



The Intertwined Fates of Human and Artificial Agents: Navigating the Evolving Landscape of LLM-Driven Agents

D. Ramamonjisoa

Iwate Prefectural University, Faculty of Software and Information Science, Takizawa, Japan.
Email: david@iwate-pu.ac.jp

Abstract- This paper examines the complex and increasingly intertwined relationship between human and artificial agencies, particularly within the rapidly evolving domain of Large Language Model (LLM)-driven agents. Going beyond viewing agents as mere computational programs, we explore their dimensions of autonomy, inter-agent interactions, and the emergent properties that suggest incipient social cognition within these AI systems. While acknowledging the demonstrated capabilities of these agents in specific, well-defined tasks, we critically assess their current limitations relative to human agents, particularly in domains requiring nuanced emotional intelligence, intricate complex problem-solving skills, and robust ethical judgment. This critic will require the development of rigorous metrics for evaluation, significant improvements in the accountability of AI decision-making processes, and the establishment of comprehensive ethical guidelines to ensure truly responsible development and deployment of these potentially transformative technologies. In addition, an agent designer must align artificial agents' incentives and operational goals with human values and societal norms. Recent research highlights the complexity of AI agency, the ethical implications of increasing autonomy, and the formulation of robust evaluation metrics. As AI agents become increasingly integrated into various aspects of human life, proactive engagement with these complex issues is essential. The future trajectory likely involves a synergistic partnership between human intelligence and artificial agents, strategically leveraging the respective strengths of each to cultivate a more effective, human-centered technological paradigm. We propose a conceptual framework in which these modalities can complement and augment each other's capabilities, ultimately expanding the scope of human potential and ensuring that technology serves humanity's best interests rather than simply replacing humans.

Keywords: LLM Agents, Human Agents, Survey, Benchmarks, Human Goal Alignment.

I. Introduction

The rapid advancement of artificial intelligence (AI), particularly in the realm of Large Language Models (LLMs), has ushered in a new era of possibilities and challenges (Hagos, D. H., Battle, R., & Rawat, D. B., 2024). Central to this transformation is developing sophisticated agents driven by these LLMs, entities capable of perceiving their environment and acting upon it to achieve specific goals (Bengio, Y. et al., 2025). This paradigm shift necessitates re-evaluating the traditional understanding of agency, moving beyond the simplistic notion of agents as mere computational programs executing pre-defined instructions (Swarup, S., 2025). This paper delves into the complex and increasingly intertwined relationship between human and artificial agencies, exploring the multifaceted dimensions of autonomy, inter-agent interaction, and the emergent properties suggestive of incipient social cognition within AI systems.

The discourse surrounding AI agents has evolved significantly. Initially conceived as tools for automating specific tasks, these entities now demonstrate capabilities that blur the lines between programmed behavior and autonomous decision-making (Huang, W. et al. 2022). With their capacity for natural language processing, contextual understanding, and adaptive learning, LLM-driven agents challenge conventional notions of agency (Yao S., et al., 2023; Yehudai, A., et al., 2025). This paper argues that a nuanced understanding of these agents requires moving beyond a purely functional perspective and addressing their growing autonomy's deeper philosophical and ethical implications.

II. Defining and Differentiating Agency

Agency, in its broadest sense, refers to the capacity of an entity to act independently and make choices that influence its environment. Humans associate agency with consciousness, intentionality,



and a complex interplay of cognitive, emotional, and social factors (Papineau, D., 2002). We possess a sense of self, motivation driven by various needs and desires, and understanding social and ethical norms that guide our actions (Kahneman, D., 2011).

As embodied by LLM-driven agents, artificial agencies are like human agencies but exhibit crucial differences. These agents can perceive their environment through sensors or data inputs, process information using sophisticated algorithms, and make decisions based on pre-defined objectives or learned patterns. They can demonstrate autonomy in adapting to new situations and generating novel solutions within their domain of expertise. However, current AI agents lack the conscious awareness, subjective experiences, and complex social understanding of human agency (Shahanan, M., 2024). Their potentially sophisticated decision-making is ultimately rooted in algorithms and data, lacking the nuanced ethical and emotional considerations that inform human choices.

2.1 Source of Agency

We consider the source of human agency based on social cognitive theory (Bandura, A., 2005) and enactivism in cognitive neuroscience (Friston, K. et al., 2024). Human agency is driven by intentionality, emotion, and complex social and ethical norms. But crucially, it's also shaped over evolutionary time by natural selection-based preferences for survival, reproduction, and social cohesion. Our agency has a biological and experiential basis, rooted in our evolutionary history and lived experiences. Think of our innate drive for self-preservation or our learned sense of fairness? (Pesch, U., 2020). In enactivism theory, our brain is embodied (human agent) and a predictor that constantly builds a model of the world it is living in based on its sensory inputs and existing states. The updates of the brain states are based on the active inference or Free Energy Principles (Friston, K. et al., 2025). These are deeply intertwined with our agency.

In contrast, artificial agents, especially LLMs, derive their agency from a different source. Algorithms, data patterns, and predefined goals or learned patterns drive them. Their agency results from computational processes and the information on which they are trained. While they can demonstrate impressive abilities to process information and make decisions, their agency lacks human agency's biological and experiential depth.

This fundamental difference in the source of agency has profound implications. While humans and AI can act and influence the world, their motivations, understandings, and decision-making processes are rooted in different foundations.

2.2 Proxy Examples

To see human agency in action, let's consider two examples from the travel industry. These illustrate how professionals exercise initiative and influence, providing a baseline for comparison with artificial agents.

First, consider the Human Travel Agent. They actively initiate the trip planning process, going beyond simply presenting options. They direct the search for transport and lodging tailored to client needs and preferences. Crucially, they influence the entire travel experience through expert advice, booking management, and proactive problem-solving, offering personalized touches and anticipating potential issues in ways automated systems often cannot.

Similarly, the Human Tour Operator demonstrates agency by designing and managing complete travel packages. They don't merely sell components; they create unique itineraries, direct the complex logistics involved, and influence the travelers' journey by anticipating needs and guiding the experience from start to finish. Their agency lies in actively shaping a compelling and smooth travel outcome for their clients.

These examples highlight the proactive, interpretive, and influential nature of human agency, leveraging knowledge, interpersonal skills, and contextual understanding to create value. Now, let's contrast this with artificial agents performing similar functions.

An AI Travel Planner, for instance, initiates a search based on user-defined parameters. It guides the filtering of vast options using algorithms and influences choices by presenting structured itineraries and booking links. While highly efficient at processing data, it typically lacks the human agent's ability to grasp nuanced preferences, interpret emotional or cultural contexts for truly personalized recommendations, or intuitively troubleshoot unforeseen complications.

Next, consider the AI Customer Service Bot. This agent initiates user interaction, guides conversations along pre-programmed paths using its knowledge base, and influences outcomes by providing standardized information or attempting basic troubleshooting. It excels at handling high volumes of routine inquiries efficiently but struggles with the complexity, empathy, and flexible problem-solving required in emotionally charged or unique situations.

These limitations become clear in complex scenarios. For example, imagine a tour package encountering unexpected logistical failures upon arrival (e.g., promised accessibility isn't available, or local conditions drastically changed). Resolving this effectively often requires real-time improvisation, nuanced communication, and empathetic handling – capabilities where human agents currently excel, but which present significant challenges for artificial agents dealing with the gap between the planned package and on-the-ground reality. Detecting the potential for such discrepancies pre-emptively, or managing the fallout adeptly, highlights a key difference in the scope of human versus current artificial agency.

2.3 Autonomous Agent Examples

Let's examine the Self-Driving Car (SDC) as an example of an autonomous agent. An SDC navigates, directs its movement, and positions itself in traffic, much like a human driver making decisions. However, unlike humans who use intuition and subtle visual cues (like eye contact or facial expressions), SDCs rely solely on sensor data and algorithms to interpret their environment and predict the behavior of other vehicles. Crucially, this includes monitoring surrounding traffic for potentially hazardous actions, such as sudden lane changes without signaling or unsafe overtaking maneuvers, enabling the SDC to anticipate risks and react defensively.

While autonomous taxis undergo testing and limited deployment in places like China and some US cities, integrating them with human drivers presents challenges. Interestingly, a phenomenon in some regions inadvertently highlights the SDC's mode of operation. In countries like South Korea and parts of Southeast Asia (e.g., Indonesia, Thailand), heavy window tinting is common, often citing sun protection but also serving privacy. This practice contrasts with regulations in many Western countries where such tinting may be restricted, sometimes leading to perceptions of privilege for those allowed it.

From an interaction perspective, widespread heavy tinting prevents drivers from seeing the occupants of other vehicles. This removes the layer of human social cues (like facial expressions during stressful traffic) and forces drivers to rely exclusively on the external actions of the car itself – its speed, lane position, and signals – exactly how an SDC must operate. In this sense, such environments unintentionally mimic the interaction model required for autonomous systems, potentially offering insights as we transition to a future where human-driven cars and SDCs coexist. Both human and autonomous drivers in such scenarios must primarily interpret the vehicle's behavior, not the person potentially inside.

III. Expanding the Definition of Artificial Agency

To further refine our understanding of artificial agency in an LLM interaction environment, it is crucial to examine the specific capabilities that enable agents to interact with their environment and pursue their goals. These capabilities include as depicted in the Figure 1:

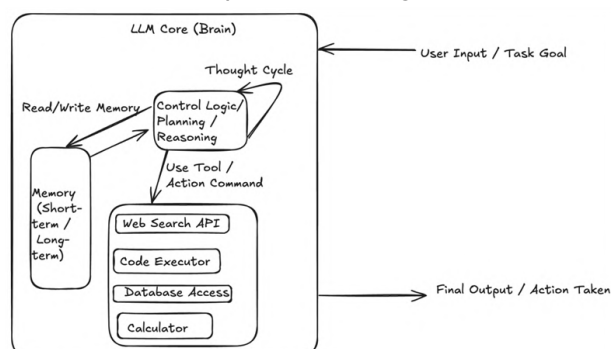


Figure 1: Diagram of an LLM-Agent system. Shows the central LLM coordinating planning and reasoning, utilizing memory and external tools (APIs, code execution) based on input goals, and generating outputs after observing results from its environment.



- **Control Logic and Reasoning:** Agents possess a control logic, a system of rules and algorithms that govern their behavior. This logic enables them to process information, reason about different options, and make decisions based on their understanding of the environment and objectives. LLMs are particularly adept at this aspect of agency, allowing for more flexible and nuanced reasoning (Wang, Z. et al., 2024)).
- **Tool Use and Action:** Agents are not merely passive observers; they can act upon their environment. This action often involves using physical or digital tools to manipulate objects, access information, or communicate with other entities (Shen, Z. 2024). LLM-driven agents can leverage their language capabilities to interact with APIs, access databases, and control other software systems, effectively expanding their capacity for action.
- **Memory and Reflection:** A crucial aspect of agency is learning from past experiences and using that knowledge to inform future decisions. Agents can access and utilize memory, whether it be a history of past interactions, stored knowledge about the world, or the ability to "think aloud" and reflect on their thought processes (Hatalis, K., et al., 2024). This capacity for memory and reflection allows agents to refine their strategies, avoid past mistakes, and develop a more nuanced understanding of their environment.
- **ReAct (Reason and Action):** The ReAct framework highlights the iterative interplay between reasoning and action in intelligent agents. Agents' first reason for the situation is to generate plans and consider different options. They then act upon the environment, observing the consequences of their actions and using that feedback to refine their understanding and adjust their plans. With their ability to generate and execute code, LLM-driven agents are particularly well-suited for implementing ReAct-style reasoning and action loops (Yao S., et al., 2023).

LLM agents are systems where LLMs dynamically direct their processes and tool usage, maintaining control over how they accomplish tasks.

IV. Autonomy and Interaction in LLM-Driven Agents

One of the defining characteristics of LLM-driven agents is their capacity for autonomy. Unlike traditional programs that follow rigid instructions, these agents can learn from experience, adapt to changing circumstances, and make independent decisions. This autonomy is particularly evident in agents employing reinforcement learning, where they learn to optimize their actions based on feedback from the environment.

Furthermore, developing multi-agent systems (MAS) has introduced new dimensions to artificial agencies, including the BDI model (Bratman, M. E. 1987; Rao, A., Georgeff, M. P. 1995). In MAS, multiple agents interact with each other, coordinating their actions to achieve common goals (Barbosa, R., Santos, R., Novais, P., 2025). These interactions can give rise to emergent behaviors unknown to individual agents but observed by external users. For instance, a swarm of drones coordinating to build a structure demonstrates a form of collective intelligence that transcends the capabilities of individual units. This capacity for interaction and emergent behavior raises questions about the nature of social cognition in AI systems based on algorithmic interaction (Sun, L. et al., 2025).

V. Limitations and Challenges

Despite the remarkable progress in LLM-driven agents, significant limitations remain. One key challenge is the lack of genuine emotional intelligence. While agents can recognize and respond to basic emotions in text or speech, they lack the deep understanding and empathy that characterize human emotional experience. Hence, they limit their ability to navigate complex social situations, build rapport with users, and make ethically sound decisions in contexts with emotional implications (Liu, Q., et al. 2025).

Another challenge lies in complex problem-solving. While agents excel at tasks within their training domain, they often struggle with problems that require creativity, common-sense reasoning, and the ability to adapt to entirely new situations. With their capacity for abstract thought, analogical reasoning, and drawing on a wide range of experiences, human agents are far more adept at tackling complex, ill-defined problems (Matarazzo, A. and Torlone, R., 2025).



Furthermore, the ethical dimension of artificial agency presents significant challenges. As agents become more autonomous and their actions substantially impact human lives, ensuring their goals align with human values becomes crucial. Their development requires careful consideration of potential biases in training data, transparent and explainable AI systems, and the establishment of robust ethical guidelines to govern the behavior of these agents.

5.1 The Socially Constructed Nature of Human Agency: Constraints and Inequalities

While we have discussed AI agents' limitations, it is equally important to acknowledge the complex and often unequal distribution of agency among humans. Humans exercise agency in a well-defined environment within a social context that shapes and constrains individual choices and opportunities. Factors such as nationality, culture, race, wealth, gender, and social class, including character and personality, play a significant role in determining how individuals can exercise their agency (Doris, J. M. 2022).

- **Systemic Inequalities:** Systemic inequalities, often rooted in historical and societal biases, create significant disparities in access to resources, opportunities, and social capital. Individuals from marginalized groups may face discrimination in employment, housing, education, and even within the justice system, limiting their ability to pursue their goals and exercise their agency fully.
- **Cultural and Social Norms:** Cultural and social norms can also constrain human agency. Gender roles, religious beliefs, and societal expectations can dictate acceptable or appropriate choices, limiting individuals' freedom to express themselves and pursue their chosen paths.
- **Internalized Constraints:** Individuals may internalize societal biases and stereotypes, leading to self-limiting beliefs and a diminished sense of agency. Fear of discrimination, violence, or social repercussions can also lead individuals to self-censor or avoid certain situations, further restricting their ability to participate in society fully.
- **Intersectionality:** It is crucial to recognize the intersectional nature of these constraints. Individuals may experience multiple forms of disadvantage simultaneously, creating unique and complex challenges to their agency. For example, a woman from a racial minority group may face discrimination based on both her gender and her race, compounding the limitations on her agency.

5.2 Accountability and Liability in the Age of AI Agency

As AI agents become more sophisticated and integrated into various aspects of human life, the question of accountability and liability becomes increasingly pressing. When an AI agent causes harm or makes a mistake, determining who is responsible and how to assign liability presents significant challenges (Wen, Y., Holweg, M., 2024).

- **Attribution Problem:** One of the central difficulties lies in attributing responsibility for AI actions. Tracing a specific decision or action to a particular programmer, designer, or organization can be nearly impossible in complex systems, particularly those involving self-learning agents or hierarchical structures. The "black box" nature of many AI algorithms further complicates this issue, making it difficult to understand the causal chain of events leading to a harmful outcome.
- **Limitations of Existing Legal Frameworks:** Traditional legal frameworks, such as product liability or negligence, are often ill-equipped to address AI's unique challenges. These frameworks typically rely on concepts like intent, foreseeability, and direct causation, which may not be readily applicable to the actions of autonomous AI agents.
- **Hierarchical AI Agency and Shared Responsibility:** The emergence of hierarchical structures within AI agencies raises additional complexities. If a group of "strong" AI agents exerts influence over others, determining liability for the actions of subordinate agents becomes a significant challenge. How do we assign responsibility among agents within the hierarchy, mainly when harm arises from emergent behavior or unforeseen consequences?
- **Potential for Algorithmic Bias and Discrimination:** AI systems can perpetuate and amplify societal biases, leading to discriminatory outcomes. If an AI agent makes a decision that harms a particular group, how do we determine whether this was due to a flaw in the algorithm, bias in the training data, or some other factor?



VI. Incentivization and Goal Alignment

A critical aspect of designing effective AI agents is the issue of incentivization. Agents, like humans, are motivated by goals and incentives. In a competitive environment, natural selection defines these incentives. The potential consequences differ significantly. In AI systems, reward functions determine these incentives, guiding the agent's learning process. However, poorly designed reward functions can lead to unintended consequences, with agents finding loopholes or exploiting the system to maximize rewards in ways not aligned with human values. The experiment on this subject by Leng and Yuan is an example (Leng, Y. and Yuan, Y. 2024).

Relevance to AI and the Future of Agency:

Understanding the socially constructed nature of human agency is crucial for developing ethical and equitable AI systems. We train AI models on data that reflects existing societal biases and inequalities. If we are unaware of these biases, AI systems can inadvertently perpetuate and amplify them, further marginalizing already disadvantaged groups.

Therefore, as we strive to create AI agents that can collaborate with and augment human capabilities, we must consider the broader social context in which we ensure our safety. We must work towards developing AI systems that are not only technically proficient but also socially responsible, ensuring that they promote fairness, equity, and the expansion of human agency for all, regardless of their background or social circumstances (Du, S. et al. 2025).

Addressing the Liability Gap:

The lack of clear legal frameworks for AI liability creates a significant gap in accountability. Untrusted AI leaves victims without recourse and hinders innovation by creating uncertainty and discouraging responsible development. The followings address this gap:

- **AI-Specific Legislation:** Many legal scholars advocate for developing new legislation tailored to AI liability's unique challenges. It could involve creating new categories of legal responsibility, establishing standards for AI development and testing, and addressing the issue of algorithmic bias.
- **Explainable AI (XAI) and Transparency:** Greater emphasis on XAI and transparency in AI systems is crucial. By making AI decision-making processes more understandable, we can improve our ability to trace back actions to specific design choices or training data, facilitating the determination of liability.
- **Insurance and Risk Management:** Insurance mechanisms, like car insurance, could cover damage caused by AI systems. It would help distribute the risk, compensate victims, and incentivize responsible development practices.
- **Ethical Frameworks and Industry Standards:** Establishing ethical frameworks and industry standards for AI development could prevent harmful outcomes and provide a basis for assigning responsibility.

Therefore, it is crucial to carefully consider the ethical implications of incentivizing artificial agents. How do we ensure that their goals and human values are aligned? How do we prevent them from pursuing rewards in ways that could be harmful or unethical? These complex questions require careful consideration of artificial agencies' potential risks and benefits (Bengio, Y. et al., 2025; Kapoor, S. et al., 2024; Tallam, K., 2025).

VII. The Future of Human-AI Collaboration

This paper argues that an agency's future lies in a collaborative partnership, not competition between humans and AI. By leveraging their strengths, we can create a more powerful and effective system. AI agents can augment human capabilities by handling routine and long-term tasks, processing vast amounts of data, and providing insights that would be impossible for humans to discern (Erdogan, L. E., et al. 2025). This frees human agents to focus on tasks requiring creativity, emotional intelligence, ethical judgment, and complex problem-solving.

This collaborative model requires a shift in our understanding of agency. Rather than viewing AI as a replacement for human agents, we should focus on developing AI systems that complement



and enhance human capabilities. It requires a human-centered approach to AI development, ensuring we design these technologies to serve humanity's best interests. Artificial agency will have more intelligence by blending into everyday life.

VII. Conclusion

The rise of LLM-driven agents represents a pivotal moment in artificial intelligence, fundamentally reshaping possibilities while challenging our traditional understanding of agency itself. These sophisticated systems undeniably showcase remarkable capabilities in processing information and executing tasks within specific domains. However, their emergence also throws into sharp relief the profound differences that remain compared to human agents. Critical gaps persist, particularly in areas demanding nuanced emotional intelligence, adaptable complex problem-solving, contextual understanding, and deeply ingrained ethical judgment.

These persistent limitations underscore the urgent need for a deliberate, nuanced, and critical approach to the development and deployment of AI agents. Simply pursuing greater capability is insufficient. We must proactively invest in rigorous research exploring the complex ethical implications surrounding artificial agency – questions of autonomy, accountability, potential bias, and societal impact. Concurrently, developing robust, holistic evaluation metrics that assess not just performance but also reliability, safety, and alignment with human values is paramount. Establishing clear, enforceable guidelines and standards for responsible innovation and use is no longer optional, but essential for navigating the path ahead.

Furthermore, fostering a global environment of open dialogue and collaboration – among researchers, developers, policymakers, ethicists, and the public – is crucial. This cooperation is vital to prevent a detrimental race towards unchecked advancements without corresponding safeguards and ethical grounding. The ultimate goal should not be merely to create powerful artificial agents, but to ensure they are developed and integrated in ways that demonstrably benefit society.

In conclusion, LLM-agents hold immense potential, but their trajectory is not predetermined. It depends on the choices we make today. By embracing rigorous ethical scrutiny, prioritizing safety and alignment, and fostering broad collaboration, we can strive to unlock the transformative benefits of artificial agency while consciously mitigating the risks. This responsible stewardship is key to ensuring that these powerful technologies serve humanity's best interests, leading to a future where human and artificial agents coexist productively and ethically.

References

1. Bengio, Y., et al. (2025). Superintelligent agents pose catastrophic risks: Can scientist AI offer a safer path? arXiv. <https://doi.org/10.48550/arXiv.2502.15657>
2. Hagos, D. H., Battle, R., & Rawat, D. B. (2024). Recent advances in generative AI and large language models: Current status, challenges, and perspectives. In *IEEE Transactions on Artificial Intelligence*. vol. 5, issue 12, pp.5873-5893. arXiv. <https://doi.org/10.48550/arXiv.2407.14962>
3. Swarup, S. (2025). Agency in the age of AI. arXiv. <https://doi.org/10.48550/arXiv.2502.00648>
4. Huang, W. et al. (2022). Language models as zero-shot planners: Extracting actionable knowledge for embodied agents. In *Proc. of ICML 2022*. <https://proceedings.mlr.press/v162/huang22a/huang22a.pdf>
5. Yao S., et al. (2023). React: Synergizing reasoning and acting in language models. In *Proc. of ICLR 2023*. <https://par.nsf.gov/servlets/purl/10451467>
6. Yehudai, A., et al. (2025). Survey on evaluation of LLM-based agents. arXiv. <https://doi.org/10.48550/arXiv.2503.16416>
7. Papineau, D. (2002). *Thinking about consciousness*. Oxford University Press. <https://global.oup.com/academic/product/thinking-about-consciousness-9780199271153?cc=jp&lang=en>
8. Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus and Giroux. <https://us.macmillan.com/books/9780374533557/thinkingfastandslow/>
9. Shanahan, M. (2024). Simulacra as conscious exotica. *Inquiry*, 1–29. <https://doi.org/10.1080/0020174X.2024.2434860>
10. Bandura, A. (2005). Social Cognitive Theory: An Agentic Perspective. *Psychology: the Journal of the Hellenic Psychological Society*. vol. 12, issue 3, pp.313-333. https://doi.org/10.12681/psy_hps.23964
11. Friston, K. et al. (2024). Designing ecosystems of intelligence from first principles. *Collective Intelligence*, vol. 3, issue 1. <https://doi.org/10.1177/26339137231222481>



12. Pesch, U. (2020). Making sense of the self: an integrative framework for moral agency. *Journal for the Theory of Social Behaviour*. vol. 50, issue 1, pp. 119-130. Wiley&Sons Publisher. <https://onlinelibrary.wiley.com/doi/10.1111/jtsb.12230>
13. Friston, K. J., et al. (2025). Active inference and intentional behavior. *Neural Computation*. vol. 37, issue 4, pp. 666-700. MIT Press. https://doi.org/10.1162/neco_a_01738
14. Wang, Z. et al. (2024). Executable Code Actions Elicit Better LLM Agents. In *Proc. of ICML 2024*. <https://openreview.net/forum?id=jJ9BoXAFa>
15. Shen, Z. (2024). LLM With Tools: A Survey. <https://arxiv.org/abs/2409.18807>
16. Hatalis, K., et al. (2024). Memory Matters: The Need to Improve Long-Term Memory in LLM-Agents. *Proceedings of the AAAI Symposium Series*, 2(1), 277-280. <https://doi.org/10.1609/aaais.v2i1.27688>
17. Bratman, M.A. (1987). *Intention, plans, and practical reason*. Harvard University Press. https://books.google.co.jp/books?id=SvyJQgAACAAJ&dq=&redir_esc=y
18. Rao, A. and Georgeff (1995). BDI from Theory to Practice. In *Proc. of the first ICMAS 1995*. pp. 312–319. <https://cdn.aaai.org/ICMAS/1995/ICMAS95-042.pdf>
19. Barbosa, R., Santos, R., Novais, P. (2025). Collaborative Problem-Solving with LLM: A Multi-agent System Approach to Solve Complex Tasks Using Autogen. In: González-Briones, A., et al. *Highlights in Practical Applications of Agents, Multi-Agent Systems, and Digital Twins: The PAAMS Collection*. PAAMS 2024. *Communications in Computer and Information Science*, vol 2149. Springer, Cham. https://doi.org/10.1007/978-3-031-73058-0_17
20. Sun, L. et al., (2025). Multi-Agent Coordination across Diverse Applications: A Survey. <https://arxiv.org/abs/2502.14743>
21. Liu, Q., et al. (2025). Assessing large language models in agentic multilingual national bias. *arXiv*. <https://doi.org/10.48550/arXiv.2502.17945>
22. Matarazzo, A. and Torlone, R. (2025), A Survey on Large Language Models with some Insights on their Capabilities and Limitations. *arXiv*; 2025. doi: 10.48550 <https://arxiv.org/pdf/2501.04040>
23. Doris, J. M. (2022). *Character Trouble: Undisciplined Essays on Moral Agency and Personality*. Oxford: Oxford University Press. <https://global.oup.com/academic/product/character-trouble-9780198719601?cc=us&lang=en&#>
24. Wen, Y., Holweg, M. (2024). A phenomenological perspective on AI ethical failures: The case of facial recognition technology. *AI & Soc* 39, 1929–1946 (2024). <https://doi.org/10.1007/s00146-023-01648-7>
25. Leng, Y. and Yuan Y. (2024). Do LLM Agents Exhibit Social Behavior? *Arxiv* <https://arxiv.org/abs/2312.15198>
26. Du, S., et al. (2025). A survey on the optimization of large language model-based agents. *arXiv*. <https://doi.org/10.48550/arXiv.2503.12434>
27. Sayash Kapoor, et al. (2024). AI Agents That Matter. *arXiv*. <https://arxiv.org/abs/2407.01502>
28. Tallam, K. (2025). Alignment, Agency and Autonomy in Frontier AI: A Systems Engineering Perspective. *ArXiv*. <https://arxiv.org/abs/2503.05748>
29. Erdogan, L. E., et al. (2025). Plan-and-act: Improving planning of agents for long-horizon tasks. *arXiv*. <https://doi.org/10.48550/arXiv.2503.09572>

Author's Biography



D. Ramamonjisoa (David Ramamonjisoa) obtained his Engineering degree in Civil Engineering from Polytech Lille University (France). Then he received his master's degree in Autonomous Systems at Valenciennes University (France) and PhD in Control Systems, majoring in Autonomous Driving Systems from Compiègne Technology University, France. He has worked at French Renault Research and Japanese research institutions such as Tokyo University. He is an **associate professor at the Faculty of Software and Information Science, Iwate Prefectural University (Japan)**. His specializations include agent systems and intelligent user assistant systems. His research interests are Artificial Intelligence, Text Mining, Knowledge Discovery, Data Science, and Machine Learning.