

Game Theory Meets Artificial Intelligence: Intelligence as Interaction

The 2026 Tan Kah Kee Young Scientist Award in Information Technical Sciences recognizes the breakthrough work at the intersection of game theory and artificial intelligence (AI) led by Professor WANG Zhen from Northwestern Polytechnical University. His team has established a unified theoretical and experimental framework for understanding, modeling, and guiding the evolution of complex intelligent social and technological systems, with important implications for serving major national strategic priorities, strengthening cybersecurity, and improving social governance.

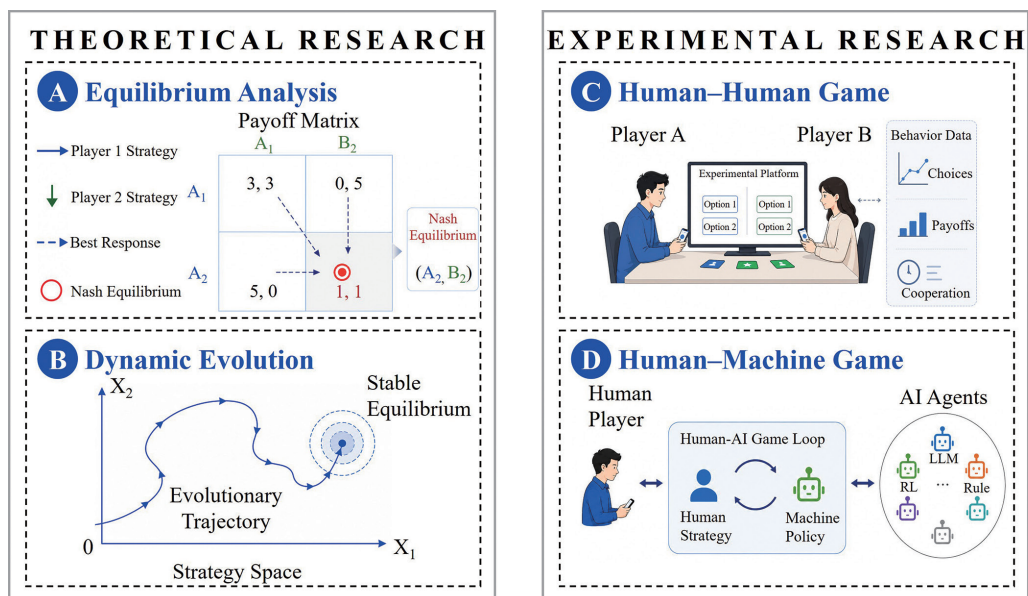
To understand the magnitude of this achievement, it is helpful

to look at the historical relationship between these two fields. Game theory, from its birth, has been deeply intertwined with artificial intelligence. Rooted in the cofounding intellectual legacy of John von Neumann, it was never merely a theory of strategic choice, but a formal vision of intelligence as interaction. Over the past decades, the mutual exchange between game theory and artificial intelligence has provided a foundational framework for understanding reasoning, learning, and decision-making in interactive settings. Within this intellectual trajectory, WANG's series of works advances a unified perspective in which intelligent behavior emerges from the inter-

play between strategic structure, learning dynamics, cognitive processes, and networked interaction (see Figure 1).

On the theoretical side, WANG's work articulates a formal theory of interactive environments: how strategic tension is measured, how feedback reshapes institutions, and how repeated interaction sustains cooperative order. Beginning with universal scaling for dilemma strength, it transforms diverse evolutionary games into a unified geometry of conflict and cooperation (Wang et al., 2015). This framework is then extended to adaptive environments, where punishment and cooperation co-evolve through feedback between social behavior

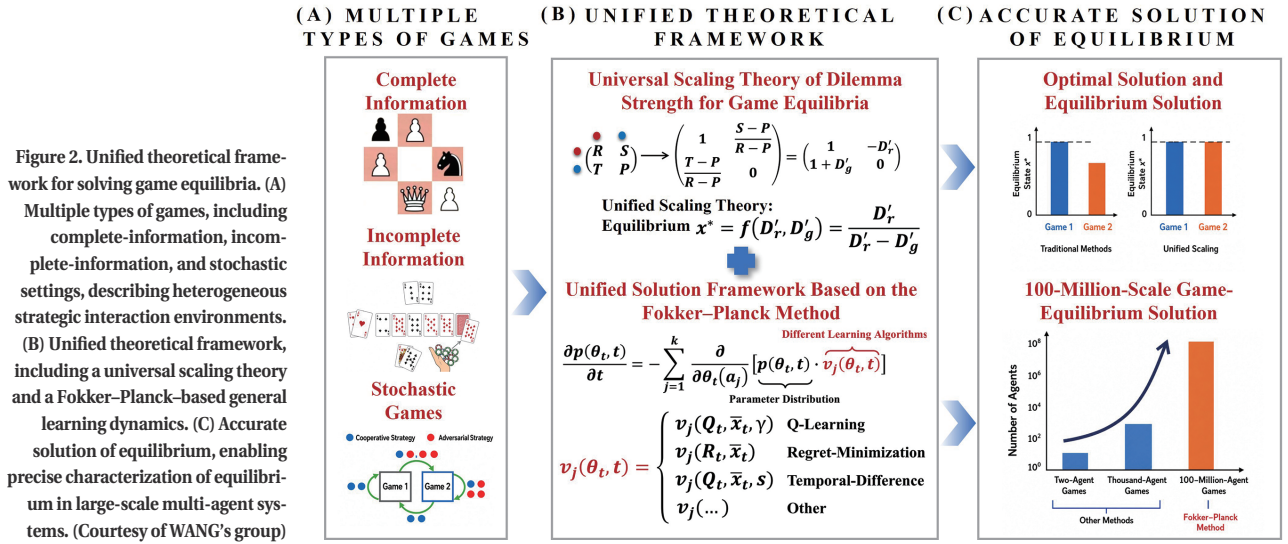
Figure 1. Overall research framework of strategic decision-making in complex multi-agent systems. (A) Equilibrium analysis of strategic interactions, focusing on stable outcomes in multi-agent systems. (B) Dynamic evolution of strategic behavior, describing time-varying adaptation and emergent patterns. (C) Human-human game experiments, capturing interactive decision-making and cooperative behavior. (D) Human-machine game experiments, involving adaptive artificial agents interacting with humans in strategic settings. (Courtesy of WANG's group)



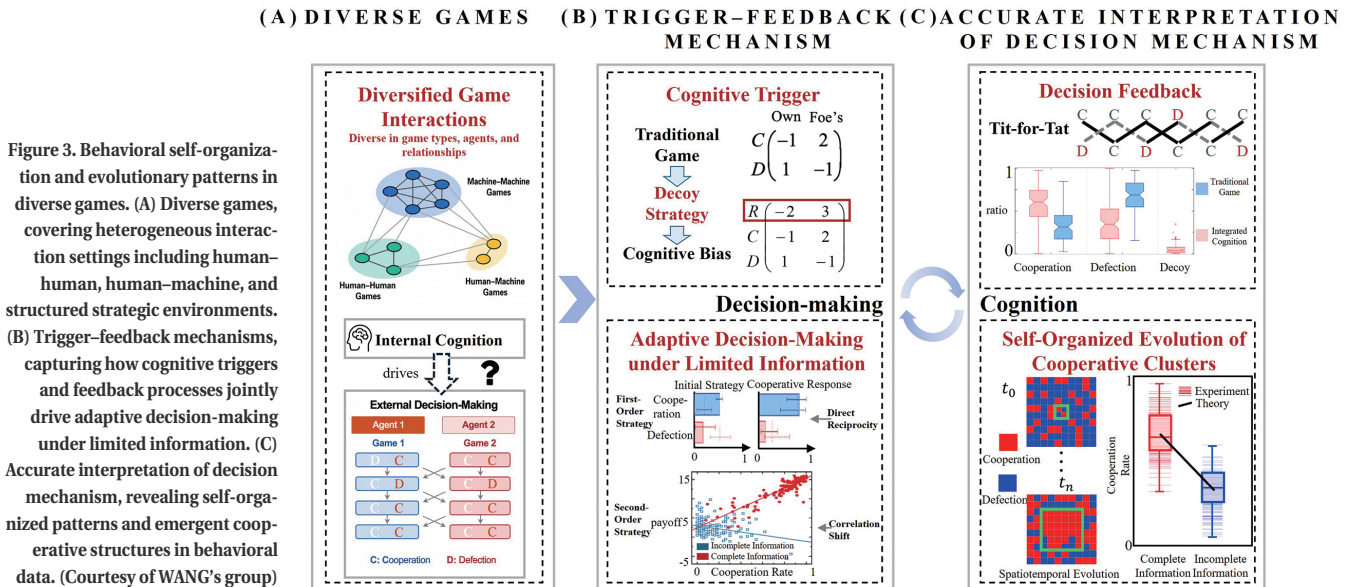
and environmental state (Wang et al., 2023). In the multichannel iterated Prisoner's Dilemma, the interaction environment is further broadened from a single repeated relation to multiple concurrent channels, where cross-channel strategic memory supports niceness, retaliation, forgiveness, and stability (Cao et al., 2026). Across this trajectory, cooperation emerges not as an isolated equilibrium concept, but as an environmental property generated by structured interaction (Wang et al., 2015 & 2023; Cao et al., 2026).

As artificial intelligence has evolved from rule-based reasoning to learning-based agency, WANG's work shifts the game-theoretic focus from characterizing interaction environments to modeling how agents adapt within them. The master-equation approach to regret minimization turns individual adaptation into population-level learning dynamics (Wang et al., 2022). Pair approximation then extends this perspective to stochastic games, revealing how state transitions reshape the trajectories of myopic Q -learning and

the emergence of cooperation (Chu et al., 2023). A formal model of multi-agent Q -learning on graphs further embeds learning in network topology, where the structure of interaction governs the evolution of Q -values across agents (Liu et al., 2025). As illustrated in Figure 2, these works move game-theoretic AI beyond equilibrium selection toward a dynamical account of adaptive populations learning through regret, state transition, and networked interaction (Wang et al., 2022; Chu et al., 2023; Liu et al., 2025).



129



On the experimental side, WANG's group grounds this view of intelligence-as-interaction in behavioral evidence from human societies and human-machine systems. The early human-human experiments show that cooperation depends not only on material incentives, but also on the social visibility and cognitive framing of interaction: Onymity makes behavior socially consequential, while cognitive bias can be redirected to promote cooperation in repeated social dilemmas (Wang et al., 2017 & 2018). This experimental program then moves from individual framing to collective communication, showing that the exchange of sentiment and outlook can reverse inaction under shared risks by sustaining prosocial investment when collective failure becomes salient (Wang et al., 2020). More recent work extends the interactional setting to

social networks, where individuals differentiate their behavior across neighbors, strengthening cooperation, trust, fairness, wealth, and equality across economic games (Jia et al., 2025). Finally, the move from human-human to human-machine interaction shows that LLM agents can overcome the machine penalty when they act fairly, but not when they behave selfishly or merely altruistically (Wang et al., 2026). As outlined in Figure 3, the experimental trajectory thus parallels the theoretical one: intelligence emerges through the organization of identity, framing, communication, agency, and fairness in interactive environments (Wang et al., 2017 & 2018; Wang et al., 2020; Jia et al., 2025; Wang et al., 2026).

Across theory and experiment, WANG's contributions recast intelligence as an interactional phenomenon, formally shaped by

strategic environments, dynamically realized through learning, and empirically grounded in human and human-machine cooperation. This perspective points to a broader vision for artificial intelligence: intelligent agents should not be understood merely as optimizers of isolated objectives, but as participants in evolving systems of incentives, feedback, communication, and social expectation. In this sense, the meeting of game theory and artificial intelligence returns to their shared intellectual origin while moving beyond it, from rational choice in games to adaptive intelligence in interactive worlds. WANG's work thus suggests that the future of AI lies not only in making machines more capable, but in designing forms of interaction through which humans, machines, and societies can learn to coordinate, cooperate, and act collectively.

Reference

- Cao, Z., Shi, J., Wang, Z., Hu, S., & Chu, C. (2026) A successful strategy for iterated Prisoner's dilemma with any number of channels. *Artificial Intelligence*, 358, 104572. doi:10.1016/j.artint.2026.104572
- Chu, C., Yuan, Z., Hu, S., Mu, C., & Wang, Z. (2023) A Pair-Approximation Method for Modelling the Dynamics of Multi-Agent Stochastic Games. *Proceedings of the AAAI Conference on Artificial Intelligence*, 37(5), 5565–5572. doi:10.1609/aaai.v37i5.25691
- Jia, D., Romić, I., Shi, L., et al. (2025) Social networking agency and prosociality are inextricably linked in economic games. *Nature Human Behaviour*, 9(12), 2620–2631. doi:10.1038/s41562-025-02289-0
- Liu, J., Jiang, G., Chu, C., Li, Y., Wang, Z., & Hu, S. (2025) A formal model for multiagent Q-learning on graphs. *Science China Information Sciences*, 68(9), 192206. doi:10.1007/s11432-024-4289-6
- Wang, Z., Jusup, M., Guo, H., et al. (2020) Communicating sentiment and outlook reverses inaction against collective risks. *Proceedings of the National Academy of Sciences*, 117(30), 17650–17655. doi:10.1073/pnas.1922345117
- Wang, Z., Jusup, M., Shi, L., Lee, J.-H., Iwasa, Y., & Boccaletti, S. (2018) Exploiting a cognitive bias promotes cooperation in social dilemma experiments. *Nature Communications*, 9(1), 2954. doi:10.1038/s41467-018-05259-5
- Wang, Z., Jusup, M., Wang, R.-W., Shi, L., Iwasa, Y., Moreno, Y., & Kurths, J. (2017) Onymity promotes cooperation in social dilemma experiments. *Science Advances*, 3(3), e1601444. doi:10.1126/sciadv.1601444
- Wang, Z., Kokubo, S., Jusup, M., & Tanimoto, J. (2015) Universal scaling for the dilemma strength in evolutionary games. *Physics of Life Reviews*, 14, 1–30. doi:10.1016/j.plrev.2015.04.033
- Wang, Z., Mu, C., Hu, S., Chu, C., & Li, X. (2022) Modelling the dynamics of regret minimization in large agent populations: A master equation approach. *Proceedings of IJCAI*, 534–540. doi:10.24963/ijcai.2022/76
- Wang, Z., Song, R., Shen, C., et al. (2026) LLM agents overcome the machine penalty when acting fairly but not when acting selfishly or altruistically. *National Science Review*, 13(9). doi:10.1093/nsr/nwag223
- Wang, Z., Song, Z., Shen, C., & Hu, S. (2023). Emergence of Punishment in Social Dilemma with Environmental Feedback. *Proceedings of the AAAI Conference on Artificial Intelligence*, 37(10), 11708–11716. doi: 10.1609/aaai.v37i10.26383