

# ***Human-Centered Orientation: A Path to Rebuilding the Core Value of Artificial Intelligence***

**Jiahui Yuan**

*Southwestern University of Finance and Economics, Chengdu, China  
18999418840@163.com*

**Abstract.** The current development of AI is heavily influenced by the computationalist paradigm, which reduces complex social situations to digital calculations and logical reasoning. This makes it hard to understand human emotions, mental states, and value judgments behind data. As a result, it causes ethical and social problems like algorithmic bias, poor explainability, and difficulty in value alignment. This paper criticizes the data-centered AI path and suggests shifting to human-centered AI. By analyzing the deep problems of computationalism in ontology, epistemology, and axiology, it explains the core meaning and theoretical framework of human-centered AI and explores practical paths from needs analysis to technical implementation. The study shows that the essence of intelligence is understanding, not just calculation. The ultimate value of technology is to improve human well-being. Only by placing human needs, emotions, and dignity at the center of technological development can we unify technical rationality with human values and build a future in which humans and AI coexist in harmony.

**Keywords:** Human-centered AI, Algorithm ethics, Algorithmic bias, Explainable AI, Value alignment

## **1. Introduction**

China emphasized in the \*Shanghai Declaration on Global AI Governance\* that we should "strengthen AI-related network security, enhance the safety and reliability of systems and applications. While respecting international and domestic legal frameworks, jointly combat manipulation of public opinion, fabrication and spread of false information, promote the development and adoption of ethical guidelines and norms for AI with broad international consensus, guide the healthy development of AI technology, and prevent its misuse, abuse, or malicious use [1]." AI technology has moved from labs to the center of social life. With computing power as support and large models as the core, AI has achieved unprecedented efficiency in information processing and content generation, greatly expanding human capabilities and playing an important role in daily life. AI demand will continue to grow, and this growth is sustainable. AI development has entered a deeper stage of inclusive access [2]. However, as AI is deeply used in high-risk areas like medical diagnosis, judicial decisions, public governance, and financial risk control, its limitations become more obvious: most AI systems operate with opaque decision-making

logic, decisions are hard to explain, they lack common-sense knowledge of the real world, and they face serious challenges in value alignment.

The root problem is that mainstream AI still works mainly through digital calculation and formal logic, so it often fails to deal well with human emotion, cultural background, social norms, and other contextual factors. Because of this, it regularly misreads metaphors, sarcasm, and culture-specific situations, which leads to judgments that do not match ordinary human understanding. The problem is practical. As AI use has expanded, public anxiety has also increased. Errors, data leaks, bias, and discrimination continue to appear, and these problems damage trust in technology. When algorithmic errors occur in high-risk fields, the consequences may be irreversible for people's health and property. At the same time, fears such as "AI will replace humans" and "technology will dominate society" have spread and weakened people's sense of agency and self-worth.

Against this background, more people have started to question technical rationality and argue for the return of human values in technological development. Greater attention is now being given to lived experience and emotional care. Academic discussion has also turned back to the relation between technology and human beings, often drawing on the philosophical work of Heidegger and Stiegler. International organizations such as the EU have issued ethical guidelines that include ideas like "trustworthy AI" and "human-centered AI" in their development frameworks. This points to a broader shift in AI development: the focus is no longer only on improving algorithmic performance, but also on making technology serve human needs and keeping human beings in the guiding position within human-AI relations.

Current AI systems, which are built on digital calculation and logical reasoning, can process information and generate decisions. Yet in complex social settings, they cannot adequately interpret the emotions, mental states, and value judgments that sit behind the data. What they handle is data itself, not the full human situation from which that data comes. As a result, they often miss the needs of the people represented by the data. This issue becomes especially clear in human-centered industries, where "human well-being is placed at the center of the production system process, providing a safe, comfortable, and motivating work environment for learning and growth, thus forming a positive cycle where human well-being and productivity reinforce each other [3]."

Starting from this limitation, this paper asks: How can we shift the basic logic of AI systems from "what can machine algorithms do" to "what do people need"? The analysis critiques the computationalism AI paradigm and considers possible directions for theoretical reconstruction and technical innovation. The aim is to bring technical rationality into closer alignment with human values and to develop a human-centered model of AI.

## **2. Critique and reflection: three major problems of computationalism**

### **2.1. Ontological problem: human existence reduced to data symbols**

Computationalism sees the whole world as a quantifiable, calculable data set. From this view, humans are no longer individuals with rich emotions, complex intentions, and vivid life experiences. Instead, they are broken down and abstracted into labeled data points, like browsing history, buying behavior, social traces, and health indicators. This simplification of human nature means AI systems cannot truly understand human existence, leading to fundamental misunderstandings.

As social platforms become more diverse and complex, user emotions play an increasingly important role in online public opinion communication [4]. From a phenomenological perspective, human existence is situated and embodied "being-in-the-world." This includes three aspects: first, "being in the world" (worldliness); second, "existing as" (the ontological meaning of "being"); third,

the "who" that undertakes this being-in-the-world (existence itself). Heidegger emphasized that this structure cannot be understood through the traditional "substance-attribute" model, but should be grasped through ontology. Emotions, intuition, creativity, and other core traits cannot be reduced to a single calculation logic. A simple "like" on social media can contain multiple emotions like affection, sympathy, politeness, or agreement. A silent record of conversation might hide complex feelings like anger, sadness, or deep thought. Computationalism can only analyze statistical correlations in data, but cannot understand the human meaning behind the data. And correlation is not causation, let alone an understanding of human nature. This ontological mismatch means AI can never "see" real people, which is the fundamental reason it struggles to fit into social life and often makes adaptive errors.

## **2.2. Epistemological problem: technical reason hides multiple ways of knowing**

The epistemological foundation of computationalism is logical positivism and instrumental reason, pursuing absolutely certain, quantifiable knowledge, with the belief that "if it's not calculable, it's not knowledge." But humans know the world in many more ways. Besides logical reasoning and scientific evidence, non-formal ways of knowing like storytelling, metaphor, intuition, and empathy are crucial for social interaction and value judgments. These types of knowing are often vague and context-dependent, cannot be coded into algorithms, yet they are essential for human society to function.

In the digital age, algorithmic power reshapes international language order in the form of "algorithmic colonialism." Linguistic technology hegemony creates systematic control through NLP, corpus colonization, and cognitive manipulation, posing serious challenges to China's national cultural security and language sovereignty [5]. When computationalism uses a single, formalized technical reason to handle social problems, it inevitably conflicts with diverse human ways of knowing. For example, in public opinion monitoring, AI can quickly count keyword frequencies but cannot understand the social emotions and critical attitudes behind satirical posts. In literary analysis, AI can recognize grammar rules but cannot grasp the emotional core behind metaphors and symbols. This epistemological blindness makes AI clumsy and rigid when dealing with real-world problems full of uncertainty and value conflicts. If technical reason expands without limit, without the balance of diverse ways of knowing, it will eventually form a new kind of "cognitive hegemony," denying forms of knowledge that algorithms cannot understand and cutting the cognitive connection between technology and human society.

## **2.3. Axiological problem: the gap between algorithm design and value-ladenness**

Computationalism often packages algorithms as value-neutral, objective tools, denying that they carry any value tendencies. But in fact, from development to application, AI is infused with the developers' value assumptions and will to power. Each step—selecting training data, setting model goals, choosing evaluation metrics, defining application scenarios—contains subjective value judgments. These hidden value-ladenness directly affects social fairness and individual rights in AI applications. Algorithmic bias is a direct product of this problem.

When algorithms are given too much authority and used in decisions like credit scoring, job hiring, crime risk prediction, and resource allocation, their hidden biases and injustices become legitimized. For example, hiring algorithms trained on historical data can inherit and amplify workplace gender and regional biases. Crime risk prediction algorithms can produce unfair judgments against minority groups due to biased data sampling. Even worse, algorithm decisions

lack transparency. When individuals face unfair treatment by algorithms, they cannot know the basis for the decision or easily appeal. Human dignity and autonomy are weakened by cold calculation logic, and AI becomes a tool that reinforces social inequality and suppresses individual rights.

### **3. Human-centered AI: meaning and theoretical framework**

In response to the multiple problems of computationalism, human-centered AI (HCAI) has emerged as an alternative paradigm. Its core idea is to take human needs, abilities, and well-being as the fundamental starting point and goal of AI development and application, and to rebuild the value logic of human-AI relationships. In essence, behind the transformation of public opinion governance in the AI era, technological change plays a decisive role. Viewing public opinion governance as an organic whole, based on an abstract analysis of the technological changes behind the transformation and the governance difficulties they bring, proposing innovative governance suggestions can help break the complexity of public opinion governance, avoid getting lost in specific problems, and move beyond talking about governance just for the sake of governance [6].

#### **3.1. Theoretical basis: from technological determinism to techno-humanism**

The theoretical basis of human-centered AI can be traced to techno-humanism in the philosophy of technology. This view rejects both uncritical faith in technological determinism and an extreme anti-technology position. Instead, it treats technology as something that should serve human needs and expand human freedom. It also draws on Don Ihde's phenomenological account of the "human-technology-world" relation, which argues that technology is not a neutral instrument. It mediates how humans encounter the world. In this sense, human-centered AI is meant to develop mediating tools that strengthen human autonomy, deepen perception, and extend human abilities, rather than autonomous agents that replace humans [6].

Research in human-computer interaction (HCI) also gives human-centered AI a practical basis. Early work on "user-centered design" later shifted toward context-sensitive design, and more recent HCAI work places human well-being at the center. Taken together, this development suggests that the main goal of technology is not only efficiency. It is also to improve lived experience and quality of life. This serves as a practical guideline for HCAI.

#### **3.2. Preserving human agency: humans as final decision-makers**

Take art theory as one example. Traditional art theory usually gives weight to human emotion, subjective experience, and cultural background, but these points are being reworked in AI art. AI art may keep some aesthetic standards from traditional art, or it may form new systems of evaluation, which challenges art concepts built around human subjectivity [7].

Within a human-centered AI paradigm, human agency remains the basic condition. AI is treated as a "co-pilot" that supports human decision-making rather than a "pilot" that takes over the human role. This matters most in fields such as healthcare, justice, and finance, where the final decision carries serious consequences. In these settings, a "human-in-the-loop" mechanism is needed so that final control stays with people. In medical diagnosis, for instance, AI can provide image analysis, compare similar cases, and offer reference suggestions. The final diagnosis and treatment plan, however, still need to be made by the doctor in light of the patient's condition and clinical judgment. The same applies in judicial trials. AI may help with case retrieval and evidence sorting, but the verdict remains the judge's responsibility.

Preserving human agency also means respecting the user's right to choose. Users should be able to refuse AI recommendations or decisions, know the main reasons behind AI outputs, and question or request correction when AI makes mistakes. Technology design should therefore allow intervention and reversal. Otherwise, technical power can become difficult to stop once it is set in motion. Technology needs to stay under human control and serve human development.

### **3.3. Understandability and institutional rebuilding: breaking the black box of algorithmic decisions**

Understandability is central to building trust between humans and AI and to keeping technology accountable. It also marks a basic difference between human-centered AI and computationalist AI. Computationalism tends to pursue extreme technical performance, sometimes at the expense of transparency. Human-centered AI starts from a different requirement. Users should be able to understand how the system reaches a decision, and its outputs should be clear enough for ordinary people to follow.

This form of understandability goes beyond technical model explanations such as feature weights and decision paths. More importantly, it requires contextualized explanations that fit how people actually think, including direct answers to questions such as "why did you make this recommendation?" or "why was my application rejected?" To do this, explainable AI (XAI) technologies need to present complex algorithmic decisions as causal logic and plain language that people can follow, rather than leaving them as a black box of algorithmic decisions [7].

Institutional reform should therefore define the rights and responsibilities of each party in human-AI collaboration more clearly and build a closed-loop system through multi-level institutional design. A starting point is special laws that clarify the tool nature and legal status of intelligent systems. In this view, intelligent systems should function as assistive tools rather than independent entities with educational responsibility. The point is practical. If an intelligent system causes direct consequences, responsibility should be shared between users and developers according to fault. Only AI that is transparent and understandable gives users a real chance to predict its behavior and step in when necessary, which makes trust-based human-AI collaboration more likely.

### **3.4. Value alignment: embedding human values into algorithms**

In terms of the institutional conditions needed for human-AI collaboration, existing academic evaluation systems and talent training models are still clearly lagging behind [8]. Value alignment is central to ensuring AI develops for good. It requires AI's goals, operational logic, and results to closely match human moral principles, social norms, and cultural values. This is even more fundamental than explainability and determines the direction of AI development.

Value alignment must work across the whole process. On the data level, we need to check and correct biases in training data, ensure data covers diverse groups and represents diverse values, and avoid data bias causing algorithmic unfairness. On the algorithm level, we need to translate ethical principles like fairness, respect, and public good into enforceable constraints and reward mechanisms for algorithms, moving beyond single-minded efficiency or traffic metrics to balance technical performance and human values. On the application level, we need multi-stakeholder ethics review mechanisms that include developers, ethicists, sociologists, and ordinary people in AI design, evaluation, and oversight, ensuring technology does not stray from humanistic values. Value alignment is not a one-time job but a continuous process that adapts as society and values change.

## 4. Implementation paths: from concept to practice

To turn human-centered ideas into practical applications, human-centered principles need to run through the whole AI lifecycle, from early needs analysis to technical deployment, with humanistic concerns built into each stage.

### 4.1. Adding human considerations in needs analysis

Traditional AI development often begins with technical potential and asks, "what can technology do?" Human-centered AI starts somewhere else. It asks, "what do people need?" Since AI is eventually used in real social settings, its design has to fit the conditions of everyday life if it is to have practical value. AI development also depends on cooperation with social governance bodies rather than technical teams working alone, and this may improve overall effectiveness in preventing ideological risks [9].

At this stage, data collection alone is not enough. More detailed methods from sociology and anthropology are needed to examine users in their actual contexts, so developers can understand lived conditions, less obvious needs, cultural background, and emotional concerns, not only surface-level functional demands.

Needs analysis should also place values such as fairness, dignity, autonomy, and inclusiveness alongside technical requirements, not after them. Through user profiling, scenario simulation, and pain point analysis, developers can identify the human problems that technology is supposed to solve in real settings. That keeps design tied to human needs from the beginning and reduces the risk of developing technology only because it is technically feasible.

### 4.2. Practical applications in key areas

#### 4.2.1. Medical AI

Human-centered medical AI is meant to do more than raise diagnostic accuracy. Its role is to support doctors, strengthen patients' ability to understand and participate, and make healthcare less impersonal. AI can rapidly analyze medical images and test data, present key information visually, and help doctors improve work efficiency. It can also convert technical medical reports and treatment plans into simpler language and charts, which makes it easier for patients to understand their condition and available treatment options and helps protect their right to know.

The difficult cases matter most. In cases involving complex or rare diseases, AI should avoid replacing human judgment when the basis for a decision is uncertain. In those situations, doctors need to make the final judgment so that errors caused by automated decision-making do not create medical risks. This use of AI may improve efficiency while still protecting trust between doctors and patients and preserving the human care that medicine depends on.

#### 4.2.2. Content recommendation and public opinion management

Traditional content recommendation systems often prioritize user time spent on platform and click-through rates. This can produce filter bubbles, repetitive content, and in some cases stronger extreme emotions, which may deepen social division. A human-centered recommendation system sets different goals. It aims to widen users' exposure to information, protect mental health, and support diversity in viewpoints.

To do this, recommendation algorithms can be adjusted to interrupt repeated information patterns. When a user has been exposed to highly similar content for a long period, the system can introduce more diverse and relatively objective material in moderation. The point is to reduce closed information loops, not to force unwanted content into every feed.

Public opinion management raises a similar issue. Here, AI should not function only as a tool for controlling information. It can help managers detect changes in social emotions, identify the concerns behind public reactions, and support conflict resolution. AI can extract the main issues within public opinion, generate objective and neutral analysis reports, and provide a basis for decision-making. Used this way, it supports more constructive social communication instead of simple information suppression.

## 5. Conclusion

In this era of rapid AI development, the growth of data scale, computing power, and model complexity is often seen as the ultimate goal of intelligence. The worship of technology inherent in computationalism can easily lead people into the mistaken belief that technology is everything, causing them to forget that the true purpose of technology is to serve humanity. But in fact, human emotions, intuition, and dignity cannot be fully quantified. Human social nature and spiritual needs cannot be computed by algorithms. AI without human warmth, no matter how powerful, will never truly be a partner to humans. The human-centered AI proposed in this paper does not mean rejecting technological progress or slowing innovation. Instead, it means rebuilding the value foundation of AI and clarifying the core rule for technology development: at every stage—needs analysis, model design, and deployment—we must put human well-being, emotions, autonomy, and dignity first. Through explainable technology, value alignment mechanisms, and human-AI collaboration, we can bring technology back to its true purpose of serving humans. AI technology will continue to advance and innovate. But as long as the fundamental human needs for being understood, respected, and treated well remain unchanged, "human-centered" will not be merely an option for AI development—it will be the sole legitimate foundation for technology's existence. Solving real problems like algorithmic bias and ethical failures, achieving unity between technical reason and human values, is the inevitable direction for AI development and the only path to a future of harmonious human-AI coexistence.

## References

- [1] Yan, J. (2025). Yang Yuanqing promotes human-centered AI for all. *Chinese Entrepreneur*, 12, 31–33.
- [2] Qi, Z., Sun, S., Zheng, G., et al. (2024). Research on human-centered human-machine collaboration quantitative model based on sEMG. *Journal of Mechanical Design and Research*, 40(3), 13–18.
- [3] Cheng, H., Yu, H., & Jiang, X. (2022). A study on the antecedents and consequences of social media information overload in major public health emergencies. *Information Science*, 12, 534–539.
- [4] Zhu, X., & Wang, X. (2025). Algorithmic colonialism and language sovereignty: China's response strategies under global language technology hegemony. *Academic Exploration*, 9, 64–78.
- [5] Zhang, X., & Jin, M. (2021). Transformation and innovation of public opinion governance in the AI era. *Journal of Intelligence*, 40(10), 66–73+165.
- [6] Gu, L. (2026). From mechanical reproduction to algorithmic generation: The technical logic and multidimensional impacts of generative AI. *New Media and Network*, 1–11.
- [7] Li, J., & Zheng, J. (2026). Research on ethical dilemmas and solutions for AI-empowered teacher education. *Academic Exploration*, (1), 172–179.
- [8] Zhou, Q., Chen, W., & Zhou, H. (2026). The two-way empowerment of generative AI and philosophy & social sciences research: Opportunities and challenges. *Zhejiang Social Sciences*, (2), 108–122+159.

- [9] Liu, W. (2024). The generation logic and scientific prevention of ideological risks in the AI era. *Journal of Hohai University (Philosophy and Social Sciences)*, 26(6), 10–19.