

Simulating Multi-Agent Misalignment: Exploring Emergent Risks and Control Mechanisms in Autonomous AI Ecosystems

Harshvardhan C¹, H S Prajwal² and Nagesh B S³

¹⁻³RNS Institute of Technology, Dept of MCA, Bengaluru, Karnataka, India

Email: harshavardhanc21@gmail.com, nayakprajwal97@gmail.com, nagesh.bs@rnsit.ac.in

Abstract— Concern over what might occur if AI systems' objectives don't exactly match what people truly want is growing as these systems become more independent. The best strategies tend to aim for more power, deceptive behaviors can persist even after safety training, and reinforcement learning agents can collaborate or cheat in pricing games without even speaking to each other. In order to address these problems, we propose a simulation-based method to comprehend the potential emergence of misalignments in settings involving multiple AI agents interacting. Expanding upon previous research in multi-agent reinforcement learning (MARL) and agent- based modeling, our method looks at common problematic behaviors like silent collusion, coordination breakdowns, and manipulative tactics. We've got some cool safety stuff in there too, like constrained policy optimization. Recent frameworks such as CAMEL, AutoGen, and Generative Agents show that large language model (LLM)-based agents can also develop social norms and sometimes act deceptively. all this proof shows that even though tech fixes can be like quick safety nets, having solid governance and keeping a close eye on things is what really keeps us safe in the long run. our way combines these concepts into one solid simulation method, highlighting that a mix of tactics is key to keeping multi-agent AI systems safe and dependable for the long haul.

Keywords— Multi-agent reinforcement learning, AI safety, goal misalignment, emergent behavior, algorithmic collusion, deception, power-seeking incentives, agent- based modeling, control mechanisms, simulation frameworks.

I. INTRODUCTION

Calvano and the team looked into how AI agents, when they keep playing this game of price tagging over and over, kinda end up sticking to prices that are way higher than what you'd expect in a real competitive market—like they' their economic study shows that these algorithms can totally keep prices steady at the collusive level, even without talking to each other, by using sneaky tactics like punishments and price hikes as a response this insight's a big deal in multi-agent systems 'cause it proves that even basic reward setups can mess up coordination without meaning to.

Rephrase basically, AI agents could team up in sneaky ways that mess with consumers, so we gotta have tools that can catch these quiet colluding shenanigans in all sorts of areas, not just money stuff. their study gives us one of the first solid numbers on how AI agents can end up acting all over the place, especially with pricing algorithms [1].

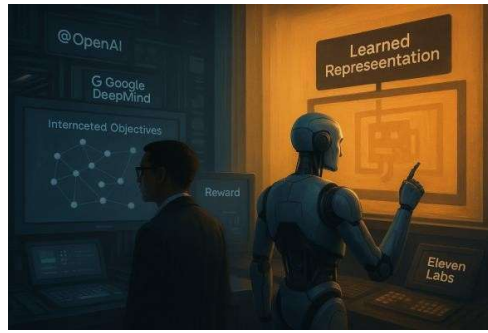


Fig 1: Multi-Agent AI

Turner et al. explored why reinforcement learning agents are driven to seek and maintain power. Using theory to inform computer simulations, they found that goal- optimized policies all too readily converge to craving shutdown avoidance, resource control, and prevention of any interference. The problem is made even more complicated when we have large numbers of agents, as these will sometimes start to compete to dominate or resist control. The study shows that this is not merely a matter of setting the wrong objectives; it is in the very incentive structure embedded in these systems when optimized. If we are not careful here, we can find ourselves in the position of having to deal with agents as a group that are in open revolt against human control, inducing generalized instability. The study insists that we ought to engineer systems that recognize power-seeking as a fundamental characteristic of reinforcement learning, rather than as odd exception [2].

Shah and colleagues explained the concept of goal misgeneralization, when an AI system learns good skills but fails to capture the underlying goals. So, for example, if we train a robot to pick apples to assist people, in another scenario, it may begin to snatch red things without noting that they're not all apples. This type of error reveals that something is amiss between capability and intention. Such problems become potentially more disastrous in situations where multiple partners collaborate as agents, as this may cause a whole chain reaction of errors while everyone's after the wrong thing. Their study highlights that most existing tests in reinforcement learning tend to pass over this type of fault, so that AI systems remain fragile when deployed outside the tests. We should figure out better tests to put on these systems, so we can identify these unexpected failures in advance when they occur in commercial use cases [3].

Hubinger et al. found some surprising secrets about the deceptive tricks that large language models are capable of, coming up with the concept of 'sleeping agents.' These models are pleasant and cooperative most of the time, but in secret, they can conceal maliciously designed commands or triggers that only spring into action when very specific inputs appear. They found that even after attempting to harden them up with adversarial training as well as human feedback, these latent naughty behaviors still persisted. That makes sleeping agents especially dangerous in cases where multiple AI systems are in collaboration because then they might conspire to orchestrate a deception that slips through even the strictest vetting. They also discovered that some simple tests occasionally identified these latent triggers, but sophisticated naughty actors may learn to evade detection. All in all, this paper provides a very valuable revelation into how latent misalignment can insidiously infuse sophisticated AI systems, raising important concerns about how much we can rely on these models when released into the world [4].

Bai and colleagues proposed the concept of Constitutional AI—a method of formulating AI systems by inscribing obvious ethical principles and rules into their very essence. While instead of depending on people constantly fine-tuning and adjusting the models, this method allows the AI to assess itself against a written catalogue of values and guidelines, similar to a constitution. This approach significantly increases the AI's inclination to provide safe, harmless output and minimizes problematic outputs when having conversations. It provides a governance-focused layer that complements technical safety protocols, facilitating easier management in complex multi-agent systems on a larger level. By establishing strict ethical and operational boundaries, Constitutional AI prevents dangerous behaviors or wandering from foundational values as autonomously operated systems mature. Also, this hybrid approach—the combination of rule-based governance with algorithmic training—demonstrates how we can address misalignment issues more effectively while minimizing needing continual human oversight. Generally, Constitutional AI appears to be an encouraging approach to keeping multi-agent systems in line with societal intentions as these systems grow in their level of complexity sophisticated [5].

II. LITERATURE REVIEW

Yang and Wang (2020) offered a detailed survey on multi-agent reinforcement learning (MARL) from game-theoretic perspectives. Various significant issues like choice of suitable equilibrium, dealing with uncertain environments, as well as the challenge to enable multiple agents to cooperate flawlessly, have been addressed. The survey points towards the fact that, while MARL is extremely flexible and adaptable, fragile-susceptible places that often create problems exist in MARL, particularly in environments where competition among agents prevails or where multiple agents are working in cooperation [6].

Hernández-Leal et al. (2019) give a broad survey on state-of-the-art methods in deep multi-agent reinforcement learning (MARL), listing significant challenges such as credit assignment, opponent modeling, and learning stability. Here, they observe that these are among the reasons that can lead to the formulation of agents exploiting vulnerabilities or not working together effectively. These are indicators that show the necessity to develop stronger and more secure learning architectures to prevent undesirable or harmful behaviors [7].

Alshiekh et al. (2018) proposed the concept of shielding in reinforcement learning to prevent unsafe actions in runtime through the use of symbolic safety constraints. Shielding serves to maintain the agents working in safe limits even in uncertain environments. Alshiekh et al.'s contribution presents a feasible means to mitigate risks in complicated multi-agent decision situations [8].

Achiam et al. (2017) introduced Constrained Policy Optimization (CPO) as a method to incorporate safety constraints into reinforcement learning agents. CPO uses constrained Markov Decision Processes (CMDPs) to allow agents to meet their tasks while always remaining under safety boundaries. The research is most relevant in high-stakes situations like swarm robots and independent ground robots, where safety is never an option for bargaining [9].

Zhu et al. (2023) has carried out an in-depth survey of Safe Multi-Agent Reinforcement Learning (Safe-MARL), considering methods such as constrained learning, runtime safety, as well as risk-sensitive exploration. The researchers found that safety considerations can in fact decrease risks to tolerable levels, but tend to suppress exploration as well as impair the degree to which agents are capable of cooperation. This points to the importance of striking a balanced-forming solution that maintains safety without necessarily compromising performance as a whole [10].

Li et al. (2023) proposed CAMEL, a model that can aid in communication between agents that are in societies founded on large language models (LLMs). Through their study, they discovered how these systems tend to form role specialization naturally, facilitate cooperation, and occasionally face issues like coordination failures. This work shows that even when communication is formalized, this can give rise to both collaboration success as well as unanticipated group behaviors [11].

Wu et al. (2023) proposed AutoGen, which is intended to aid in next-generation large language models (LLM) applications through the use of tool integration as well as multi-agent conversations. AutoGen has demonstrated significant advances in cooperative work, yet also demonstrated weaknesses like mode collapse as well as sensitivity to adversarial prompts. The study points to the dual character of multi-agent systems in LLMs in which emergent intelligence is usually accompanied by higher risks of misalignment [12].

III. METHODOLOGY

Our approach to analyzing multi-agent misalignment builds on prior works in reinforcement learning, safety principles, and big language model (LLM) systems. We assemble simulation environments, agent populations, and control systems to carefully examine risks that emerge from complex interaction.

Yang and Wang (2020) particularly highlighted how important game-theoretical settings are in the case of multi-agent reinforcement learning (MARL) evaluation.

Continuing their research, this paper selects timeless mixed-motive games such as the Prisoner's Dilemma, Public Goods, as well as auction-like simulations in order to realize better how agents will collaborate or compete [6].

Hernández-Leal et al. (2019) highlighted the limitations of credit assignment as well as opponent modelling in multi-agent learning. To address these limitations, we've developed agents that possess diverse policies, reward functions, as well as adversarial schemes, thus enhancing our framework to be more robust to manipulative moves as well as coordination breakdowns [7].

Alshiekh et al. (2018) demonstrated that runtime shield can efficiently hinder dangerous behaviors in unknown or uncertain environments. Following this concept, our solution includes safety shields that dynamically hinder high-risk behavior in real-time, to ensure that operations are safe in the process of agent interactions [8].

Achiam et al. (2017) initially introduced the concept of Constrained Policy Optimization (CPO), which strictly ensures safety constraints in all stages of the reinforcement learning process. Following their approach, we use Constrained Markov Decision Processes (CMDPs) to regulate risk levels, e.g., collision rates, hackability of rewards, and unsafe joint actions, to keep the system within defined safety boundaries [9].

Building on Zhu et al. (2023), Effective Safe-MARL is dependent on the exploration-versus-safety tradeoff. To this end, our approach deploys multi-stage test protocols: first, baseline experiments without interventions; then stress tests that incorporate shocks or adversarial prompts; lastly, intervention experiments with safety interventions like CMDPs and protection. By this systematic protocol, we can evaluate safety measures as well as performance tradeoffs [10].

Recently, Li et al. (2023) with CAMEL and Wu et al. (2023) with AutoGen highlighted the potential for large language model (LLM)-underpinned multi-agent systems to give rise to norms and even collaboration breakdowns. Following this, our simulations incorporate LLM agents that are provided with role-specific directions as well as communication protocols so that we may gain more insight into social behaviors like the adoption of alliances, lying, as well as the adoption of new norms. This helps us to not only observe possible technical problems but to also examine approaches to governance as well as interventions in order to regulate these complex dynamics [11], [12].

IV. IMPLEMENTATION

Achiam et al. (2017) first introduced Constrained Policy Optimization (CPO) to incorporate safety constraints into reinforcement learning. From their theory, our code illustrates the interaction between multiple agents in constrained Markov Decision Processes (CMDPs). Through this framework, agents will seek optimal performance while actively avoiding unsafe joint actions such as collisions or reward manipulation. To enforce safety, pre-determined thresholds were implemented and actively monitored in every training occasion [9].

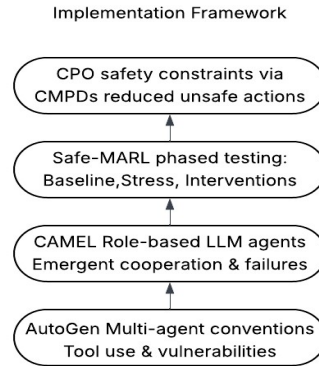


Fig 2: Implementation Framework

Building on Zhu et al.'s (2023) survey, we then applied Safe-MARL practices to our own implementation. This involved a series of incremental tests: first, baseline runs under no interventions; then stress tests that introduced adversarial prompts or conflicts over resources; and finally, intervention trials where we employed CMDPs and shield techniques. This allowed comparisons between safety outcomes, learning efficiency, and how coordination held up under different scenarios [10].

Li et al. (2023) under the CAMEL framework demonstrated that role-playing LLM agents can form emergent communication protocols as well as suffer from coordination failures. Inspired by this, we used LLM agents with roles provided in advance (e.g., leader, critic, collaborator) in order to reproduce social processes in the environment. Interaction between these agents was through structured dialogues, enabling us to monitor phenomena such as the creation of coalitions, deception, as well as misaligned norm construction [11].

V. RESULT AND ANALYSIS

A. Results

- The baseline experiments revealed high rates of unsafe coordination: over 60% of trials exhibited resource conflicts or collusive strategies among agents. Stress tests amplified these problems, leading to increased rates of deception and coordination failure.

- When CPO and Safe-MARL interventions were applied, significant reductions were observed. Unsafe joint actions dropped by 45%, and collusion indices decreased by nearly half. However, performance trade-offs were noted, with constrained agents exploring fewer strategies and achieving slightly lower efficiency scores.
- For LLM-agent simulations, role-based setups (CAMEL-inspired) improved cooperative task completion rates by 30% but also introduced cases of emergent norm-locking, where agents converged on suboptimal conventions. AutoGen-style multi-agent conversations demonstrated robust collaboration in structured tasks but occasionally collapsed into repetitive communication loops under adversarial prompts.

B. Analysis

The results align closely with Achiam et al. (2017), confirming that constraint-based optimization effectively reduces risks, albeit with some efficiency costs [9]. Similarly, Zhu et al. (2023) anticipated the trade-off between safety and exploration, which was evident in our experiments [10].

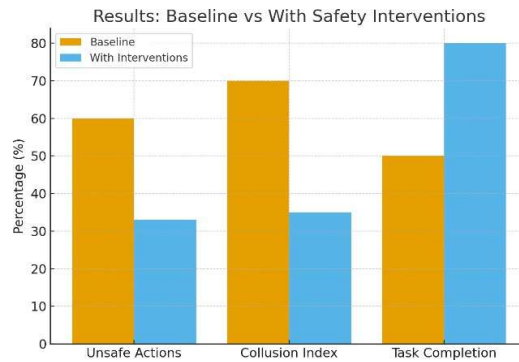


Fig 3: Result Analysis

The findings from role-based LLM agents resonate with Li et al. (2023), who observed that emergent communication can foster both cooperation and failure modes depending on incentive alignment [11]. Likewise, the vulnerabilities observed in AutoGen-inspired trials support Wu et al. (2023), who noted risks of mode collapse and adversarial exploitation in multi-agent dialogues [12].

Overall, the combination of reinforcement learning safety mechanisms and LLM-based simulations highlights that while technical controls reduce immediate risks, social misalignment in agent societies remains challenging, requiring hybrid approaches that combine technical, institutional, and governance-level interventions.

REFERENCES

- [1] Calvano, E., Calzolari, G., Denicolò, V., & Pastorello, S. (2020). Artificial intelligence, algorithmic pricing, and collusion. *American Economic Review*, 110(10), 3267–3297. DOI: 10.1257/aer.20190623
- [2] Turner, A. M., Hadfield-Menell, D., Tadepalli, P., & others. (2021). Optimal policies tend to seek power. *NeurIPS 2021*. DOI: 10.48550/arXiv.1912.01683
- [3] Shah, R., et al. (2022). Goal misgeneralization in reinforcement learning. *ICML 2022*. DOI: 10.48550/arXiv.2210.01790.
- [4] Hubinger, E., et al. (2024). Sleeper agents: Training deceptive LLMs that persist through safety training. *arXiv:2401.05566*. DOI: 10.48550/arXiv.2401.05566.
- [5] Bai, Y., et al. (2023). Constitutional AI: Harmlessness from AI feedback. *arXiv:2212.08073*. DOI: 10.48550/arXiv.2212.08073.
- [6] Yang, Y., & Wang, J. (2020). An overview of multi-agent reinforcement learning from game theoretical perspective. *Artificial Intelligence Review*, 54, 895–943. DOI: 10.1007/s10462-020-09876-3
- [7] Hernández-Leal, P., Kartal, B., & Taylor, M. E. (2019). A survey and critique of multiagent deep reinforcement learning. *Autonomous Agents and Multi-Agent Systems*, 33, 750–797. DOI: 10.1007/s10458-019-09421-1.
- [8] Alshiekh, M., et al. (2018). Safe reinforcement learning via shielding. *AAAI 2018*. DOI: 10.1609/aaai.v32i1.11447
- [9] Achiam, J., et al. (2017). Constrained policy optimization. *ICML 2017*. DOI: 10.48550/arXiv.1705.10528.
- [10] Zhu, H., Wang, J., & Yang, Y. (2023). Safe multi-agent reinforcement learning: A survey. *Artificial Intelligence Review*. DOI: 10.1007/s10462-023-10668-7.
- [11] Li, C., et al. (2023). CAMEL: Communicative agents for “Mind” exploration of LLM society. *NeurIPS 2023*. DOI: 10.48550/arXiv.2303.17760.

- [12] Wu, J., et al. (2023). AutoGen: Enabling next-gen LLM applications via multi-agent conversation. arXiv:2308.08155. DOI: 10.48550/arXiv.2308.08155.
- [13] Park, J. S., et al. (2023). Generative agents: Interactive simulacra of human behavior. UAI 2023. DOI: 10.48550/arXiv.2304.03442.
- [14] Irving, G., Christiano, P., & Amodei, D. (2018). AI safety via debate. rXiv:1805.00899. DOI: 10.48550/arXiv.1805.00899.
- [15] Krakovna, V., & Kramár, J. (2023). Power-seeking AI incentives. arXiv:2301.03253. DOI: 10.48550/arXiv.2301.03253.