

Curricula and Multi-Agent Drafting for Reinforcement-Learned Software Reasoning

Abstract

Reinforcement Learning For Software Reasoning has become a test of how well researchers can connect performance with evidence quality, resource limits, and transfer across settings. This comparative analysis considers organizing intermediate drafts and code lineage into learnable reasoning trajectories. The comparison integrates 2 focal papers with 13 independently retrieved publications verified through persistent DOI or publisher records. The analysis is organized around draft coordination, curriculum design, lineage graphs, reward hacking, and generalization. Instead of assuming that metrics from unlike protocols as commensurate, the review compares research framing, method assumptions, and test envelope. Across the literature, a robust inference is that advances in reinforcement learning for software reasoning become credible when representation, objective, and evaluation protocol are evaluated together and when uncertainty about distribution shift is reported explicitly. The comparative structure connects method selection to decision consequence and recurring validity threats, and proposes a research agenda centered on comparable baselines, sensitivity analysis, and retained provenance.

Keywords

Reinforcement Learning For Software Reasoning; Draft Coordination; Curriculum Design; Lineage Graphs; Reward Hacking; Generalization

1. Introduction

Research on reinforcement learning for software reasoning links scientific ambition and practical constraint. The scientific task demands not only stronger nominal performance but also explanations of the mechanisms that support observed behavior. The boundary is clearest when considering organizing intermediate drafts and code lineage into learnable reasoning trajectories. Evidence that appears persuasive within a specific substrate, benchmark, population, spatial unit, or workload may be fragile outside its original boundary. For this reason, analysis must preserve the unit of analysis and the decision context. Equally important is distinguishing a mechanism from a correlation, a benchmark advantage from a deployment advantage, and an interpretable diagnostic from a causal explanation.

This review uses five analytical lenses: draft coordination, curriculum design, lineage graphs, reward hacking, generalization. These dimensions matter because they jointly cover what changes, how change is observed, and which conditions must persist. The comparison begins from representation, objective, and evaluation protocol. Under that principle, technical creativity remains only one part of credibility; the comparison also audits the baseline, uncertainty treatment, and resource or hazard boundary. This contribution is a literature synthesis and is not presented as a new experiment and supplies no synthetic performance claim.

2. Methodological Foundations

The analytical chain starts from the input evidence, the transformation applied to it, the output used for evaluation, and the decision that follows. In reinforcement learning for software reasoning, the chain is commonly fragmented even though errors propagate across them. A powerful model can intensify errors embedded in the initial evidence; an apparently accurate output can still be poorly calibrated for the final decision. For that reason, ablation evidence should accompany aggregate performance. A useful evaluation makes explicit the constants, perturbations, and uncontrolled factors, then selects metrics that correspond to the intended use rather than to convenience alone.

A second premise concerns transfer boundaries. Draft coordination defines the local design problem, while generalization determines whether the design can survive outside that boundary. Bridging the two are curriculum design and lineage graphs connect those ends by exposing trade-offs that a single score conceals. At this point, null findings and sensitivity evidence maps the conditions under which the explanation remains viable. For applied work, documentation of deployment constraints is therefore part of the scientific result, not an administrative afterthought.

3. Architecture and Learning Objectives

The target study YAA-02, DRAFT-RL: Multi-Agent Chain-of-Draft Reasoning for Reinforcement Learning-Enhanced LLMs [1], addresses a concrete part of reinforcement learning for software reasoning through multi-agent chain-of-draft reasoning optimized with reinforcement learning. It identifies a representation, mechanism, or evaluation boundary needed to compare claims. Against the focus on organizing intermediate drafts and code lineage into learnable reasoning trajectories, the paper functions as a bounded point of comparison: transfer may depend on its sensing regime, preparation, workload, population, or benchmark. The synthesis maintains the declared scope and triangulates the underlying assumptions.

The target study YAA-03, Evo-CuRL: Curriculum-Aware Reinforcement Learning over Code Lineage Graphs for Software Engineering Reasoning [2], addresses a concrete part of reinforcement learning for software reasoning through curriculum-aware reinforcement learning over code-lineage graphs. It identifies a representation, mechanism, or evaluation boundary needed to compare claims. Against the focus on organizing intermediate drafts and code lineage into learnable reasoning trajectories, the paper functions as a bounded point of comparison: transfer may depend on its sensing regime, preparation, workload, population, or benchmark. The synthesis maintains the declared scope and triangulates the underlying assumptions.

Across the focal studies, draft coordination should be interpreted jointly with curriculum design. A design can improve one yet redistribute uncertainty rather than remove it. A structured comparison takes the form of a matrix whose rows are operating conditions and whose columns are evidence quality, computation or process cost, robustness, and consequence of failure. The matrix is designed to prevent an attractive mechanism from being detached from the circumstances that support it. It also clarifies which ablations are informative: ablation design should target causal attribution instead of numerical proliferation.

Lineage graphs links technical design with operational choice. A defensible assessment must contrast nominal operation with boundary tests and report more than means, and test plausible alternative explanations. Where distribution shift is central, calibration or sensitivity is as important as ranking. Where costs or hazards are asymmetric, utility-weighted assessment is more revealing than undifferentiated accuracy. These choices preserve methodological differences; they make the reasons for their differences auditable.

4. Comparative Evaluation

A complementary strand is visible in Reinforcement Learning Model Based on Regret for Multi-Agent Conflict Games [3]; Stabilizing independent multi-agent reinforcement learning via curriculum-based iterative self-play [4]; Multi-hop temporal knowledge graph reasoning with multi-agent reinforcement learning [5]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of reinforcement learning for software reasoning. Together they illuminate draft coordination: some emphasize representations or mechanisms, whereas others foreground validation and application conditions. The useful comparison is not a leaderboard across incompatible settings, but a structured reading of what each study controls, leaves variable, and treats as a meaningful outcome. Divergence can then reveal a change in scale, data process, endpoint, or operating constraint.

For triangulation, the literature includes Abmarl: Connecting Agent-Based Simulations with Multi-Agent Reinforcement Learning [6]; Curriculum-Learning-Guided Multi-Agent Deep Reinforcement Learning for N-1 Static Security Prevention an... [7]; Curriculum-assisted multi-agent reinforcement learning for scalable V2X resource allocation [8]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of reinforcement learning for software reasoning. Together they illuminate curriculum design: some emphasize representations or mechanisms, whereas others foreground validation and application conditions. The useful comparison is not a leaderboard across incompatible settings, but a structured reading of what each study controls, leaves variable, and treats as a meaningful outcome. Divergence can then reveal a change in scale, data process, endpoint, or operating constraint.

The design space is broadened by Legal loophole detection model based on multi-agent reinforcement learning [9]; A multi-agent reinforcement learning framework for educational resource management optimisation [10]; Rule-augmented curriculum transfer reinforcement learning for multi-agent UAV adversarial tasks [11]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of reinforcement learning for software reasoning. Together they illuminate lineage graphs: some emphasize representations or mechanisms, whereas others foreground validation and application conditions. The useful comparison is not a leaderboard across incompatible settings, but a structured reading of what each study controls, leaves variable, and treats as a meaningful outcome. Divergence can then reveal a change in scale, data process, endpoint, or operating constraint.

A first comparison set connects Legal loophole detection model based on multi-agent reinforcement learning [12]; Collaborative optimisation of multi-agent reinforcement learning in enterprise digital supply chain [13]; A multi-agent reinforcement learning framework for educational resource management optimisation [14]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of reinforcement learning for software reasoning. Together they illuminate reward hacking: some emphasize representations or mechanisms, whereas others foreground validation and application conditions. The useful comparison is not a leaderboard across incompatible settings, but a structured reading of what each study controls, leaves variable, and treats as a meaningful outcome. Divergence can then reveal a change in scale, data process, endpoint, or operating constraint.

The neighboring evidence base brings together RAMARL: Robustness Analysis with Multi-Agent Reinforcement Learning - Robust Reasoning in Autonomous Cyb... [15]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of reinforcement learning for software reasoning. Together they illuminate generalization:

some emphasize representations or mechanisms, whereas others foreground validation and application conditions. The useful comparison is not a leaderboard across incompatible settings, but a structured reading of what each study controls, leaves variable, and treats as a meaningful outcome. Divergence can then reveal a change in scale, data process, endpoint, or operating constraint.

5. Deployment and Governance

Implementation can follow a staged workflow. First, define the decision and its failure cost. Second, specify the evidence and operating envelope. Third, choose a method whose inductive assumptions match that evidence. Fourth, test reward hacking and generalization under controlled perturbations. Finally, retain provenance so that an unexpected outcome can be traced to data, configuration, material batch, environmental condition, or user interaction. The workflow connects organizing intermediate drafts and code lineage into learnable reasoning trajectories connected to observable requirements.

Across applications, responsible progress in reinforcement learning for software reasoning depends on interfaces as much as components. Interfaces determine how uncertainty moves between measurement and inference, between algorithm and operator, or between microscopic design and system behavior. A shared protocol should state acceptance conditions and what would overturn a claim. When those agreements are explicit, performance tuning can proceed while preserving visible validity and safety requirements.

6. Limitations and Future Directions

Several limitations qualify the synthesis, including divergent terms, methods, and documentation depth. Metadata confirmation secures the citation trail but does not reproduce measurements or results. Venue and version information may evolve. The supplied focal strings remain preserved while questionable normalization is shown only as a candidate.

Research can advance by seeking to preregister comparison protocols where feasible, disclose reusable configurations and examine transfer beyond the nominal regime in deployment constraints. Research should also document why failures occur in addition to how often. For reinforcement learning for software reasoning, a valuable shared benchmark would combine standard-condition utility, efficiency, confidence quality, and fault recovery. Such a benchmark would reward designs that remain useful when evidence is incomplete or conditions change.

7. Conclusion

The evidence concerning reinforcement learning for software reasoning becomes clearer when treated as a linked evidence chain rather than a collection of isolated scores. The focal studies show how organizing intermediate drafts and code lineage into learnable reasoning trajectories; the additional literature reveals the assumptions required to interpret those contributions. Across draft coordination, curriculum design, lineage graphs, reward hacking, and generalization, a durable conclusion is the need for alignment: the representation or mechanism must match the problem, assessment must align with action, just as deployment claims align with validation scope. Alignment makes transfer more cautious, replication more feasible, and failure more instructive.

References

1. Li, Y., Liu, M., Wang, H., Zhang, Y., Ma, Y., & Tan, W. (2026). DRAFT-RL: Multi-Agent Chain-of-Draft Reasoning for Reinforcement Learning-Enhanced LLMs. *Proceedings of the AAAI Conference on Artificial Intelligence*, 40(35), 29530-29537.

2. Tan, W., Li, Y., & Liang, W. (2026, June). Evo-CuRL: Curriculum-Aware Reinforcement Learning over Code Lineage Graphs for Software Engineering Reasoning. In *Proceedings of the 2026 International Conference on Multimedia Retrieval* (pp. 1327-1335).
3. XIAO, Z., & ZHANG, S. Y. (2009). Reinforcement Learning Model Based on Regret for Multi-Agent Conflict Games. *Journal of Software*, 19(11), 2957-2967. <https://doi.org/10.3724/sp.j.1001.2008.02957>
4. Akgün, O. (2026). Stabilizing independent multi-agent reinforcement learning via curriculum-based iterative self-play. *Neurocomputing*, 704, 134819. <https://doi.org/10.1016/j.neucom.2026.134819>
5. Bai, L., Chen, M., & Xiao, Q. (2024). Multi-hop temporal knowledge graph reasoning with multi-agent reinforcement learning. *Applied Soft Computing*, 160, 111727. <https://doi.org/10.1016/j.asoc.2024.111727>
6. Rusu, E., & Glatt, R. (2021). Abmarl: Connecting Agent-Based Simulations with Multi-Agent Reinforcement Learning. *Journal of Open Source Software*, 6(64), 3424. <https://doi.org/10.21105/joss.03424>
7. Zhang, X., Li, Z., Quan, X., Cheng, K., & Yu, Y. (2026). Curriculum-Learning-Guided Multi-Agent Deep Reinforcement Learning for N-1 Static Security Prevention and Control. *Energy Engineering*, 123(9), 1-10. <https://doi.org/10.32604/ee.2025.073912>
8. Morshed, M., & Zaman Chowdhury, M. (2026). Curriculum-assisted multi-agent reinforcement learning for scalable V2X resource allocation. *Physical Communication*, 76, 103062. <https://doi.org/10.1016/j.phycom.2026.103062>
9. Ma, D. (2026). Legal loophole detection model based on multi-agent reinforcement learning. *International Journal of Reasoning-based Intelligent Systems*, 18(9). <https://doi.org/10.1504/ijris.2026.10076379>
10. Feng, Y. (2026). A multi-agent reinforcement learning framework for educational resource management optimisation. *International Journal of Reasoning-based Intelligent Systems*, 18(7). <https://doi.org/10.1504/ijris.2026.10075912>
11. Liu, Y., Liu, S., Wu, Y., Jing, B., Yan, S., & Xu, Y. (2025). Rule-augmented curriculum transfer reinforcement learning for multi-agent UAV adversarial tasks. *Engineering Research Express*, 7(4), 0452h2. <https://doi.org/10.1088/2631-8695/ae2b50>
12. Ma, D. (2026). Legal loophole detection model based on multi-agent reinforcement learning. *International Journal of Reasoning-based Intelligent Systems*, 18(9), 51-64. <https://doi.org/10.1504/ijris.2026.152191>
13. Zhou, H. (2026). Collaborative optimisation of multi-agent reinforcement learning in enterprise digital supply chain. *International Journal of Reasoning-based Intelligent Systems*, 18(11). <https://doi.org/10.1504/ijris.2026.10077115>
14. Feng, Y. (2026). A multi-agent reinforcement learning framework for educational resource management optimisation. *International Journal of Reasoning-based Intelligent Systems*, 18(7), 32-43. <https://doi.org/10.1504/ijris.2026.151422>
15. Saad, A., & Håkansson, A. (2022). RAMARL: Robustness Analysis with Multi-Agent Reinforcement Learning - Robust Reasoning in Autonomous Cyber-Physical Systems. *Procedia Computer Science*, 207, 3662-3671. <https://doi.org/10.1016/j.procs.2022.09.426>