

Alertness to AI Algorithmic Biases in the Healthcare System

Rufin SIME^a, Clovis FOGUEM^{b,c}, Samuel FOSSO WAMBA^d, Bernard KAMSU-FOGUEM^a

^aUTTOP (Université de technologie Tarbes Occitanie Pyrénées), 47 Avenue Azereix, Tarbes 65016, France

^bAuban Moët Hospital (Groupement Universitaire de Champagne), 137 Rue de l'Hôpital Auban - Moët, Épernay 51200, France

^cMETRICS (Université de Lille) HAUB, 1 Place de Verdun, Lille 59045, France

^dTBS Business School, 1 Place Alphonse Jourdain, Toulouse 31068, France

ABSTRACT

Artificial intelligence (AI) is assisting practitioners in performing certain tedious tasks with computer algorithms, particularly those involving large volumes of data, bringing recent progress in medicine. Technical and social biases can undermine confidence in these algorithms, so the algorithms still require human validation. Especially when the approach by which AI systems are made can cause ethical concern. This work aims to review how the concepts of biases are integrated in different countries. Particularly, the approach of the ethical dimension and regulation that France, the United States (US), and China have when developing AI should be studied to understand better how it can lead to biases and ethical interrogations. The purpose is to assess the basis on which algorithms are created in the present healthcare system. Results show (i) a necessity to enrich AI systems, increase their explainability, and recommendations. (ii) Countries' approaches to AI's opportunities align with their beliefs.

Keywords: Artificial intelligence (AI); Deep Learning; Machine Learning; Bias; Ethic, Ethical Issues Social Bias, Technical Bias; Health; Medical Field

I. INTRODUCTION

Artificial intelligence (AI)'s increase in usage because of its ability to improve a significant variety of areas to the work of humans has come with its fair share of advantages and disadvantages (Lecun et al., 2015). In fact, AI algorithms are often created with biases that can hinder its efficiency and efficacy (Zhang et al., 2004). While there is an increase in AI models being used in healthcare such as the AI models Food and Drug Administration (FDA) approved to screen diseases like diabetic retinopathy (Shah et al, 2024). The ethical and explainable aspects of medical AI can be a central issue in the approval of AI models. Artificial intelligence development in said fields must go through varying laws, rules, decrees, acts, et, before being allowed for general use.

According to the Stanford University Human-Centered Artificial Intelligence 2025 report, the US and China are the main leaders in AI development and research. In 2024, the US produced forty notable AI models, fifteen were produced by China, and three were produced by France (Europe) (Maslej et al., 2025).

The US and China are leaders in this field, although the two countries have a non-convergent approach to data management. Indeed, the US liberal approach to the treatment of data and China state controlled approach to the treatment of data and development of AI. While these visions seem to enable better AI model development, they do not prioritize the

equity and ethical properties of the algorithm, especially in fields such as healthcare. For instance, there does not seem to exist any study in which AI has been used to detect rare diseases. That is why we are interested in applying AI to the detection and care of rare diseases. A disease is considered rare if it affects approximately 1 in 2000 people (Haendel, M., et al., 2020). The small size of data sample for rare diseases make detecting them a tedious job for physicians. Especially with the diverse manifestations of a rare disease which can lead to misdiagnosis and difficulty identifying the correct disease. The application of AI to rare diseases detection could potentially lead to a broader early detection of a disease thanks to AI's access to a larger set of information than physicians.

The presence of biases in AI is a serious problem in healthcare that needs to be addressed. Misdiagnosis can lead to an increase in disease detection time that can go from several months to several years in some cases. While research has been done on the biases associated with AI and their mitigation, the increase in AI usage to a broader and more diverse population does not seem to have been taken into account in its potential to lead to biases and ethical issues in AI implementation for medicine.

That is why this work aims to review laws and regulation that are the basis for the creation and use of AI algorithms in the US (Algorithmic accountability act), China (Cybersecurity Law, Personal information protection law and algorithmic regulation framework) and France (Decree n°2017-330 on the 17th of March 2024).

II. MATERIAL AND METHODS

Three countries were selected to review the impact of laws and regulations on AI development. The US was selected because it is the leading country for AI research and model production. China was selected because it is second only to the US as far as AI research and development is concerned. France was selected because it is the country in which this research is directed and it was also the third leader in notable AI models created in 2024 as shown in the figures below (Maslej et al., 2025).

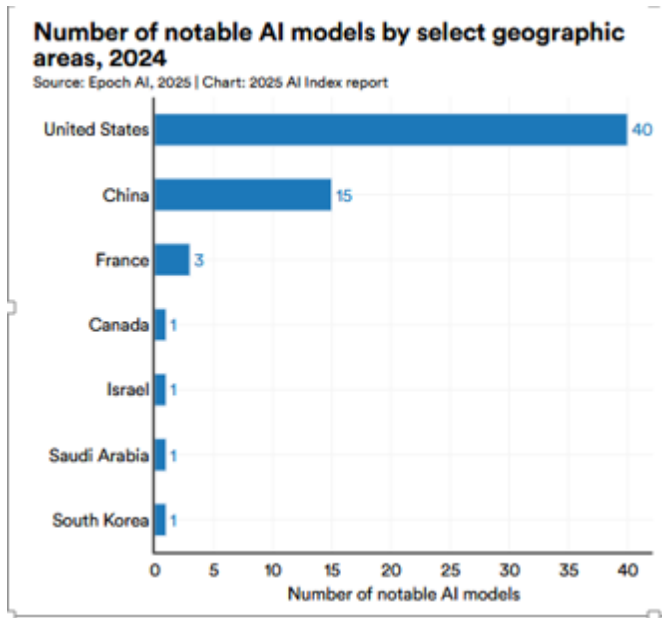


Figure 1 : Notables AI models by geographic area. Source: (Maslej, et al., 2025)

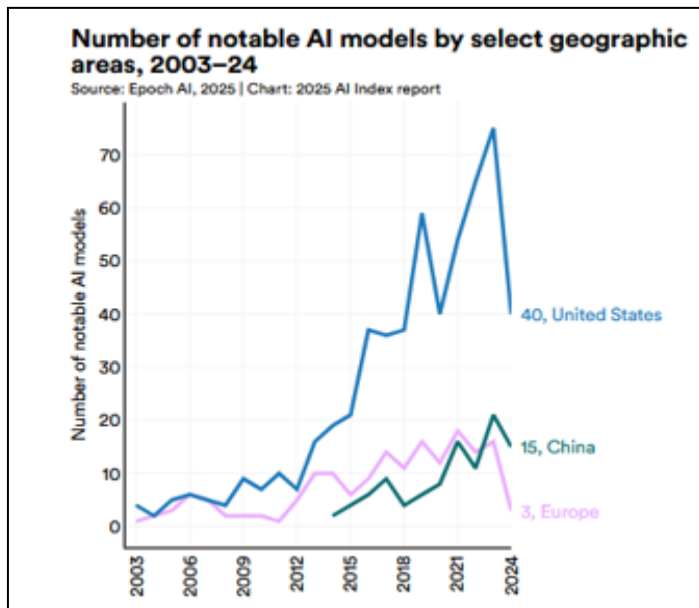


Figure 2: Notables AI models by year and geographic are Source: (Maslej, et al., 2025)

The main laws, decrees, and acts selected for these countries were those that targeted AI development. The algorithmic accountability act (US, 2022), The Cybersecurity Law (China, 2017), The Personal Information Protection Law (China, 2021) and the Decree n°2017-330 of March 17th 2024 (France, 2024). These are some of the main concepts that shape the use, collection and application of data and AI algorithm development. The review of these will provide a view on the steps shaping AI biases and ethical aspects.

In healthcare, biases can stem from ethical issues in two related health sectors. For example, in the US some states allow

collecting data without patients' consent or knowledge which could pose an ethical concern. There are also biases in the representation of the overall population in which some subsets can be overrepresented or underrepresented, related to disparities in data collection among the population.

III. RESULT

In this section, an analysis of the review results of the different regulations is proposed. The US and China are leading in the race to AI development overall despite their drastically different visions. However, advancing AI must be guided by a more equitable and ethical approach.

The US is the leader in AI technologies research and development. The laws that they abide to are fairly liberal and lax which allows for the rapid development of AI and the easy access to data. These advantages also come with drawbacks such as the use of users' personal data without their explicit consent. For example, this was case in the Cambridge Analytica scandal in which Facebook users data were collected and used without their consent for political purpose. This approach overall, often does not therefore favor the production of equitable and ethical AI.

The republic of China on the other hand has a government encouraged AI development and research approach. This has enabled China to rapidly develop AI over just a few years (going from below Europe in 2015 to above Europe in 2023, in less than 10 years). The laws regarding AI development and use tend to be more state oriented than citizen oriented. This could, in some cases, lead to ethical and freedom issues, such as the obligation for Chinese citizens to have social media with their full name and identification, for the government to control the type of content users post. This approach can also be associated to restriction or even control of freedom similar to Bentham's concept of panopticon. A concept designed by Jeremy Bentham (philosopher and social theorist) a few centuries earlier (1791) as a revolutionary type of prison in which prisoners are theoretically always under the watch of guards. An approach that was associated with prison and that did not seem to respect human rights.

France and the European Union used a more ethical and transparent approach to AI development and requirements through the General Data Protection Regulation and the European Union Artificial Intelligence Act. While it does not permit the boom of AI in the same dimension as the US or China, France's approach seems to be the one that would be the most equitable and ethical.

The following table will show the main ideas in the shaping of creating and developing AI systems in different countries.

Table 1: Algorithm regulation in the main countries of notable AI

<i>Countries of regulation</i>	<i>Types of Algorithmic regulation</i>		
	France Decree n°2017-330 of March 17th 2024 (2024)	US Algorithmic accountability act (2022)	China Cybersecurity Law (2017) Personal information protection Law (2021)
	Algorithmic contribution in the decision making process	Conducting impact risk-assessments for algorithms.	Data Localization and protection
	The data used by the algorithm and their sources	Evaluating discrimination, bias, or any negative societal effects.	Critical Infrastructure Protection
	The parameters used for the decision making process and they are used to reach the decision	The bill would apply to companies using algorithms for hiring, credit scoring, and criminal justice applications, among other things.	Real name registration system
	The operations realized by the process		

IV. DISCUSSION

AI bias in the health sector can significantly impact patient care and outcomes (Cross, Choma and Onofrey, 2024). Biases in AI systems arise from skewed training data, where certain demographic groups, such as ethnic minorities or rare conditions, are underrepresented. This can lead to inaccurate predictions, misdiagnoses, or suboptimal treatment recommendations (Muralidharan et al., 2024). Furthermore, technical biases, such as flawed data collection or incomplete datasets amplify disparities (Ferrara, 2024),(Gorenc, 2025). Social biases among healthcare providers also influence how AI recommendations are interpreted. There may also be biases inherent in people's behaviors, attitudes or adherence to different treatment options. To ensure fairness and efficacy, AI systems in healthcare must be rigorously tested, diverse in training data, and regularly audited for biases to prevent harm. One solution to correct algorithmic biases is through regulation (Institut de Montaigne, 2020). While there is an advance in national laws or continental law to contain certain biases. The absence of an international unified regulation can allow for a technical leeway that could create bias from one country to another while following said countries regulatory

system. For instance the US has a state by state regulatory approach that makes each US state responsible for its regulation of AI application as it sees fit.

From a regulatory standpoint, the application of AI in China raises ethical questions. Indeed, it reflects the concept of the panopticon. The panopticon describes a system of surveillance where individuals are always visible to a central authority, without the inmates knowing whether or not they are being watched. The power dynamics are built on the idea that the possibility of being observed at any time leads to self-regulation, essentially internalizing control. However this does not breed well from an ethical standpoint, an ethical AI should have Beneficence Nonmaleficence, Autonomy, Justice, Explicability (Flordi & Cows, 2019). The integration of AI into China's societal control mechanisms closely mirrors Bentham's idea of the panopticon, it creates an environment of constant surveillance, shaping behavior and promoting self-regulation. This combination of technology, policy, and social pressure embodies a form of soft authoritarianism, where control is maintained not through brute force, but through the pervasive, invisible power of AI systems that govern and monitor citizens' lives. While the development of AI in China

has been fast tracked it can be seen as built on a huge ethical dilemma.

From a literature standpoint, five of the most relevant articles written by researchers on the topic of AI biases that are looking into the adoption and trust of AI systems show the importance of AI explainability, (i) to mitigate biases and ethical disparities at different levels of AI development. (ii) to improve trust, hence adoption of AI algorithm. The first article written by (Isaac, Wang, Napper & Marsh, 2024) was about highlighting the salience of biases in the decision-making process, particularly regarding age-appropriateness in the adoption of AI integration. The second article was written by (Liu, Liu, Zheng, Xu, & Wang, 2025), it highlights the importance of AI explainability to improve the implementation and adoption of AI for healthcare professionals, particularly in post-treatment understanding of AI care. The third article was written by (Stutz et al, 2025), it emphasizes the problem of standard methods for evaluating medical AI, which do not take into account the decision-making biases of healthcare experts and can therefore make the results obtained by AI biased. The fourth article was written by Albahri et al (2023). It focuses on trust and explainability in medical AI, while showing key challenges, evaluating current methods, and providing actionable insights to guide more ethical, robust, and reliable AI implementations in healthcare. The fifth article was written by Dai et al. (2025), which highlights a challenge in the adoption of medical AI in healthcare in China. Despite high awareness and optimism, actual use of medical AI remains low among Chinese medical staff, suggesting the need for targeted system improvements, risk reduction, and supporting infrastructure to bridge the adoption gap.

In the case of the US, which is at the opposite spectrum of China, it is also possible to see ethical issues stemming from the regulation. AI regulation in the US has also raised a wide range of ethical issues, from bias and discrimination to privacy violations and the potential for social unrest due to job displacement. The absence of strong, unified regulations means that ethical considerations often take a backseat to technological progress and corporate interests. Balancing the rapid development of AI with responsible, ethical regulations remains one of the country's greatest challenges. Without addressing these ethical concerns head-on, AI could deepen social inequalities and cause unintended harm.

V. CONCLUSION

This paper aimed to present the cultural aspects of AI development and how they can lead to bias, with a focus on the medical field. Some countries have an easier time developing significant AI models than others. While the US and China are at opposite ends of the spectrum of AI application (Liberal approach versus state-controlled approach) in society, we observed in both countries a faster and greater advancement of AI than in Europe. However, causal links cannot be asserted despite the advantages these approaches have shown in the development of AI; however they also raise some ethical concerns. In the US, the freedom given to companies can allow them to exploit citizens' information and data without their consent from one state to another. In China, the state controlled approach of AI development makes it so that Chinese citizens'

action are always under the scope of the government. It can lessen the freedom of action and impede freedom of speech. That is why, an international regulation on AI application could be a step to reduce disparities within AI developing countries regarding their ethical and equity properties (Partow-Navid & Slusky, 2023). Especially in the case of medical AI, which can have global risks associated with it. Global risks that require international communication to be mitigated (Batool, Zowghi, & Bano, 2025). Some of the main information to take out from this research on the topic of AI implementation in healthcare are:

- 1- AI is different between the US, China, and Europe (it means that there is a cultural bias in the development of AI).
- 2- The Bias created by these cultures can lead to ethical issues in the data storage, treatment and communication according to the laws applied for AI in these countries
- 3- There is a need for a uniformed regulation to reduce the disparity between biases and allow for ethical awareness not to promote the Bentham panopticon concept.
- 4- To preserve equity and human rights it is primordial to adapt AI regulations to prevent a panopticon society in which everybody's actions are observed through AI use.

Nevertheless, further studies on the subject are warranted.

REFERENCES

1. A.S. Albahri, Ali M. Duham, Mohammed A. Fadhel, Alhamzah Alnoor, Noor S. Baqer, Laith Alzubaidi, O.S. Albahri, A.H. Alamoodi, Jinshuai Bai, Asma Salhi, Jose Santamaria, Chun Ouyang, Ashish Gupta, Yuantong Gu, and Muhammet Deveci. (2023). A systematic review of trustworthy and explainable artificial intelligence in healthcare: Assessment of quality, bias risk, and data fusion. *Inf. Fusion* 96, C (Aug 2023), 156–191. <https://doi.org/10.1016/j.inffus.2023.03.008>
2. Blumenthal-Barby, J. S., & Krieger, H. (2015). Cognitive biases and heuristics in medical decision making: a critical review using a systematic search strategy. *Medical decision making : an international journal of the Society for Medical Decision Making*, 35(4), 539–557. <https://doi.org/10.1177/0272989X14547740>
3. Cross, J. L., Choma, M. A., & Onofrey, J. A. (2024). Bias in medical AI: Implications for clinical decision-making. *PLOS digital health*, 3(11), e0000651. <https://doi.org/10.1371/journal.pdig.0000651>
4. Dai Q, Li M, Yang M, Shi S, Wang Z, Liao J, Li Z, E W, Tao L, Tang Y. (2025). Attitudes, Perceptions, and Factors Influencing the Adoption of AI in Health Care Among Medical Staff: Nationwide Cross-Sectional Survey Study. *J Med Internet Res* 2025;27:e75343. URL: <https://www.jmir.org/2025/1/e75343>. DOI: 10.2196/75343
5. Decree n° 2017-330 of 14 march 2017. (2017). French Republic official journal. 24th march 2017
6. Haendel, M., Vasilevsky, N., Unni, D., Bologa, C., Harris, N., Rehm, H., Hamosh, A., Baynam, G., Groza, T., McMurphy, J., Dawkins, H., Rath, A., Thaxton, C., Bocci, G., Joachimiak, M. P., Köhler, S., Robinson, P. N., Mungall, C., & Oprea, T. I. (2020). How many rare diseases are there?. *Nature reviews. Drug discovery*, 19(2), 77–78. <https://doi.org/10.1038/d41573-019-00180-y>
7. Institut des montaigne. (2020). Algorithmes : contrôle des biais. [white paper]. <https://www.institutmontaigne.org/ressources/pdfs/publications/algorithms-controle-des-biais-svp.pdf>
8. Lecun, Y., Bengio, Y., Hinton, G. (2015). Deep Learning. *Nature*. 521. 436-444. <https://doi.org/10.1038/nature14539>
9. Liu, Y., Liu, C., Zheng, J., Xu, C., & Wang, D. (2025). Improving Explainability and Integrability of Medical AI to Promote Health Care Professional Acceptance and Use: Mixed Systematic Review. *Journal of medical Internet research*, 27, e73374. <https://doi.org/10.2196/73374>

10. Maslej, N., Fattorini, L., Perrault, R., Gil, Y., Parli, V., Kariuki, N., Capstick, E., Reuel, A., Brynjolfsson, E., Etchemendy, J., Ligett, K., Lyons, T., Manyika, J., Niebles, J., Shoham, Y., Wald, R., Walsh, T., Hamrah, A., Santarlasci, L., Lotufo, J.B., Rome, A., Shi, A., & Oak, S. (2025). Artificial Intelligence Index Report 2025. ArXiv, abs/2504.07139.
11. Mathew S. Isaac, Rebecca Jen-Hui Wang, Lucy E. Napper, and Jessecae K. Marsh. (2024). To err is human: Bias salience can help overcome resistance to medical AI. *Comput. Hum. Behav.* 161, C (Dec 2024). <https://doi.org/10.1016/j.chb.2024.108402>
12. National People's Congress. (2021). Personal Information Protection Law of the People's Republic of China." China Law Translate, 20 Aug. 2021, <https://www.chinalawtranslate.com/en/Personal-Information-Protection-Law/>
13. Shah, S. A., Sokol, J. T., Wai, K. M., Rahimy, E., Myung, D., Mruthyunjaya, P., & Parikh, R. (2024). Use of Artificial Intelligence-Based Detection of Diabetic Retinopathy in the US. *JAMA ophthalmology*, 142(12), 1171–1173. <https://doi.org/10.1001/jamaophthalmol.2024.4493>
14. Standing Committee of the National People's Congress. (2017). Cybersecurity Law of the People's Republic of China. 7 Nov. 2016, [effective June 1, 2017], DigiChina
15. Stutz, D., Cemgil, A. T., Roy, A. G., Matejovicova, T., Barsbey, M., Strachan, P., Schaekermann, M., Freyberg, J., Rikhye, R., Freeman, B., Matos, J. P., Telang, U., Webster, D. R., Liu, Y., Corrado, G. S., Matias, Y., Kohli, P., Liu, Y., Doucet, A., & Karthikesalingam, A. (2025). Evaluating medical AI systems in dermatology under uncertain ground truth. *Medical image analysis*, 103, 103556. <https://doi.org/10.1016/j.media.2025.103556>
16. U.S Congress. Algorithmic Accountability Act of 2022. (2022)
17. Zhang, J., Patel, V. L., Johnson, T. R., & Shortliffe, E. H. (2004). A cognitive taxonomy of medical errors. *Journal of biomedical informatics*, 37(3), 193–204. <https://doi.org/10.1016/j.jbi.2004.04.004>
18. Ferrara, E. (2024). Fairness and Bias in Artificial Intelligence: A Brief Survey of Sources, Impacts, and Mitigation Strategies. *Sci*, 6(1), 3. <https://doi.org/10.3390/sci6010003>
19. Gorenc, N. (2025). AI embedded bias on social platforms. *International Review of Sociology*, 35(2), 317–336. <https://doi.org/10.1080/03906701.2025.2489056>
20. Muralidharan, V., Adewale, B. A., Huang, C. J., Nta, M. T., Ademiju, P. O., Pathmarajah, P., Hang, M. K., Adesanya, O., Abdullateef, R. O., Babatunde, A. O., Ajibade, A., Onyeka, S., Cai, Z. R., Daneshjou, R., & Olatunji, T. (2024). A scoping review of reporting gaps in FDA-approved AI medical devices. *NPJ digital medicine*, 7(1), 273. <https://doi.org/10.1038/s41746-024-01270-x>
21. Floridi, L., & Cowls, J. (2019). A Unified Framework of Five Principles for AI in Society. *Harvard Data Science Review*, 1(1). <https://doi.org/10.1162/99608f92.8cd550d1>
22. Batool, A., Zowghi, D. & Bano, M.(2025). AI governance: a systematic literature review. *AI Ethics* 5, 3265–3279. <https://doi.org/10.1007/s43681-024-00653-w>
23. Partow-Navid, P., & Slusky, L. (2023). The Need for International AI Activities Monitoring. *Journal of International Technology and Information Management (JITIM)*, 31(3), 14. <https://iima.org/wp/jitim/>