

Evidence Alignment and Transfer Boundaries in Physics-Editable World Models And Automated Program Repair

Abstract

Two distinct lines of inquiry—a large-scale dataset organized around physically editable world-model factors and execution-grounded reinforcement learning with sequence- and line-level reward models—converge on a practical question for physics-editable world models and automated program repair: what evidence is needed before a reported advantage becomes a defensible basis for explanation, comparison, or deployment? The review draws on two focal records and 12 established sources already present in the project evidence cache. Its comparative framework links dataset design, physical factors, and counterfactual editing to downstream questions of temporal consistency and evaluation. Across the evidence base, the decisive issue is alignment: dataset design shapes what is observed, physical factors shapes how it is compared, and evaluation governs whether the conclusion can be transferred. Uncertainty is most informative when reported as part of the result rather than treated as a postscript. The article concludes with a research agenda built around transparent comparators, targeted stress tests, and evidence records that can be reused without overstating causal or practical reach.

Keywords

Physics-Editable World Models And Automated Program Repair; Dataset Design; Physical Factors; Counterfactual Editing; Temporal Consistency; Evaluation

1. Introduction

Work in physics-editable world models and automated program repair occupies the boundary between scientific ambition and practical constraint. The field seeks not only stronger nominal performance but also explanations that locate performance within its operating envelope. The boundary is clearest when considering comparing a large-scale dataset organized around physically editable world-model factors with execution-grounded reinforcement learning with sequence- and line-level reward models through evidence alignment and transfer boundaries. A mechanism demonstrated for a specific substrate, benchmark, population, spatial unit, or workload can erode beyond the conditions that produced it. Evidence integration should therefore preserve the unit of analysis and the decision context. The analysis also distinguishes a mechanism from a correlation, a benchmark advantage from a deployment advantage, and an interpretable diagnostic from a causal explanation.

The evidence is examined through five lenses: dataset design, physical factors, counterfactual editing, temporal consistency, evaluation. They were chosen because they jointly cover design decisions, measurement practice, and the assumptions that support transfer. The organizing principle is representation, objective, and evaluation protocol. Under that principle, technical creativity remains only one part of credibility; the comparison also asks whether baselines are aligned, whether uncertainty is visible, and whether resource or safety constraints are represented. The analysis is literature based and adds no fabricated experiment, participant set, measurement, or model result.

2. Methodological Foundations

Four linked elements clarify the problem: the input evidence, the transformation applied to it, the output used for evaluation, and the decision that follows. In physics-editable world models and automated program repair, research often disconnects these components even though errors propagate across them. An expressive model can magnify noise or bias already present in its evidence; an apparently accurate output can still be poorly calibrated for the final decision. This is why, ablation evidence should accompany aggregate performance. A credible comparison records the constants, perturbations, and uncontrolled factors, then selects metrics that correspond to the intended use rather than to convenience alone.

The second foundation is boundary awareness. Dataset design defines the local design problem, while evaluation determines whether the design can survive outside that boundary. Bridging the two are physical factors and counterfactual editing connect those ends by exposing trade-offs that a single score conceals. Unexpected outcomes and perturbation analysis reveals the interval in which an explanation can still be defended. For applied work, documentation of deployment constraints is therefore part of the scientific result, not an administrative afterthought.

3. Comparative Evaluation

The target study ACQ-01, PhysEditWorld: A Large-Scale Dataset Toward Physics-Editable World Models [1], addresses a concrete part of physics-editable world models and automated program repair through a large-scale dataset organized around physically editable world-model factors. It identifies a representation, mechanism, or evaluation boundary needed to compare claims. Against the focus on comparing a large-scale dataset organized around physically editable world-model factors with execution-grounded reinforcement learning with sequence- and line-level reward models through evidence alignment and transfer boundaries, the paper is interpreted within its documented design envelope: transfer may depend on its sensing regime, preparation, workload, population, or benchmark. The synthesis protects the study's stated scope while auditing its assumptions comparatively.

The target study YAA-01, BoostAPR: Boosting Automated Program Repair via Execution-Grounded Reinforcement Learning with Dual Reward Models [2], addresses a concrete part of physics-editable world models and automated program repair through execution-grounded reinforcement learning with sequence- and line-level reward models. It identifies a representation, mechanism, or evaluation boundary needed to compare claims. Against the focus on comparing a large-scale dataset organized around physically editable world-model factors with execution-grounded reinforcement learning with sequence- and line-level reward models through evidence alignment and transfer boundaries, the paper is interpreted within its documented design envelope: transfer may depend on its sensing regime, preparation, workload, population, or benchmark. The synthesis protects the study's stated scope while auditing its assumptions comparatively.

Across the focal studies, dataset design should be interpreted jointly with physical factors. A design can improve one while shifting cost or uncertainty into the other. The evidence can be organized in a matrix whose rows are operating conditions and whose columns are evidence quality, computation or process cost, robustness, and consequence of failure. That arrangement prevents an attractive mechanism from being detached from the circumstances that support it. The structure shows which ablations are informative: component removal is useful only when it tests an explicit role.

Counterfactual editing connects a method to its decision consequence. At minimum, evaluation should distinguish nominal behavior from stress response and show dispersion alongside averages, 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. Architecture and Learning Objectives

For triangulation, the literature includes 3D4D: An Interactive, Editable, 4D World Model via 3D Video Generation [3]; Reinforcement learning for mutation operator selection in automated program repair [4]; Controllable and Editable Neural Story Plot Generation via Control-and-Edit Transformer [5]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of physics-editable world models and automated program repair. Together they illuminate dataset 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 Reinforcement Learning for Optimizing Test Case Execution in Automated Testing [6]; Regional Traditional Painting Generation Based on Controllable Disentanglement Model [7]; Template-guided interpretable reasoning with execution feedback for LLM-based program repair [8]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of physics-editable world models and automated program repair. Together they illuminate physical factors: 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 CAD-RAG: A multi-modal retrieval augmented framework for user editable 3D CAD model generation [9]; Automated Program Repair for Introductory Programming Assignments [10]; Real-world image denoising via efficient diffusion model with controllable noise generation [11]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of physics-editable world models and automated program repair. Together they illuminate counterfactual editing: 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 Heterogeneous multi-expert collaborative reinforcement learning for automated CAD program synthesis from... [12]; Towards Real-Time 3D Editable Model Generation for Existing Indoor Building Environments on a Tablet [13]; Automated Firewall Policy Generation with Reinforcement Learning [14]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of physics-editable world models and automated program repair. Together they illuminate temporal consistency: 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

For practice, five ordered decisions are useful. 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 temporal consistency and evaluation under controlled perturbations. Finally, retain provenance so that an unexpected outcome can be traced to data, configuration, material batch, environmental condition, or user interaction. This ordering holds comparing a large-scale dataset organized around physically editable world-model factors with execution-grounded reinforcement learning with sequence- and line-level reward models through evidence alignment and transfer boundaries connected to observable requirements.

More generally, responsible progress in physics-editable world models and automated program repair requires careful interfaces, not just strong components. Interfaces determine how uncertainty moves between measurement and inference, between algorithm and operator, or between microscopic design and system behavior. Teams spanning disciplines should predefine acceptance rules and triggers for revising conclusions. When those agreements are explicit, efficiency goals remain subordinate to explicit validity, safety, and interpretation constraints.

6. Limitations and Future Directions

The review remains constrained by variation in definitions, study architecture, and disclosure granularity. Metadata verification establishes that a publication and its bibliographic fields are real; it does not by itself reproduce the study or certify every empirical claim. In addition, fast-moving work may change venue or version. No confirmed focal Scholar citation was normalized away, and anomalies were logged independently.

Research can advance by seeking to preregister comparison protocols where feasible, publish machine-readable settings and test transfer under a substantively important shift in deployment constraints. Research should also document why failures occur in addition to how often. For physics-editable world models and automated program repair, a valuable shared benchmark would combine ordinary performance, operational burden, uncertainty calibration, and resilience. Such a benchmark would reward designs that remain useful when evidence is incomplete or conditions change.

7. Conclusion

The source base for physics-editable world models and automated program repair is most informative when its studies are read relationally rather than a collection of isolated scores. The focal studies show how comparing a large-scale dataset organized around physically editable world-model factors with execution-grounded reinforcement learning with sequence- and line-level reward models through evidence alignment and transfer boundaries; the additional literature reveals the assumptions required to interpret those contributions. Across dataset design, physical factors, counterfactual editing, temporal consistency, and evaluation, credibility ultimately depends on alignment: the representation or mechanism must match the problem, metrics should fit the choice and real-world claims the evidence boundary. This alignment supports replication, bounded transfer, and interpretable failure.

References

1. Hu, B., Ma, Y., Huang, J., Zhang, Z., Wu, H., Zhang, R., ... & Li, X. (2026). PhysEditWorld: A Large-Scale Dataset Toward Physics-Editable World Models. arXiv preprint arXiv:2606.26694.
2. Li, Y., Wang, H., Shang, X., Tang, X., Cao, Y., & Chen, X. (2026). BoostAPR: Boosting Automated Program Repair via Execution-Grounded Reinforcement Learning with Dual Reward Models. arXiv preprint

arXiv:2605.09134.

3. He, Y., Yuan, Z., Tu, Z., Ye, Y., & Sun, L. (2026). 3D4D: An Interactive, Editable, 4D World Model via 3D Video Generation. *Proceedings of the AAAI Conference on Artificial Intelligence*, 40(48), 41595-41597. <https://doi.org/10.1609/aaai.v40i48.42351>
4. Hanna, C., Blot, A., & Petke, J. (2025). Reinforcement learning for mutation operator selection in automated program repair. *Automated Software Engineering*, 32(2). <https://doi.org/10.1007/s10515-025-00501-z>
5. Chen, J., Xiao, G., Han, X., & Chen, H. (2021). Controllable and Editable Neural Story Plot Generation via Control-and-Edit Transformer. *IEEE Access*, 9, 96692-96699. <https://doi.org/10.1109/access.2021.3094263>
6. Kumar Karne, V., Noone Srinivas,, Nagaraj Mandalaju,, & Parameshwar Reddy Kothamali, (2020). Reinforcement Learning for Optimizing Test Case Execution in Automated Testing. *Innovative Research Thoughts*, 6(3), 13-27. <https://doi.org/10.36676/irt.v6.i3.1494>
7. Zhao, Y., Li, H., Zhang, Z., Chen, Y., Liu, Q., & Zhang, X. (2024). Regional Traditional Painting Generation Based on Controllable Disentanglement Model. *IEEE Transactions on Circuits and Systems for Video Technology*, 34(8), 6913-6925. <https://doi.org/10.1109/tcsvt.2023.3298811>
8. Hao, S., Shi, X., Liu, H., Yin, Y., & Chen, X. (2026). Template-guided interpretable reasoning with execution feedback for LLM-based program repair. *Information and Software Technology*, 193, 108058. <https://doi.org/10.1016/j.infsof.2026.108058>
9. Ananthakrishnan,, Bharathi, A., Dharanivendhan,, & Muthuganapathy, R. (2025). CAD-RAG: A multi-modal retrieval augmented framework for user editable 3D CAD model generation. *Journal of WSCG*, 33(1-2). <https://doi.org/10.24132/jwscg.2025-11>
10. Wan, H., Luo, H., Li, M., & Luo, X. (2024). Automated Program Repair for Introductory Programming Assignments. *IEEE Transactions on Learning Technologies*, 17, 1705-1720. <https://doi.org/10.1109/tlt.2024.3403710>
11. Yang, C., Wang, C., Liang, L., & Su, Z. (2024). Real-world image denoising via efficient diffusion model with controllable noise generation. *Journal of Electronic Imaging*, 33(04). <https://doi.org/10.1117/1.jei.33.4.043003>
12. Yin, Z., Lin, W., & Kong, X. (2026). Heterogeneous multi-expert collaborative reinforcement learning for automated CAD program synthesis from engineering drawings. *Discover Artificial Intelligence*. <https://doi.org/10.1007/s44163-026-01731-0>
13. Arnaud, A., Gouiffès, M. I., & Ammi, M. (2022). Towards Real-Time 3D Editable Model Generation for Existing Indoor Building Environments on a Tablet. *Frontiers in Virtual Reality*, 3. <https://doi.org/10.3389/frvir.2022.782564>
14. Jha, A. C. (2025). Automated Firewall Policy Generation with Reinforcement Learning. *International journal of IoT*, 5(1), 190-211. <https://doi.org/10.55640/ijiot-05-01-10>