

## COMPARISON & ELECTION BETWEEN AI MODELS: HETERODOX ANALYSIS OF CONDITIONS

*Antonio Sánchez-Bayón, Moshe Yanovskiy y Yehoshua Socol*

### Resumen

Revisión heterodoxa sobre el desarrollo en la implementación cotidiana de los algoritmos de inteligencia artificial, sus sesgos y riesgos para la vida humana por falta de transparencia o efecto caja negra. Se centra la atención aquí en la evaluación de un problema que ha influido en el desarrollo de la inteligencia artificial, como es el dilema ético-económico de caja negra, junto con su paradoja. La atención a dicho problema, sobre si prima la transparencia sobre el rendimiento algorítmico (y cómo se valora, con sus sesgos y riesgos), ello permite comprender la paradoja conducente a la actual dicotomía entre el mundo anglosajón y europeo continental. Mediante un estudio bibliométrico-narrativo y crítico-hermenéutico, junto con los marcos teóricos y metodológicos de la Escuela Austriaca y los Neoinstitucionalistas (dada su experiencia en el análisis de otras cajas negras, como el Estado, el Sector público o la economía de bienestar), desde este trabajo se ofrece una exposición y explicación del problema, su alcance y si cabe esperar una futura convergencia de posiciones al respecto.

**Palabras clave:** inteligencia artificial; transparencia; sostenibilidad; dilema de caja negra; sesgos y riesgos; enfoques heterodoxos.

### Abstract

A heterodox review of the development of everyday implementations of artificial intelligence algorithms, their biases, and risks to human life due to a lack of transparency or the black box effect. This paper focuses on the evaluation of a problem that has influenced the development of artificial intelligence: the ethical-economic dilemma of the black box, along with its paradox. Focusing on this problem—whether transparency prevails over algorithmic performance (and how it is valued, with its biases and risks)—allows us to understand the paradox that leads to the current dichotomy between the Anglo-Saxon and continental European worlds. Through a bibliometric-narrative and critical-hermeneutic study, along with the theoretical and methodological frameworks of Austrian Economics and Neo-Institutional Economics (given their experience in analyzing other black boxes, such as the State, the public sector, and welfare economics), this paper offers an exposition and explanation of the problem, its scope, and whether a future convergence of positions on the matter can be expected.

**Keywords:** artificial intelligence; transparency; sustainability; black-box dilemma; biases and risks; heterodox approaches.

**JEL:** A14, B5, O3, P16, Z1

## **Introduction**

Artificial intelligence (AI), as it is known today, has its origins in the collaboration between university professors and the military for encryption work (i.e., Enigma Project: Government Code & Cypher School, with the participation of Turing). After World War II, its development began in universities with public funding (Huang et al., 2023; Gofman & Jin, 2024; Neumann et al., 2024). A popular example, because it was the official origin of the term AI itself (in 1956, at Dartmouth College, with public funding), was the "Dartmouth Summer Research Project on Artificial Intelligence," an event organized by scholars such as McCarthy, Minsky, Rochester, and Shannon (Doroudi, 2023). The term AI encompasses several fields (with attention paid here to the relationships between economics, engineering, and applied ethics). It is often used to refer to the ability of machines to imitate human cognitive functions: learning from experience, adapting to new tasks, or performing functions such as image, voice, or sound recognition; language translation; etc., even decision-making (LaGrandeur, 2024; Singla, 2024). AI is based on algorithms and models that allow computers to process information and solve problems autonomously (Tan et al., 2024). However, this has several implications, which differ depending on the model used for its development. Initially, after World War II, the situation was similar in the West, with the study of AI promoted at the university level and state research centers (i.e., machine learning and the Turing test, 1950). However, after decades of slow progress (effectively, but not efficiently), with the arrival of globalization and the intensification of digitalization (Sánchez-Bayón, 2020 and 2021), companies began to take an interest in their development and applications (Sánchez-Bayón, 2025a). Thus, a division emerged between the European model (inspired by public interventionism, open source and a focus on transparency and ethics) versus the American model (driven by private initiative, without open source and oriented towards efficient results). To understand why the American model (with its private entrepreneurship and business orientation) has prevailed since November 2022 (Sánchez-Bayón et al., 2024a-b and 2025), it is necessary to analyze how the scientific community has addressed the issue, the great

debate of which has been framed in the following terms: AI algorithms are progressively implemented in image processing, natural language processing (NLP), clinical decision support, law enforcement, and other areas. At the same time, many concerns were raised regarding ethical issues and potential risks of applied AI to solving life-related problems (Awad et al, 2018; Benkler, 2019; Biller-Andorno and Biller, 2019; Ngiam and Khor, 2019; Sánchez-Bayón, 2025b), with hundreds of papers published on this topic, especially during the last few years (before 2022 and the breaking point for the US model with LLMs: ChatGPT, Gemini, Grok, etc., Xie & Avila, 2025). In this work, we address one of the frequently mentioned concerns: many AI algorithms are “black boxes” so that the user has no idea why the machine chooses one solution or another (Benkler, 2019; Petkus et al., 2020; Wang et al., 2023; Marcus and Teuwen, 2024). The concern about being a black box is fully applicable to deep neural networks (DNNs) with a large number of parameters (up to 108 and more). However, other AI methods (logistic regression, decision tree, support vector machines, etc.) are generally considered incomparably more transparent (the first example of AI is probably to be traced back to radar-based proximity fuses during World War II). While DNNs became true game changers in fields, such as NLP and image processing, which superiority in other fields (i.e., in biomedical research and applications) is questioned (Wang et al., 2023). The advantage of non-DNN AI models as more transparent is questioned by some authors. For example, Lipton (2017) scrupulously examines various aspects of human understanding of the work of an AI system. The author considers the transparency of the AI system at different levels: the entire model (simulable, Teufel et al., 2023; Chen et al., 2023; Chaudhary, 2024), the individual components (decomposable), and the training algorithm (algorithmic transparency, Grimmelikhuijsen, 2023; Cheong, 2024) and explains how key components of the AI system relate to human understanding of how the system obtains its results. For interpreting the results, which could further contribute to the “informativeness” of the entire system, should also be considered (Romanova, 2025). Lipton argues that although simpler non-DNN models have more understandable algorithms, the entire class of these models does not demonstrate an obvious advantage. The author also shows that analysis by human experts fits the definition of a “black box” fairly well. Regardless,

the public perception of non-DNN ML methods as more transparent can cause a significant advantage in their competition with DNNs. Therefore, this perception itself can become a major advantage, possibly somewhat outweighing the performance degradation, as "transparent" algorithms are seen as much more compatible with "human-involved" solutions. In this work, thanks to the analysis of narrative bibliometrics (Jahin et al, 2023) and the empirical illustration of the theorems of Austrian economics (with their translation to the classroom, improving learning, Alonso et al, 2024; Sánchez-Bayón, 2015), we propose to verify and quantify the differential hypothesis (between the US model and the EU model), that a "black box" reputation is widely considered as a major drawback of AI algorithms, directly influencing implementation opportunities (especially in life and health-related sectors).

## **Materials and Methods**

This study is based on heterodox approaches (Sánchez-Bayón, 2020 and 2025c), which apply analytical elements from: a) narrative bibliometrics (Torres, 2023; Rivas et al., 2024), beyond the traditional systematic literature review (Tahiru, 2021; Ammar, 2025; Zhu et al., 2025); and b) the theory of the Austrian School of Economics-ASE (Menger, 2007[1871]; Huerta de Soto, 2000), such as the theorem on the impossibility of economic calculation under socialism (Mises, 2000[1922] and 1949) – currently revised by Boettke, 2000; Huerta de Soto, 2010, etc. - and some other main principles of political economy (Menger, 2007[1871]; Sánchez-Bayón, 2025c). The debate over the impossibility of economic calculation theorem is a defining element of Austrian economics thinking and has distinguished it from other schools (Huerta de Soto, 2000 and 2008; Smith, 2024); moreover, this theorem has experienced a revival under the management of the last crisis (Sánchez-Bayón et al., 2023 and 2024c). The economic calculation theorem, or the impossibility of socialism, has been discussed and applied by scholars in this heterodox tradition in a wide range of contexts and future research (i.e., public management of digitalization in the tourism industry, Sánchez-Bayón et al., 2024c). This paper focuses on the comparison between two main models: a) the American entrepreneurial model (based on private entrepreneurship, with a closed code

that prioritizes the efficiency of the results); b) the European academic model (based on public intervention, with an open code that prioritizes transparency and ethics in the processes). Likewise, the Mises theorem is related to the Menger-Hayek theorem (Menger, 2007[1871]; Hayek, 1988) on institutional evolution (American model) and constructivism (European model), and the Huerta de Soto-Sánchez-Bayón theorem (Huerta de Soto et al., 2021) on dynamic processes, entrepreneurship and well-being with empirical illustration (Alonso et al., 2024), in addition to addressing secondary effects such as the black box. According to these theorems, a heterodox interpretation of the rise of AI in 2022 is possible (Floridi, 2024; Sánchez-Bayón et al., 2025), favoring the US model (leaving behind the EU's public academic model); this paper examines whether this event was a coincidence or causality, based on these economic principles (Sánchez-Bayón, 2025c).

A search in the *Web of Science Core Collection* by the key sequence “artificial intelligence ethics” (until 2022, with the boom of USA model, Floridi, 2024; Sánchez-Bayón et al, 2025), it was founded above 600 results with one–two papers per year in 1990-2000, up to 117 papers in 2018 and 189 papers in 2019 (later was the COVID-19 boom and in 2022 the AI boom). Since our aim is to suggest recommendations for present AI development (between EU university model vs. USA business model), the modern trends are of primary importance. Therefore, we decide to limit our consideration to the papers published from the beginning of 2017 until 2022 when the search was performed (between the crisis recoveries to AI boom, Challoumis, 2024; Noncheva & Baykin, 2025). The total number of 400 papers were identified, out of them 267 in peer-review journals including highest-rank journals with impact factors 30-70 (i.e., New England Journal of Medicine, The Lancet Oncology, Nature, Science; Awad et al, 2018; Benkler, 2019; Ngiam & Khor, 2019; Biller-Andorno & Biller, 2019). We performed primary screening of the papers (using personal judgment) based on the abstracts. Out of total 400, we chose 198 papers (49%) of interest to the subject of the algorithm transparency. These papers were accessed and manually scored according to the following four-grade scale:

1. The paper does not mention the issue of algorithm’s transparency

2. Algorithm transparency problem is mentioned but not pursued
3. There is special focus on the AI transparency problem
4. The AI transparency problem principal for the article

We did not manage to formulate any formal way of scoring, so we must admit that it was somewhat subjective. We present an example of scoring, considering four articles published in the most influential journals (impact factor 30-70) and scored 1–4, respectively.

1. Awad et al. (2018) in *Nature* presents outcomes of impressive sociological research of choice of people of various cultures' in situations like the trolley problem (in the context of autonomous vehicles). The algorithm choice (transparent vs. 'black box') is not mentioned.
2. Detailed analysis of AI (machine learning) currently applied in clinical oncology (Ngiam & Khor, 2019) in *Lancet Oncology* just mentions the desirability of "doctors' understanding of how machine learning tools produce predictions."
3. Biller-Andorno and Biller (2019) in *New England Journal of Medicine* devote a special section ('Morality, Transparency, Humanity') to the transparency issue.
4. Benkler (2019) in *Nature* perceives expanding application of AI systems as a serious threat to the society and designates non-transparent ('black box') algorithms as a key problem.

The score of papers screened during the preliminary consideration was set to 1.

The 2019 impact factors (IF) of the peer-review journals were recorded according to the *Journal Citation Reports™* (JCR, 2019). For conference proceedings and other non-peer-review publications, IF=0 was set. We also recorded the number of times each publication was cited. However, taking in account that most publications are very recent, we did not consider the number of citations as a meaningful parameter. All the data processing was performed using MATLAB™ software ([www.mathworks.com](http://www.mathworks.com)). The Supplementary Material contains the table in Excel™ format (<https://>

[papers.ssrn.com/sol3/papers.cfm?abstract\\_id=5283013](https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5283013)), with data about authors, paper title, publication details (journal, volume etc.), times cited, journal IF, the transparency score 1-4, and DOI (digital object identification).

## **Results & discussion**

During the academic debates between 2017-2022 (previous to AI boom), the main concerns of the authors were:

1. Military applications
2. Autonomous vehicles
3. Legal and moral responsibility for the actions of AI systems
4. Health care AI applications. For example, Mazurowski (2020) points out the conflict of interests of some radiologists when competing the AI consulting systems.
5. Data privacy (first and foremost – in healthcare)
6. Surveillance (AI systems for facial recognition etc.) and corresponding threats of abuses and personal rights violations (see important case study by Andreeva et al, 2019).
7. Last but not least, authors are addressing public perception of AI systems as affected by ‘black box’-type algorithms. Floridi et al, 2018 (one of the most cited paper in the sample) points out the ‘explainability’ of the algorithm as a crucial factor of AI success. Cath (2018) insists that extensive governmental regulations and control are the key factors for public trust in AI systems. Dietvorst et al. (2018), they argue that producers could overcome the ‘algorithm aversion’ phenomenon (caused by lack of transparency) by providing the user with opportunity of algorithm correction, even slight.



Out of 400 papers, 291, 53, 34 and 22 were scored 1,2,3, and 4, respectively. The results are shown at Fig. 1. Only 27% of the papers on AI ethics address the issue of AI transparency. However, in our opinion this number is somewhat misleading. Namely, most papers scored “1” (actually, all but two) deal not with AI specifically but essentially with ethical issues of the society, its moral values and social order. Therefore, we conclude that out of papers properly dealing with artificial intelligence ethics (2+53+34+22), nearly all mention the transparency issue, and more than half (34+22) pursue it. It may be also meaningful to note, that papers published in journals with higher IF tend to address the issue of the AI transparency much more frequently. Fig. 2 presents median journal IF vs. transparency importance score. We can suggest therefore that algorithm transparency is a major issue in the ethical context of AI.

The list of 400 analyzed papers is available at: <https://drive.google.com/file/d/1aUyxGvwS4Hz0d717LVmRk2clnabIIIn2V/view?usp=sharing>

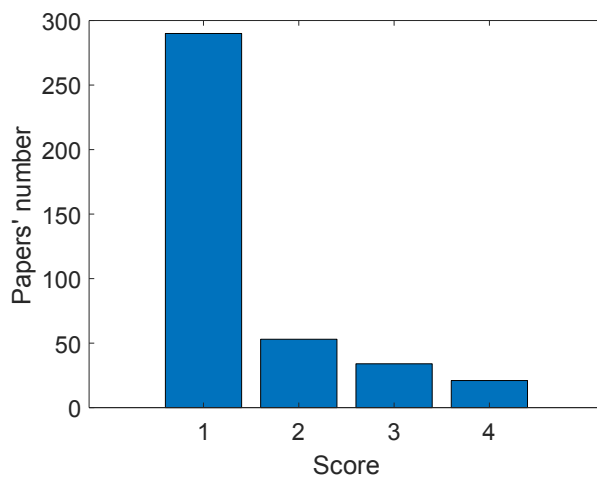


Fig. 1. Algorithm's transparency importance score distribution.  
1 – not mentioned, 2 – mentioned but not pursued, 3 – special focus, 4 – principal topic.

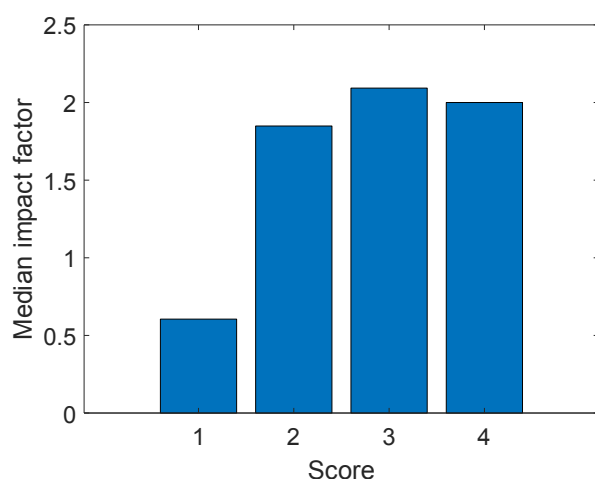


Fig. 2. Median impact factor of journals according to the transparency score of published papers. 1 – not mentioned, 2 – mentioned but not pursued, 3 – special focus, 4 – principal topic.

Although scientific and academic production is biased in favor of the university model, given the drift of the publishing system toward the knowledge industry, ultimately, funding is required. In the United States, funding comes from the business world, while in Europe, funding comes from the public sector. Thus, the debate persists and remains unresolved.

Right now, the most relevant AI model is the USA business approach, but before the AI boom in 2022, the main Western AI model was the EU university model (with public funds for research, Foffano et al, 2023), because there was a concern on ethics issues and the algorithm's transparency. There was an initiative to preserve this model in Europe (European Parliament, 2023), but there are not enough AI into the EU, and the biggest big-tech are in USA (Bollerman, 2025). Our meta-analysis confirmed that algorithm's transparency is discussed as an important topic in scientific literature dealing with ethics of AI and its acceptance by the society. Therefore, an important practical recommendation can be formulated: In life & health-related tasks, in every case where hardly interpretable AI system has no obvious performance superiority over a better interpretable, the latter should be preferred.

The black box risk persists today (Marcus & Teuwen, 2024) because the AI mainstream is focused in USA business approach and its DNNs, as game-changers in the field for

efficiency, latency, etc. However, the EU university approach is still relevant for sectors such as bioethics, because the algorithm's transparency and the simulation of human decisions are basic aspects. In economic terms, there are many opportunity costs with the USA model, because the ethic limits are requested for dignity by the Human Rights International Law (as *ius cogens* or imperative law for everybody), and it is the most security way to improve the AI (to control the AI in favor of human beings, to control the technological monopolies, etc.). In this sense, before 2022, there was a biggest concern on algorithm's transparency and its work under presumption of 'human in the loop' solutions.

For future research lines, it is intended to delve deeper into the models, with attention to special aspect, like AI university model in USA (Oh & Sanfilippo, 2025); AI education and the integration of diversity and disability, as well as analyzing which model is successful in other parts of the World (Al-Rashaida et al, 2025; Buragohain & Chaudhary, 2025; Dumitru et al, 2025).

## Conclusions

AI was developed in the academic field in the 1950s, but eventually moved into the business world with the beginning of globalization, giving rise to two contrasting models. As indicated, on the one hand, the American business model, based on private entrepreneurship, with closed source code and prioritizing efficient results, and on the other hand, the European model, based on public intervention, with open source code and prioritizing transparency and ethics in processes. With the AI boom of 2022, it might seem that the USA business model has prevailed over the EU university model, but the debate remains open: is a more efficient but closed model preferable or a more transparent model that simulates human action? In the bioethical field (the study of life and healthcare), it is key to address this model. In this sense, the most notable criticisms come from the European university model, but for it to develop more effectively, it needs to become more competitive, building bridges with the business world, rather than opposing it. This issue must be addressed as soon as possible, as the risk is greater, as evidenced by the asymmetry between American and European LLMs. Furthermore,

it's also worth extending the debate to other parts of the world to see what creative proposals they offer in this regard.

## References

- Alonso MA, Sánchez-Bayón A & Gallego-Morales D (2024). Enhancing Visual Literacy and Data Analysis Skills in Macroeconomics Education: A Beveridge Curve Analysis Using FRED® Data. In: Valls Martínez M & Montero J (eds) *Teaching Innovations in Economics* (pp. 51-76). Springer, Cham., [https://doi.org/10.1007/978-3-031-72549-4\\_3](https://doi.org/10.1007/978-3-031-72549-4_3)
- Al-Rashaida M, Moustafa A, Mohsen W (...) & Khan A (2025). AI Strategies for Inclusive Education Resources and Support, AI in *Learning, Educational Leadership, and Special Education* (pp. 349-380). IGI Global. <https://doi.org/10.4018/979-8-3373-0573-8.ch012>
- Ammar A (2025). systematic and bibliometric review of artificial intelligence in sustainable education: Current trends and future research directions, *Sustainable Futures*, 10, Article 101033. <https://doi.org/10.1016/j.sftr.2025.101033>
- Andreeva O, Ivanov V, Nesterov A & Trubnikova T (2019). Facial Recognition Technologies in Criminal Proceedings: Problems of Grounds for the Legal Regulation of Using Artificial Intelligence. *Tomsk State University Journal* 449, 201–212.
- Awad E, Dsouza S, Kim R (...) & Rahwan I (2018). The Moral Machine Experiment. *Nature*, 563, 59–64. <http://dx.doi.org/10.1038/s41586-018-0637-6>
- Benkler Y (2019). Don't let industry write the rules for AI. *Nature* 569, 161.
- Biller-Andorno N & Biller A (2019). Algorithm-Aided Prediction of Patient Preferences - An Ethics Sneak Peek. *New England Journal of Medicine*, 381(15), 1480-1485. <https://doi.org/10.1056/NEJMms1904869>
- Boettke P (2000). *Socialism and the market: The socialist calculation debate revisited*. London: Routledge
- Bollerman M (2025). Digital Sovereigns Big Tech and Nation-State Influence. arXiv. <https://doi.org/10.48550/arXiv.2507.21066>
- Buragohain D & Chaudhary S (2025). Navigating ChatGPT in ASEAN Higher Education: Ethical and Pedagogical Perspectives, *Computer Applications in Engineering Education*, 33(4). <https://doi.org/10.1002/cae.70062>
- Cath C (2018). Governing artificial intelligence: ethical, legal and technical opportunities and challenges. *Philosophical Transactions of the Royal Society A*, 376(2128), Article 20180080. <https://doi.org/10.1098/rsta.2018.0080>

- Challoumis C (2024). Charting the course-The impact of AI on global economic cycles. In *XVI International Scientific Conference* (pp. 103-127). Copenhagen: ISG Konf..
- Chaudhary G (2024). Unveiling the black box: Bringing algorithmic transparency to AI. *Masaryk University Journal of Law and Technology*, 18(1), 93-122.
- Chen Y, Zhong R, Ri N (...) & McKeown K (2023). Do models explain themselves? counterfactual simulatability of natural language explanations. arXiv <https://doi.org/10.48550/arXiv.2307.08678>
- Cheong B (2024). Transparency and accountability in AI systems: safeguarding wellbeing in the age of algorithmic decision-making. *Frontiers in Human Dynamics*, 6, Article 1421273.
- Dietvorst B, Simmons J & Massey C (2018). Overcoming Algorithm Aversion: People Will Use Imperfect Algorithms If They Can (Even Slightly) Modify Them. *Management Science* 64(3). <https://doi.org/10.1287/mnsc.2016.2643>
- Doroudi S (2023). The Intertwined Histories of Artificial Intelligence and Education. *Int J Artif Intell Educ* 33, 885–928. <https://doi.org/10.1007/s40593-022-00313-2>
- Dumitru C, Abdulsahib G, Khalaf O & Bennour A (2025). Integrating artificial intelligence in supporting students with disabilities in higher education: An integrative review, *Technology and Disability*. <https://doi.org/10.1177/10554181251355428>.
- European Parliament. (2023). *EU AI Act: First regulation on artificial intelligence*. <https://www.europarl.europa.eu/topics/en/article/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>
- Floridi L (2024). Why the AI hype is another tech bubble. *Philosophy & Technology*, 37(4), Article 128.
- Floridi L, Cowls J, Beltrametti M (...) Vayena E (2018). AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations. *Minds & Machines* 28, 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
- Foffano F, Scantamburlo T & Cortés A (2023). Investing in AI for social good: an analysis of European national strategies. *AI & society*, 38(2), 479-500.
- Gofman M, & Jin Z (2024). Artificial intelligence, education, and entrepreneurship. *The Journal of Finance*, 79(1), 631-667. <https://doi.org/10.1111/jofi.13302>

- Grimmelikhuijsen S (2023). Explaining why the computer says no: Algorithmic transparency affects the perceived trustworthiness of automated decision-making. *Public Administration Review*, 83(2), 241-262.
- Hayek F (1988). *The fatal conceit*. Chicago: The University of Chicago.
- Huerta de Soto J (2008). *The Austrian School: Market Order and Entrepreneurial Creativity*. Cheltenham: Edward Elgar Publishing Ltd.
- Huerta de Soto J (2010). *Socialism, Economic Calculation and Entrepreneurship*. Cheltenham: Edward Elgar Publishing Ltd
- Huerta de Soto J, Sánchez-Bayón A & Bagus P (2021). Principles of Monetary & Financial Sustainability and Wellbeing in a Post-COVID-19 World: The Crisis and Its Management. *Sustainability*, 13(9), Article 4655 (1-11). <https://doi.org/10.3390/su13094655>
- Jahin M, Naife S, Saha A & Mridha M (2023). *AI in supply chain risk assessment: A systematic literature review and bibliometric analysis*. arXiv. <https://doi.org/arXiv.2401.10895>
- Journal Citation Reports (2019) *Impact factor into Journal Citation Reports*. <https://clarivate.com/webofsciencegroup/solutions/journal-citation-reports>
- LaGrandeur K (2024). AI and Reverse Mimesis: From Human Imitation to Human Subjugation?. *Mimetic Posthumanism: Homo Mimeticus 2.0 in Art, Philosophy and Technics*, 5, 321.
- Lipton Z (2017). *The Mythos of Model Interpretability*. arXiv. <https://arxiv.org/abs/1606.03490>
- Marcus E & Teuwen J (2024). Artificial intelligence and explanation: How, why, and when to explain black boxes. *European Journal of Radiology*, 173, Article 111393.
- Mazurowski M (2020). Artificial Intelligence in Radiology: Some Ethical Considerations for Radiologists and Algorithm Developers *Academic Radiology* 27(1), 127-129. <https://doi.org/10.1016/j.acra.2019.04.024>
- Menger C (2007). *Principles of Economics*. Auburn: Mises Institute. (Original work published 1871)
- Mises L (2000). *Socialism: An Economic and Sociological Analysis*. Auburn: Mises Institute. (Original work published 1922)
- Mises L (1949). *Human action. A treatise on Economics*. New Haven: Yale University Press.
- Neumann O, Guirguis K, & Steiner R (2024). Exploring artificial intelligence adoption in public organizations: a comparative case study. *Public*

- Management Review*, 26(1), 114-141. <https://doi.org/10.1080/14719037.2022.2048685>
- Ngiam K & Khor W (2019). Big data and machine learning algorithms for health-care delivery. *The Lancet Oncology*, 20(5): E262-E273. [http://dx.doi.org/10.1016/S1470-2045\(19\)30149-4](http://dx.doi.org/10.1016/S1470-2045(19)30149-4)
- Noncheva D & Baykin A (2025). Innovate approaches to crisis management and economic recovery: the role of artificial intelligence. *ББК 65.9 (4Укр) 261-18я431 Ф 59*, 455.
- Oh S & Sanfilippo M (2025). Responsible AI in academia: policies and guidelines in US universities, *Information and Learning Sciences*, <https://doi.org/10.1108/ILS-03-2025-0042>
- Petkus H, Hoogewerf J & Wyatt JC. (2020). What do senior physicians think about AI and clinical decision support systems: Quantitative and qualitative analysis of data from specialty societies. *Clinical Medicine*, 20(3), 324-328. <https://doi.org/10.7861/clinmed.2019-0317>. PMID: 32414724; PMCID: PMC7354034.
- Rivas E, Núñez M, Rodríguez J & Rubio M (2024). Revisión de la producción científica sobre Storytelling mediado por tecnología entre 2019 y 2022 a través de SCOPUS. *Texto Livre*, 17, Article e51392.
- Romanova A (2025). *Analysis of Interfaces Informativeness Issues in the Development of Autonomous Artificial Intelligence Systems for Corporate Management*. <https://doi.org/10.2139/ssrn.5347689>
- Sánchez-Bayón A (2015). Filosofía del aula inteligente del S. XXI: críticas urgentes y necesarias. *Bajo Palabra*, 10, 259-269. <https://doi.org/10.15366/bp2015.10.022>
- Sánchez-Bayón A (2020). Renovación del pensamiento económico-empresarial tras la globalización, *Bajo Palabra*, 24, 293-318. <https://doi.org/10.15366/bp.2020.24.015>
- Sánchez-Bayón A (2021). The digital economy review under the technological singularity: technovation in labour relations and entrepreneur culture. *Sociología y Tecnociencia*, 11(2), 53-80. [https://doi.org/10.24197/st.Extra\\_2.2021.53-80](https://doi.org/10.24197/st.Extra_2.2021.53-80)
- Sánchez-Bayón, A. (2025a). ¿Cómo innovar en aprendizaje de gestión digital de riqueza y bienestar? Experiencia con monedas digitales socio-empresariales. *AROEC*, 8(1), 1-32
- Sánchez-Bayón A (2025b). Bioética y biojurídica: una revisión veinte años después. *Encuentros Multidisciplinares*, 27(79), 1-16



- Sánchez-Bayón A (2025c). Revisión de las relaciones ortodoxia-heterodoxia en la Economía y la transición digital. *Pensamiento*, 81(314), 523-550. <https://doi.org/10.14422/pen.v81.i314.y2025.012>
- Sánchez-Bayón A, Urbina D, Alonso-Neira MA & Arpi R (2023). Problema del conocimiento económico: revitalización de la disputa del método, análisis heterodoxo y claves de innovación docente. *Bajo Palabra*, (34), 117–140. <https://doi.org/10.15366/bp2023.34.006>
- Sánchez-Bayón, A., Alonso-Neira, M.A., Morales, D. (2024a). Aprender a emprender con IA y método de talento digital: Revisión de responsabilidad social universitaria. *Iberoamerican Business Journal*, SI 1 ( 1 ) , 4 8 – 6 3 . <https://doi.org/10.22451/5817.ibj2024.Spec.Ed.vol1.1.11094>
- Sánchez-Bayón A, Alonso MA, Miquel AB & Sastre FJ (2024b). Aprendizaje creativo e innovación docente sobre RSC 3.0, ODS y divisas alternativas. *Encuentros Multidisciplinares*, 78, 1-13
- Sánchez-Bayón A, Sastre FJ & Sánchez LI (2024c). Public management of digitalization into the Spanish tourism services: a heterodox analysis. *Review of Managerial Science*, 18(4), 1-19. <https://doi.org/10.1007/s11846-024-00753-1>
- Sánchez-Bayón A, Miquel-Burgos AB & Alonso-Neira MA (2025). Experience of learning technovation for i-entrepreneurship training: how to prepare the students for digital economy? *Estrategia y Gestión Universitaria*, 13(1), Article e8765. <https://doi.org/10.5281/zenodo.14908364>
- Singla A (2024). Cognitive Computing Emulating Human Intelligence in AI Systems. *Journal of Artificial Intelligence General Science (JAIGS)*, 1(1), Article e38. <https://doi.org/10.60087/jaigs.v1i1.38>
- Smith A, Walsh J, Long J (...) Fisher C (2020). Standard machine learning approaches outperform deep representation learning on phenotype prediction from transcriptomics data. *BMC Bioinformatics*, 21, Article 119. <https://doi.org/10.1186/s12859-020-3427-8>
- Smith D (2024). *Austrian Economics*. Piamonte: Amazon Italy.
- Tahiru F (2021). AI in education: A systematic literature review. *Journal of Cases on Information Technology (JCIT)*, 23(1), 1-20.
- Tan K, Wu J, Zhou H, Wang Y & Chen J (2024). Integrating advanced computer vision and ai algorithms for autonomous driving systems. *Journal of Theory and Practice of Engineering Science*, 4(01), 41-48.
- Teufel J, Torresi L & Friederich P (2023). Quantifying the intrinsic usefulness of attributional explanations for graph neural networks with artificial

simulatability studies. In *World Conference on Explainable Artificial Intelligence* (pp. 361-381). Cham: Springer Nature Switzerland.

Torres D (2023). Entre métricas y narraciones: definición y aplicaciones de la Bibliometría Narrativa. *Anuario ThinkEPI*, 17. <https://doi.org/10.3145/thinkepi.2023.e17a30>

Turing A (1950) Computing Machinery and Intelligence. *Mind* 49, 433-460

Wang L, Chen X, Zhang L, Li L, Huang Y, Sun Y, Yuan X, Davis C, Henstock P, Hodge P, Maciejewski M, Mu X, Ra S, Zhao S, Ziemek D & Fisher C. (2023). Artificial intelligence in clinical decision support systems for oncology. *International Journal of Medical Sciences*, 20(1), 79-86. <https://doi.org/10.7150/ijms.77205>.

Xie Y & Avila S (2025). The social impact of generative LLM-based AI. *Chinese Journal of Sociology*, 11(1), 31-57.

Zhu H, Sun Y & Yang J (2025). Towards responsible artificial intelligence in education: a systematic review on identifying and mitigating ethical risks, *Humanities and Social Sciences Communications*, 12(1). <https://doi.org/10.1057/s41599-025-05252-6>.