

Safe Exploration in Embodied Reinforcement Learning

Aaron Phillips - Corresponding author

Abstract

This review examines safe exploration in embodied reinforcement learning. The organizing question is how agents can learn informative behaviour while respecting physical, social, and operational constraints. Ten related scholarly sources are synthesized through a decision-centered framework spanning problem definition, mechanism, measurement, evaluation, implementation, and governance. The review does not invent experiments, pooled estimates, or unreported quantitative results. It instead evaluates the strength and transferability of the available evidence, with particular attention to treating simulated constraint satisfaction as proof of physical-world safety. The resulting framework links technical or empirical performance to explicit use conditions and identifies tests that should precede wider adoption in robotics and autonomous embodied systems.

Keywords: safe reinforcement learning, embodied AI, exploration, robotics, constraints

Synthesis lens	Principal question
Mechanism	the chain connecting evidence inputs, safe exploration in embodied reinforcement learning, intermediate outputs, and downstream decisions
Evaluation	construct validity, comparative performance, uncertainty, transfer across settings, and decision utility
Primary risk	treating simulated constraint satisfaction as proof of physical-world safety

1. Introduction

Research on safe exploration in embodied reinforcement learning sits at the boundary between a promising component and a consequential decision. The core question is how agents can learn informative behaviour while respecting physical, social, and operational constraints. Answering it requires more than assembling favorable results. It requires a chain of evidence that makes the problem boundary, mechanism, measurement, comparator, uncertainty, and accountable decision visible to the reader.

This article presents a structured narrative review of ten related publications. Sources were selected for topical relevance, traceable publication metadata, and complementary coverage of technical, empirical, and governance questions. No meta-analytic pooling is attempted because the included studies differ in designs, populations, system boundaries, and outcomes. The purpose is decision-oriented synthesis: to identify what each source can support, where transfer remains uncertain, and which evidence would justify the next stage of use.

The central risk is treating simulated constraint satisfaction as proof of physical-world safety. A top-line metric or broad policy label can appear decisive while concealing the conditions that produced it. Accordingly, the synthesis separates observation from interpretation and implementation from validation. This separation is especially important when the same system creates benefits for one user or institution while transferring cost, risk, or administrative burden to another.

2. Evidence Base and Problem Boundary

A defensible review begins by defining the unit of analysis. For the present topic, the relevant boundary includes inputs, institutional or technical context, the mechanism producing an output, the person or system acting on that output, and the downstream outcome. Evidence falls outside the boundary when it measures only an intermediate proxy yet is used to claim an end-to-end benefit.

Alshiekh et al. (2018) addresses "Safe Reinforcement Learning via Shielding." In this synthesis, that work is used as a source on problem definition within safe exploration in embodied reinforcement learning. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for robotics and autonomous embodied systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Brunke et al. (2022) addresses "Safe Learning in Robotics: From Learning-Based Control to Safe Reinforcement Learning." In this synthesis, that work is used as a source on conceptual boundary within safe exploration in embodied reinforcement learning. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for robotics and autonomous embodied systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Zanon and Gros (2020) addresses "Safe Reinforcement Learning Using Robust MPC." In this synthesis, that work is used as a source on measurement within safe exploration in embodied reinforcement learning. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for robotics and autonomous embodied systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Together, these works provide complementary entry points, but their conclusions should not be flattened into a single ranking. Differences in data generation, setting, temporal horizon, and outcome definition

may be substantively meaningful. A review should therefore preserve those differences and state where analogy, rather than direct replication, connects one domain to another.

3. Mechanism and System Design

The operative mechanism is the chain connecting evidence inputs, safe exploration in embodied reinforcement learning, intermediate outputs, and downstream decisions. This formulation discourages component-only reasoning. A model, platform, policy, assay, or infrastructure service acquires practical meaning only when its interfaces and use conditions are specified. The evidence chain should describe what information is available, what transformation occurs, which assumptions make that transformation credible, and who is authorized to act.

Gu et al. (2024) addresses "A Review of Safe Reinforcement Learning: Methods, Theories, and Applications." In this synthesis, that work is used as a source on system mechanism within safe exploration in embodied reinforcement learning. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for robotics and autonomous embodied systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Marvi and Kiumarsi (2020) addresses "Safe reinforcement learning: A control barrier function optimization approach." In this synthesis, that work is used as a source on representation choice within safe exploration in embodied reinforcement learning. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for robotics and autonomous embodied systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Basso et al. (2021) addresses "Dynamic stochastic electric vehicle routing with safe reinforcement learning." In this synthesis, that work is used as a source on failure analysis within safe exploration in embodied reinforcement learning. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for robotics and autonomous embodied systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

These studies also show why modular claims require interface tests. A component may perform well while the complete pathway fails because inputs drift, users interpret outputs differently, recovery semantics are ambiguous, or institutional incentives change. Design documentation should therefore include not only the intended pathway but also foreseeable deviations, degraded modes, and conditions under which the system should stop or defer.

4. Evaluation and Uncertainty

Evaluation should center on construct validity, comparative performance, uncertainty, transfer across settings, and decision utility. Average performance remains useful, but it should be accompanied by distributions, error categories, relevant subgroups or operating regimes, and uncertainty at the point where a decision changes. Comparators must isolate the value of the proposed addition rather than compare a mature intervention with an intentionally weak baseline.

Thananjeyan et al. (2021) addresses "Recovery RL: Safe Reinforcement Learning With Learned Recovery Zones." In this synthesis, that work is used as a source on comparative evaluation within safe exploration in embodied reinforcement learning. The connection is deliberately bounded: the cited study informs the

evidence map, but it does not by itself establish readiness for robotics and autonomous embodied systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Yang et al. (2020) addresses "Safe reinforcement learning for dynamical games." In this synthesis, that work is used as a source on uncertainty within safe exploration in embodied reinforcement learning. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for robotics and autonomous embodied systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

External validation should introduce meaningful shifts in setting, time, population, workload, market structure, or environmental conditions. A nominally independent sample can still be too similar to test transfer. Conversely, a failed transfer is scientifically informative when the changed conditions are measured and the mechanism of failure is examined rather than hidden in an aggregate score.

5. Translation and Governance

Translation to robotics and autonomous embodied systems depends on more than technical readiness. Decision rights, documentation, maintenance, monitoring, appeal, and recovery are part of the intervention. The governance requirement is traceable evidence, named responsibility, version control, monitoring, appeal, and explicit revalidation. Each control should be connected to a named failure mode and should identify an owner, evidence source, review interval, and escalation path.

Zhang et al. (2021) addresses "Safe Reinforcement Learning With Stability Guarantee for Motion Planning of Autonomous Vehicles." In this synthesis, that work is used as a source on implementation within safe exploration in embodied reinforcement learning. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for robotics and autonomous embodied systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

GarcíaJavier and FernándezFernando (2015) addresses "A comprehensive survey on safe reinforcement learning." In this synthesis, that work is used as a source on governance within safe exploration in embodied reinforcement learning. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for robotics and autonomous embodied systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

A credible implementation plan also defines sunset and revalidation conditions. New data, software updates, institutional restructuring, population change, or shifts in incentives can invalidate previously reasonable assumptions. Monitoring should therefore test the validity of the evidence chain, not merely whether a service remains online or a headline metric remains stable.

6. Research Agenda

Future studies should preregister or otherwise predefine the intended decision, principal outcomes, exclusions, and analysis plan when feasible. They should report missingness, sensitivity analyses, negative or null findings, and the cost of errors to differently situated stakeholders. Reproducible artifacts should include sufficient metadata to reconstruct data provenance, model or analytical versions, parameter choices, and the sequence from evidence to conclusion.

Cross-disciplinary work needs a shared object of explanation rather than a collection of adjacent vocabularies. For safe exploration in embodied reinforcement learning, that shared object is the complete decision pathway. Technical, clinical, social, environmental, or economic expertise should each change the design or interpretation of the study. When evidence remains incomplete, the useful output is a bounded hypothesis, a transparent uncertainty statement, or a request for additional measurement.

7. Conclusion

The literature supports a disciplined approach to safe exploration in embodied reinforcement learning: define the decision, preserve system boundaries, test consequential failure modes, connect uncertainty to action, and assign accountable control. The strongest conclusion is not that one method is universally superior. It is that credible use in robotics and autonomous embodied systems requires an auditable link from source evidence to design choice, evaluation result, and governed decision.

Declarations

Ethics approval: Not applicable. This article synthesizes previously published literature and reports no new research involving humans, animals, human tissue, or identifiable sensitive data.

Consent to participate and consent for publication: Not applicable.

Funding: No external funding is reported for this commissioned synthesis.

Competing interests: The authoring group declares no competing interests for this commissioned synthesis.

AI-assisted writing: Generative AI supported drafting, source organization, and language editing. Accountable authors and editors must verify every claim, citation, and disclosure before publication.

Data, code, and materials availability: No new dataset or code was generated. All evidence is identified in the reference list.

Licence: This manuscript is prepared for publication under the Creative Commons Attribution 4.0 International licence (CC BY 4.0).

References

- Alshiekh, M., Bloem, R., Ehlers, R., Könighofer, B., Niekum, S., & Topcu, U. (2018). Safe Reinforcement Learning via Shielding. Scholarly publication. <https://doi.org/10.1609/aaai.v32i1.11797>
- Basso, R., Kulcsár, B., Sanchez-Díaz, I., & Qu, X. (2021). Dynamic stochastic electric vehicle routing with safe reinforcement learning. *Transportation Research Part E Logistics and Transportation Review*, 157, 102496. <https://doi.org/10.1016/j.tre.2021.102496>
- Brunke, L., Greeff, M., Hall, A. W., Yuan, Z., Zhou, S., Panerati, J., & Schoellig, A. P. (2022). Safe Learning in Robotics: From Learning-Based Control to Safe Reinforcement Learning. *Annual Review of Control Robotics and Autonomous Systems*, 5(1), 411-444. <https://doi.org/10.1146/annurev-control-042920-020211>
- GarcíaJavier, & FernándezFernando (2015). A comprehensive survey on safe reinforcement learning. *Journal of Machine Learning Research*. <https://doi.org/10.5555/2789272.2886795>
- Gu, S., Yang, L., Du, Y., Chen, G., Walter, F., Wang, J., & Knoll, A. (2024). A Review of Safe Reinforcement Learning: Methods, Theories, and Applications. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 46(12), 11216-11235. <https://doi.org/10.1109/tpami.2024.3457538>
- Marvi, Z., & Kiumarsi, B. (2020). Safe reinforcement learning: A control barrier function optimization approach. *International Journal of Robust and Nonlinear Control*, 31(6), 1923-1940. <https://doi.org/10.1002/rnc.5132>
- Thananjeyan, B., Balakrishna, A., Nair, S., Luo, M., Srinivasan, K., Hwang, M., Gonzalez, J. E., Ibarz, J., Finn, C., & Goldberg, K. (2021). Recovery

- RL: Safe Reinforcement Learning With Learned Recovery Zones. *IEEE Robotics and Automation Letters*, 6(3), 4915-4922.
<https://doi.org/10.1109/lra.2021.3070252>
- Yang, Y., Vamvoudakis, K. G., & Modares, H. (2020). Safe reinforcement learning for dynamical games. *International Journal of Robust and Nonlinear Control*, 30(9), 3706-3726. <https://doi.org/10.1002/rnc.4962>
- Zanon, M., & Gros, S. (2020). Safe Reinforcement Learning Using Robust MPC. *IEEE Transactions on Automatic Control*, 66(8), 3638-3652.
<https://doi.org/10.1109/tac.2020.3024161>
- Zhang, L., Zhang, R., Wu, T., Weng, R., Han, M., & Zhao, Y. (2021). Safe Reinforcement Learning With Stability Guarantee for Motion Planning of Autonomous Vehicles. *IEEE Transactions on Neural Networks and Learning Systems*, 32(12), 5435-5444.
<https://doi.org/10.1109/tnnls.2021.3084685>