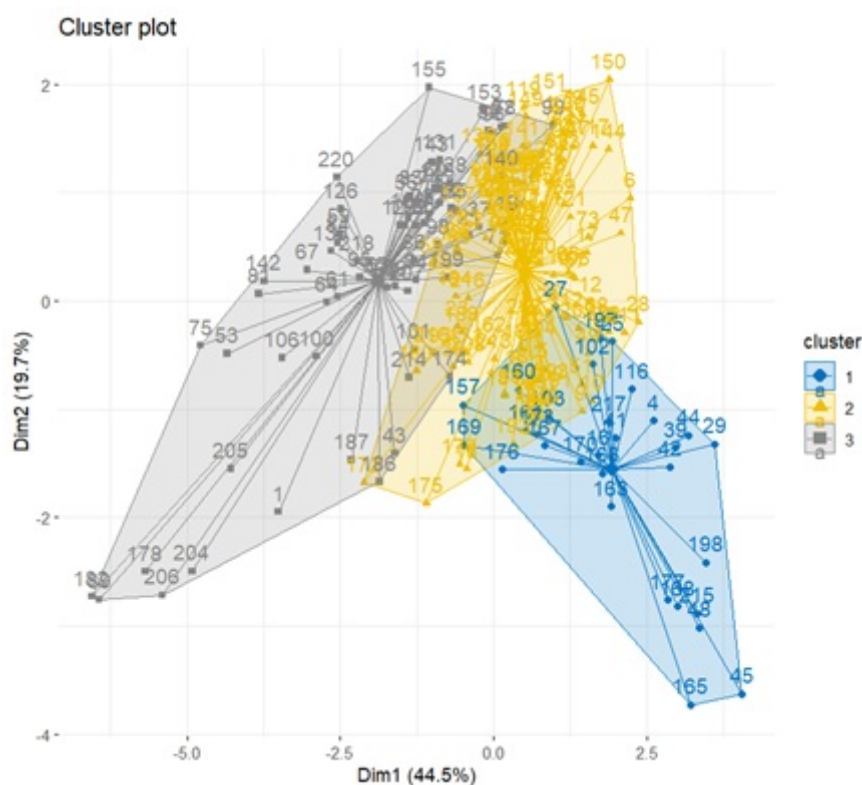


FIGURA 2.
“Disposición de los clústeres en el espacio”

Fuente: Aminahuel *et al.*, 2025.

Estos agrupamientos estadísticos que realizó el algoritmo revelaron perfiles territoriales distintivos cuando se observaron las variables que utilizó el algoritmo de manera automática para delinear los perfiles se determinaron:

- Clúster 1 (Urbano / Alto desarrollo): municipios con bajo analfabetismo y baja proporción indígena, pero con los niveles más altos de acceso a internet y telefonía celular.
- Clúster 2 (Periférico/Intermedio): municipios con un perfil más complejo, que mostró bajos niveles de población indígena (similar al clúster 1), pero con alta ruralidad y bajo acceso a internet (similar al clúster 3). A este se lo denominó *clúster de municipios intermedios*.
- Clúster 3 (Rural / Desafíos sociales): municipios con altos niveles de analfabetismo, alta población indígena y alta proporción de hogares rurales, junto con un acceso limitado a la tecnología.

La fiabilidad y precisión de esta segmentación se validó estadísticamente mediante el análisis de varianza (Anova), que demostró que la mayoría de las variables clave (como analfabetismo, población indígena y acceso a internet) presentaban diferencias promedio altamente significativas entre los grupos ($p < 0.001$). Este análisis computacional proporcionó una estructura de datos clara y validada sobre la cual se puede iniciar la fase interpretativa.

Finalmente, con el fin de demostrar la capacidad de estos clústeres para clasificar y predecir la pertenencia de nuevos municipios, se desarrolló un modelo de regresión logística multinomial con regularización Lasso. Esta técnica se seleccionó debido a su ventaja en la prevención del sobreajuste estadístico, es decir, cuando el modelo está alineado tan estrechamente con los datos de entrenamiento que no sabe cómo responder ante datos nuevos, lo cual es crítico en modelos predictivos con múltiples variables predictoras como ocurre en las bases de datos de ciencias sociales. A diferencia de un modelo de regresión tradicional que incluye todas las

variables, Lasso impone una penalización (L1) sobre la magnitud de los coeficientes de regresión, forzando a los coeficientes de las variables menos relevantes a cero. Esto implica, en la práctica, que Lasso ofreció una herramienta adicional para comprender cuál era el impacto real de cada variable en la diferenciación-definición de los perfiles de municipios o, dicho de otro modo, qué variable incidía más en la determinación de los clústeres. Este modelo fue evaluado utilizando validación cruzada de 10 pliegues. Los resultados mostraron una precisión predictiva alta, lo que se corroboró con una matriz de confusión que exhibió errores de clasificación mínimos. Esta alta precisión reafirmó la validez de los clústeres y demuestra un potencial práctico como posiblemente efectivo para la clasificación de nuevos/otros municipios que no fueron usados para entrenar el modelo. No obstante, la buena precisión del modelo predictivo, aunque es un indicador positivo de la separabilidad-definición de los clústeres, sugiere la posibilidad de que exista una separación casi perfecta entre los grupos en el espacio de las variables consideradas. Si bien esto simplifica la clasificación, también puede implicar que el modelo, en un entorno de datos real y con nuevas variables, podría no mantener exactamente el mismo nivel de rendimiento —principalmente al comprender que se trata de variables sociales—.

Limitaciones algorítmicas y criterios de aplicación

El trabajo expuesto demuestra de forma práctica las limitaciones inherentes a la aplicación de algoritmos en datos sociales y la importancia de las validaciones metodológicas. Aunque el K-medias y el clúster jerárquico resultaron efectivos, se observaron solapamientos entre los grupos, particularmente en el clúster 2, lo que indica que las fronteras entre los tipos de municipios no son perfectamente discretas. Como conclusión del trabajo, las autoras (Aminahuel *et al.*, 2025) explicitaron que, al cambiar las métricas de distancia (Euclídea vs. Manhattan) o el algoritmo (K-medias vs. jerárquico), la composición de los clústeres variaba ligeramente, al colocar algunos municipios en un grupo y no en otro, especialmente para aquellos que se encontraban en estadios “difusos”, por ejemplo, los municipios intermedios, difíciles de clasificar. En el caso de análisis, se trataba de gobiernos locales, pero, al extrapolar esta problemática —propia de la dinámica de los datos sociales en la que existe mucha colinealidad entre variables— a otro tipo de unidades de observación, como, por ejemplo, personas, se puede aseverar que los algoritmos poseen limitaciones predictivas y, justamente por ello, las autoras del experimento debieron validar de manera cruzada, utilizar Lasso, hacer regresión multinomial, etc., para evitar el sobreajuste del modelo.

En ese sentido, la implementación del modelo predictivo con regresión logística multinomial demostró una precisión sorprendentemente alta (0.9911) en la clasificación de los municipios. Si bien esto validaba la fuerte separabilidad de los clústeres en el espacio de variables consideradas, la conclusión resultó doblemente cautelosa:

- Riesgo de sobreajuste (separación perfecta): una precisión tan elevada puede sugerir una separación casi perfecta de los grupos, lo que simplifica la clasificación. Sin embargo, en un entorno de datos reales futuros donde las variables sociales son inherentemente cambiantes, el modelo podría no mantener este nivel de rendimiento, lo que requiere cautela en la generalización.
- La aplicación del Lasso fue fundamental para refinar el modelo. Al forzar a cero los coeficientes de variables menos relevantes, Lasso no solo previno el sobreajuste, sino que también identificó las variables socioeconómicas más esenciales para la diferenciación: la ruralidad para el clúster 1 y la población indígena y el analfabetismo para el clúster 3. Esto constituyó un modelo predictivo más interpretable.

En suma, la descripción del experimento realizado explicita que el agrupamiento predictivo con algoritmos de *machine learning* para datos sociales presenta dificultades evidentes para la toma de decisión. Esto requiere

una construcción metodológica que depende más de las elecciones del analista (métricas, número de clústeres, criterio de enlace, conocimiento de la base de datos, delimitación de tipologías para casos intermedios).

Propuesta de integración cualitativa

Para evitar el reduccionismo inherente al dataísmo, los perfiles definidos por los algoritmos deben ser dotados de significado social, político e histórico. La triangulación metodológica sugeriría las siguientes acciones para profundizar los hallazgos obtenidos en el experimento:

- Contextualización de los clústeres extremos: realizar estudios de caso cualitativos en municipios representativos del clúster 3 (Rural/Indígena). Esto implicaría emplear la observación participante y entrevistas en profundidad *in situ*. El objetivo no sería validar la estadística, sino comprender las razones estructurales detrás del bajo acceso tecnológico y educativo, explorando la agencia comunitaria y los contextos sociopolíticos (como la presencia de políticas de desarrollo específicas o los efectos de la historia de la tierra) que el censo no captura.
- Análisis documental cualitativo de las políticas públicas: el clúster 2 (Periférico/Intermedio), con su mezcla de baja población indígena y alta ruralidad, es el más ambiguo. Se podría realizar un análisis documental cualitativo sobre las políticas públicas o planes de desarrollo territorial implementados en esos municipios. Esto revelaría si su perfil intermedio es resultado de una transición económica, de políticas de conectividad focalizadas o de dinámicas migratorias que la regresión logística multinomial no puede interpretar causalmente.
- Validación con expertos y referentes territoriales: los resultados de la modelización (los tres clústeres y las variables clave) deberían ser presentados a actores locales, funcionarios públicos y académicos expertos en desarrollo ecuatoriano. Su validación cualitativa (¿estos tres grupos tienen sentido en la realidad de Ecuador?) aseguraría que la estructura estadística descubierta tenga resonancia y utilidad política, para verificar si el agrupamiento se alinea con las categorías administrativas o identidades territoriales preexistentes.

La superación del reduccionismo inherente al dataísmo exige la implementación de una triangulación metodológica que eleve el valor interpretativo de los hallazgos algorítmicos. En la práctica, esto implica un proceso de hibridación consciente en el cual los resultados del *machine learning* sean utilizados como la estructura inicial para la investigación cualitativa. Por ejemplo, la segmentación de municipios ecuatorianos mediante algoritmos de agrupamiento (K-medias o clúster jerárquico) identifica clústeres estadísticos —perfiles poblacionales discretos—, pero son la teoría social y el abordaje cualitativo los que les otorgan significado social, político e histórico. Esta combinación se puede materializar a través de inmersiones en la agencia social (mediante entrevistas en profundidad y observación participante en los clústeres extremos) y el análisis documental (estudiando las políticas públicas que definen los perfiles ambiguos), para así capturar las dinámicas de desigualdad, contexto y agencia que los datos censales no pueden expresar.

De esta forma, la combinación de técnicas permite que la escalabilidad del *big data* se fusione con la densidad interpretativa del *small data*. Solo esta sinergia entre el patrón estadístico y la explicación contextual —lo que algunos denominan la construcción de una *comprensión densa*— garantiza que las CSC sean un campo con potencia investigativa al servicio de las preguntas fundamentales de las ciencias sociales.

Sugerencias metodológicas: algoritmos con mayor concordancia

La experiencia del análisis censal de Ecuador permite proyectar una serie de sugerencias metodológicas orientadas a fortalecer la integración entre las técnicas de inteligencia artificial, aprendizaje automático y los

abordajes cualitativos clásicos de las ciencias sociales. Se propone un marco flexible que pretende articular las capacidades exploratorias de los algoritmos con la profundidad interpretativa de los métodos cualitativos en un diálogo continuo entre los datos y los contextos sociales que los producen.

En primer lugar, los resultados derivados de modelos de *machine learning* —como los algoritmos de agrupamiento (K-medias, clúster jerárquico) o la regresión logística multinomial— pueden ser comprendidos como dispositivos de descubrimiento. Estos permiten identificar regularidades, tipologías o zonas de interés que luego deben ser resignificadas mediante la investigación cualitativa y la teoría social. Los clústeres, lejos de constituir verdades empíricas acabadas, funcionan como hipótesis heurísticas que orientan la selección de casos o territorios para la observación participante, la realización de entrevistas en profundidad o el análisis documental. La lógica de esta integración no es confirmatoria, sino comprensiva: las técnicas cualitativas permiten indagar los factores históricos, institucionales, culturales o políticos que los algoritmos no alcanzan a representar.

En segundo lugar, la relación entre ambos enfoques debe concebirse como un proceso iterativo y dialógico, en el cual las fases cuantitativa y cualitativa se retroalimentan. Los resultados de campo —por ejemplo, los relatos de actores locales o los documentos de políticas públicas— pueden revelar dimensiones omitidas por el modelo o reinterpretar la importancia de ciertas variables, eso abre la posibilidad de ajustar y reentrenar los modelos computacionales. Esta dinámica circular promueve una triangulación genuina en la que los algoritmos cumplen una función de detección de patrones, mientras las metodologías cualitativas aportan densidad teórica y contexto sociohistórico.

En el contexto de la investigación social aplicada, la elección algorítmica y los procedimientos de validación deben orientarse a la interpretabilidad y la solidez estadística. Para trabajar con datos sociales, se recomienda la siguiente concordancia:

- *Clustering* (agrupamiento): la combinación de K-medias (rápido y eficiente) y el clúster jerárquico (que genera el dendrograma para visualizar la estructura de los datos) es ideal para la fase exploratoria. No obstante, siempre deben ser validados con métodos que evalúen la calidad del agrupamiento (como la silueta o el método del codo).
- Modelos predictivos (clasificación): la regresión logística multinomial es altamente recomendable porque, a diferencia de los modelos de caja negra (como ciertas redes neuronales), ofrece interpretabilidad a través de sus coeficientes. Esto permite al científico social entender cómo contribuye cada variable a la predicción, al vincular el modelo matemático con la teoría social (Sosa Escudero, 2019).
- Selección de variables: el uso de técnicas de regularización como Lasso es esencial para modelos con múltiples predictores. Su capacidad para reducir a cero variables redundantes o irrelevantes resulta en un modelo final más parsimonioso y robusto, lo cual destaca las características sociales verdaderamente discriminantes.

Para asegurar la autenticidad de los resultados, las siguientes validaciones resultan fundamentales, dependiendo del caso:

- Validación estadística externa: después del agrupamiento no supervisado, el análisis de varianza (Anova) es crucial para confirmar que las diferencias de medias entre los clústeres son estadísticamente significativas en las variables discriminantes. Esto asegura que el *software* no solo ha encontrado grupos, sino que esos grupos tienen perfiles promedio genuinamente distintos.
- Validación predictiva: para los modelos de clasificación, el uso de la validación cruzada (*k-fold cross-validation*) es vital. Esta técnica evalúa la generalizabilidad del modelo. Esto garantiza que su rendimiento no se debe al sobreajuste sobre los datos de entrenamiento, sino que se mantendría en un nuevo conjunto de datos.

- Validación de la opacidad: la inclusión de métodos que promuevan la explicabilidad (como Lasso) garantiza que el modelo final no sea una caja negra, sino una herramienta interpretable que pueda ser auditada y utilizada de manera responsable en la formulación de políticas públicas.

A modo de resumen, en la tabla 1 se formaliza la propuesta, a la vez que se incluyen combinaciones de algoritmos y técnicas de *machine learning*, ciencia de datos —como el caso implementado— e IA en combinación con técnicas y métodos cualitativos.

TABLA 1.
Algoritmos con mayor concordancia con técnicas cualitativas para ciencias sociales

Objetivo analítico	Técnica de ML / ciencia de datos / IA	Técnica cualitativa asociada	Propósito de la combinación
Identificar patrones o tipologías sociales	<i>Clustering</i> (<i>K-means</i> , jerárquico, DBSCAN)	Estudios de caso; observación participante	Contextualizar clústeres y comprender causas estructurales
Detectar variables explicativas clave	Regresión multinomial con Lasso	Entrevistas a actores y expertos	Validar sentido social de variables seleccionadas
Explorar discursos en grandes corpus	NLP / <i>topic modeling</i> / <i>sentiment analysis</i>	Análisis del discurso y narrativo	Conectar patrones digitales con significados culturales
Validar resultados predictivos	Validación cruzada, <i>fairness metrics</i>	Talleres participativos	Evaluar la legitimidad social de los modelos
Contrastar diferencias entre grupos	Anova / pruebas no paramétricas	Comparación de casos	Interpretar diferencias desde contextos sociales
Modelar mecanismos sociales	Simulación basada en agentes / modelos generativos	Reconstrucción de procesos sociales	Explorar cómo interacciones micro generan patrones macro
Analizar estructuras relacionales	Análisis de redes complejas	Etnografía relacional	Comprender dinámicas de poder, circulación y vínculos
Estudiar dinámicas temporales	Series temporales / modelos secuenciales	Historias de vida / análisis longitudinal	Vincular cambios cuantitativos con trayectorias sociales
Detectar sesgos algorítmicos	Auditoría algorítmica / <i>fairness ML</i>	Análisis crítico institucional	Identificar impactos sociales del modelado
Explorar escenarios sociales	Modelos predictivos / simulaciones contrafácticas	Deliberación participativa	Evaluar consecuencias sociales de decisiones modeladas
Reflexividad del dato	<i>Data provenance</i> / análisis de <i>datasets</i>	Metodología crítica del dato	Reconocer condiciones sociales de producción de datos

Fuente: elaboración propia en diálogo con propuestas de triangulación metodológica y modelización en ciencias sociales computacionales (Denzin, 1970; Cioffi-Revilla, 2014).

La experiencia analítica presentada evidencia que los métodos de las ciencias sociales computacionales operan más como herramientas de descubrimiento que como mecanismos explicativos autosuficientes. El agrupamiento algorítmico identifica regularidades formales, pero su significado social emerge únicamente

cuando se lo sitúa en contextos históricos, institucionales y culturales. Este hallazgo dialoga con debates contemporáneos de la sociología digital y analítica que advierten sobre el riesgo de reducir la complejidad social a patrones cuantificables, ya que el valor epistemológico de la modelización reside más en su capacidad para abrir preguntas interpretativas que en clausurar explicaciones. Reconocer estas limitaciones metodológicas exige su integración crítica dentro de marcos teóricos complementarios. Finalmente, la transparencia en las decisiones técnicas, la explicitación de supuestos y la triangulación cualitativa emergen como condiciones necesarias para evitar reduccionismos epistemológicos.

Consideraciones finales

La articulación entre técnicas cualitativas e IA, *machine learning* y ciencia de datos en el marco de las CSC abre un campo fértil para construir una comprensión más densa y situada de los patrones descubiertos algorítmicamente. Mientras los modelos de *clustering* o regresión permiten identificar regularidades estructurales y agrupamientos latentes en los datos, las aproximaciones etnográficas, documentales o participativas pueden reintroducir las dimensiones histórica, simbólica y política que las métricas por sí solas no pueden captar. En este sentido, la combinación metodológica no busca validar la estadística, sino dialogar con ella: los algoritmos pueden orientar la mirada hacia casos o zonas de interés y las técnicas cualitativas pueden devolverle al dato su densidad social, al mostrar cómo los procesos de diferenciación territorial o exclusión digital (en los municipios, por ejemplo) se viven, se narran y se negocian cotidianamente.

En última instancia, esta integración puede dar lugar a un modelo híbrido de investigación, en el que las técnicas computacionales funcionen como dispositivos de descubrimiento y las técnicas cualitativas, como herramientas de interpretación social. Este enfoque mixto permite superar la dicotomía entre explicación estructural y comprensión contextual para así avanzar hacia un conocimiento más reflexivo sobre los procesos de territorialización de la desigualdad. De esta manera, los algoritmos dejan de ser cajas negras para convertirse en interlocutores epistemológicos dentro de una conversación más amplia sobre los modos de producir verdad en la era de los datos.

Para concluir, se considera que la investigación aplicada en CSC debe operar bajo un principio de conciencia dual: aprovechar el poder de la IA y el *big data* para la predicción a escala, pero siempre someter sus resultados a la validación estadística y a la interpretación cualitativa, regresando siempre a la teoría social para asegurar que la tecnología esté al servicio de resolver las problemáticas sociales de nuestro siglo.

Referencias

- Abarzúa Cutroni, N. (2022). El Big Data como fenómeno social, político y cultural: reflexiones epistemológicas desde las ciencias sociales. *Revista de Estudios Culturales Latinoamericanos*, 27(2), 157-176.
- Abeliuk, A. (2021). *Algorithmic bias in social media recommendations: a quantitative analysis of X (ex-Twitter)*. Universidad de Chile.
- Aminahuel, A., Donadio, F. y Musso, M. (2025). *Clasificación de municipios ecuatorianos mediante métodos de clustering y regresión logística multinomial con Lasso* [Trabajo final de diplomatura]. Universidad Nacional de Río Cuarto, Argentina.
- Aminahuel, A. y Rodríguez, M. (2024). Notas críticas sobre políticas de comunicación y regulaciones en el capitalismo de plataformas en América Latina. *Revista Correspondencias y Análisis*, 19, 92-117. <https://ojs.correspondenciasy analisis.com/index.php/Journalcya/article/view/497>
- Anderson, C. (2008). The end of theory: The data deluge makes the scientific method obsolete. *Wired Magazine*. <https://www.wired.com/2008/06/pb-theory/>

- Barbu, D. (2014). Microtargeting and political marketing in the age of Big Data. *Journal of Political Communication*, 31(4), 423-440.
- Castells, M. (1996). *La era de la información. Economía, sociedad y cultura. Vol. I: La sociedad red*. Alianza Editorial.
- Celis Bueno, C. (2021). Sesgos maquínicos y política de los datos. *Revista Chilena de Estudios Culturales Digitales*, 15(2), 1-17.
- Cioffi-Revilla, C. (2010). Computational social science. *Wiley Interdisciplinary Reviews: Computational Statistics*, 2(3), 259-271. <https://doi.org/10.1002/wics.95>
- Cioffi-Revilla, C. (2014). *Introduction to computational social science: Principles and applications*. Springer.
- Crawford, K. (2022). *Atlas de la inteligencia artificial. Poder, política y costos planetarios*. Fondo de Cultura Económica.
- Day Montarcé, M. (2021). *Geomática y análisis espacial en las ciencias sociales contemporáneas*. Universidad Complutense de Madrid.
- Denzin, N. K. (1970). *The research act: A theoretical introduction to sociological methods*. Aldine.
- Durand, C. (2022). *Tecnofeudalismo. Crítica de la economía digital*. La Cebra Ediciones y Kaxilda.
- Everaert-Desmedt, N. (2004). *Charles Sanders Peirce: el pensamiento simbólico y la semiótica contemporánea*. Fondo de Cultura Económica.
- Fumagalli, A. (2010). *Bioeconomía y capitalismo cognitivo*. Traficantes de Sueños.
- Gitelman, L. (2013). *Raw data is an oxymoron*. MIT Press.
- Goodfellow, I., Bengio, Y. y Courville, A. (2016). *Deep learning*. MIT Press.
- Gualda, E. (2022). Métodos mixtos y triangulación en la investigación social contemporánea. *Revista Española de Sociología*, 31(3).
- Gualda Caballero, E., Taboada Villamarín, A. y Rebollo Díaz, C. (2023). 'Big data' y ciencias sociales: una mirada comparativa a las publicaciones de antropología, sociología y trabajo social. *Gazetade Antropología*, 39(1), artículo 09. <https://doi.org/10.30827/Digibug.79779>
- Hastie, T., Tibshirani, R. y Friedman, J. (2015). *The elements of statistical learning: Data mining, inference, and prediction*. Springer.
- Joler, V. y Pasquinelli, M. (2021). The nooscope manifested: AI as instrument of knowledge extractivism. *AI & Society*, 36(1), 1263-1280.
- Keuschnigg, M., Lovsjö, N. y Hedström, P. (2018). *Analytical sociology and computational social science*. *Journal of Computational Social Science*, 1(1), 3-14. <https://doi.org/10.1007/s42001-017-0006-5>
- King, G. (2011). Ensuring the data-rich future of the social sciences. *Science*, 331(6018), 719-721.
- Lazer, D., Pentland, A., Adamic, L., Aral, S., Barabási, A., Brewer, D., Christakis, N., Contractor, N., Fowler, J., Gutmann, M., Jebara, T., King, G., Macy, M., Roy, D. y Van Alstyne, M. (2009). Computational social science. *Science*, 323(5915): 721-723. <https://www.sciencemag.org/lookup/doi/10.1126/science.1167742>
- Lipton, Z. C. (2018). The mythos of model interpretability. *Communications of the ACM*, 61(10), 36-43.
- Marres, N. (2017). *Digital sociology: The reinvention of social research*. Polity Press.
- Mejías, U. A. y Couldry, N. (2019). Colonialismo de datos: repensando la relación de los datos masivos con el sujeto contemporáneo. *Virtualis*, Joler, V. y Pasquinelli, M. (2021). The nooscope manifested: AI as instrument of knowledge extractivism. *AI & Society*, (18), 78-97. <https://www.revistavirtualis.mx/index.php/virtualis/article/view/289>
- Morales, S. (2019). Derechos digitales y regulación de internet: aspectos claves de la apropiación de tecnologías digitales. En S. Morales (ed.), *Tecnologías digitales: miradas críticas de la apropiación en América Latina* (pp. 35-50). Siglo del Hombre Editores. <https://www.jstor.org/stable/j.ctvt6rmh6.5>
- Nohlen, D. (2007). *Instituciones políticas en su contexto: las virtudes del método comparativo*. Rubinzal-Culzoni Editores.
- O'Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown Publishing. https://www.researchgate.net/publication/314165204_Cathy_O'Neil_Weapons_of_Math_Destruction

- ruction_How_Big_Data_Increases_Inequality_and_Threatens_Democracy_New_York_Crown_Publishers_2016_272p_Hardcover_26_ISBN_978-0553418811
- Parra Saiani, J. y Piovani, J. (2021). *Métodos mixtos en investigación social: entre la estadística y la interpretación*. Universidad Nacional de La Plata.
- Pearl, J. (2018). *The book of why: The new science of cause and effect*. Basic Books.
- Rodríguez, M. (2019). Estar *online* para influir *offline*: comunicación y participación ciudadana en comunidades sociodigitales en Argentina y Brasil. En M. Zerega y M. Gómez Coglianó (coords.), *Mundos digitales: paradojas de la vida digital* (pp. 20-41). Universidad Casa Grande. https://issuu.com/casagrande55/docs/libro_mundos_digitaes_paradojas_de_la_vida_digita
- Rouvroy, A. (2013). El/Los fin(es) de la crítica: Conductismo de datos versus debido proceso. En M. Hildebrandt y K. De Vries (eds.), *Privacidad, debido proceso y el giro computacional*. Routledge.
- Rouvroy, A. y Berns, T. (2013). Gouvernamentalité algorithmique et perspectives d'émancipation. *Réseaux*, 177(1), 163-196.
- Rouvroy, A. y Berns, T. (2016). Gubernamentalidad algorítmica y perspectivas de emancipación. ¿La disparidad como condición de individuación a través de la relación? (E. Feuerhake, trad.). *Adenda Filosófica*, (1), 88-116. <https://www.academia.edu/30732187/>
- Sadin, E. (2020). *La inteligencia artificial o el desafío del siglo. Anatomía de un antihumanismo radical*. Caja Negra Editora.
- Schab, M. (2020). *Vigilancia masiva y manipulación de la opinión pública* [Trabajo final de especialización]. Universidad de Buenos Aires, Argentina. http://bibliotecadigital.econ.uba.ar/download/tpos/1502-1700_SchabJI.pdf
- Sosa Escudero, W. (2019). *Big data*. Siglo XXI Editores.
- Srnicek, N. (2019). *Capitalismo de plataformas*. Caja Negra Editora.
- Turkle, S. (2011). *Alone together: Why we expect more from technology and less from each other*. Basic Books.
- Van Dijck, J. (2016). *La cultura de la conectividad*. Siglo XXI Editores.
- Weber, M. (2014). *Economía y sociedad: esbozo de sociología comprensiva*. Fondo de Cultura Económica.

Notas

- * Artículo de investigación
- 1 El autor de referencia sostiene que, al analizar la producción académica sobre estudios comparados en América Latina, se observa una mayor predominancia de estudios cualitativos.
- 2 Traducción propia, versión *online* original del artículo disponible en Keuschnigg *et al.* (2018).
- 3 Problemática actual de los modelos de inteligencias artificiales generativas mercantiles.

Licencia Creative Commons CC BY 4.0

Acerca de la autora

Aminahuel Aimé es docente investigadora de la Universidad Nacional de Río Cuarto, becaria posdoctoral del Consejo Nacional de Investigaciones Científicas y Técnicas de Argentina (Conicet-AR), ha cursado el Posdoctorado en Inteligencia Artificial Aplicada a las Ciencias Sociales en la Facultad de Ciencias Económicas de la Universidad Nacional de Córdoba. Es doctora en Administración y Política Pública de la Universidad Nacional de Córdoba, tiene un posgrado en Ciencia de Datos de la Universidad Nacional de Río Cuarto y es licenciada en Ciencia Política de la misma institución.

Cómo citar: Aimé, A. (2025). Ciencia de datos, *big data* e inteligencia artificial en las ciencias sociales computacionales. Crítica epistemológica y aportes para la investigación. *Signo y Pensamiento*, 44. <https://doi.org/10.11144/Javeriana.syp44.cdbd>