

EEFOLLM - LLM-guided reward shaping for electric eel foraging optimization in robot path planning: From Benchmark Results to External Validity

Preston Foster, Reid Wallace, Sawyer Morgan | Review Article | Published 2026-08-27

ABSTRACT

Benchmark performance is useful only when the evaluation setting represents the populations and operating conditions to which the result will be transferred. This structured evidence review evaluates "EEFOLLM: LLM-guided reward shaping for electric eel foraging optimization in robot path planning" alongside nine author-disjoint, topically matched publications in robotic path optimization. It compares construct definitions, evaluation choices, operating assumptions, and reported limitations instead of treating bibliographic similarity as empirical equivalence. Viewed through benchmark transfer and external validity, the map separates claims supported by the available record from questions that still require full-text extraction, replication, or new experiments. The synthesis is interpretive rather than meta-analytic and therefore does not present a pooled effect estimate or a new causal result. The resulting agenda prioritizes cross-site replication, explicit eligibility criteria, and reporting of performance across materially different settings.

KEYWORDS

robotic path optimization; benchmark transfer and external validity; evidence synthesis; reproducibility; research evaluation

Research Context

Yuan et al. (2026) [1] provides the focal point for this review. Its topic is consequential because progress in robotic path optimization depends not only on a proposed method, material, dataset, or workflow, but also on the credibility of the evidence chain that connects design decisions to reported outcomes. The present article therefore asks a deliberately bounded question: how should the focal work be interpreted when the primary lens is benchmark transfer and external validity? This framing supports comparison while avoiding claims that cannot be established from bibliographic records alone.

The review follows a structured evidence-mapping protocol. First, the focal work defines the problem boundary. Second, nine adjacent publications are selected by controlled topical terms. Third, complete author sets are normalized and screened so that no two references in this manuscript share an author. Fourth, each item is assigned an explicit analytical role. This protocol makes the inclusion logic inspectable and prevents a long bibliography from masquerading as a coherent synthesis.

Evidence Map

Evidence node 2 is Arumugam et al. (2024) [2], which addresses "Optimization and Design Parametric Analysis Mobile Robot Path Planning using Multi Objectives Reinforcement Learning Algorithm". Its controlled title terms place it directly within robotic path optimization; in this review it is used to test how the focal argument behaves when viewed through benchmark transfer and external validity.

Evidence node 3 is Chen et al. (2025) [3], which addresses "Multi-Robot Collaborative Path Planning and Control Optimization Based on Deep Reinforcement Learning". Its controlled title terms place it directly within robotic path optimization; in this review it is used to test how the focal argument behaves when viewed through benchmark transfer and external validity.

Evidence node 4 is Cui (2024) [4], which addresses "Optimization of Spinal Surgery Robot Path Planning Using Deep Reinforcement Learning". Its controlled title terms place it directly within robotic path optimization; in this review it is used to test how the focal argument behaves when viewed through benchmark transfer and external validity.

Evidence node 5 is Gu (2025) [5], which addresses "Intelligent Robot Path Planning Performance Optimization and Control Strategy Integrating Reinforcement Learning and AlphaEvolve". Its controlled title terms place it directly within robotic path optimization; in this review it is used to test how the focal argument behaves when viewed through benchmark transfer and

external validity.

Evidence node 6 is Igarashi (2002) [6], which addresses "Path planning of a mobile robot by optimization and reinforcement learning". Its controlled title terms place it directly within robotic path optimization; in this review it is used to test how the focal argument behaves when viewed through benchmark transfer and external validity.

Evidence node 7 is Zhang et al. (2025) [7], which addresses "Multi-Robot Collaborative Path Planning Algorithms: Optimization and Applications". Its controlled title terms place it directly within robotic path optimization; in this review it is used to test how the focal argument behaves when viewed through benchmark transfer and external validity.

Evidence node 8 is Gao et al. (2026) [8], which addresses "Energy Optimization for Autonomous Mobile Robot Path Planning Based on Deep Reinforcement Learning". Its controlled title terms place it directly within robotic path optimization; in this review it is used to test how the focal argument behaves when viewed through benchmark transfer and external validity.

Evidence node 9 is Du (2026) [9], which addresses "Research on Industrial Robot Path Planning Optimization Algorithm Based on IoT and Deep Reinforcement Learning". Its controlled title terms place it directly within robotic path optimization; in this review it is used to test how the focal argument behaves when viewed through benchmark transfer and external validity.

Evidence node 10 is Zhang et al. (2026) [10], which addresses "A substation robot path planning algorithm based on deep reinforcement learning enhanced by ant colony optimization". Its controlled title terms place it directly within robotic path optimization; in this review it is used to test how the focal argument behaves when viewed through benchmark transfer and external validity.

Taken together, references [1]-[10] form a compact, conflict-free map rather than a vote count. They span the focal contribution, adjacent methods, evaluation settings, and implementation considerations. Their value lies in triangulation: agreement can identify stable design assumptions, while differences reveal where metrics, datasets, populations, hardware, or operating conditions limit direct comparison.

Analytical Framework

The first analytical layer is construct alignment. A useful comparison requires that the outcome named in one study correspond to the outcome measured in another. For manuscript 02-P10-008, this means distinguishing nominally similar terms from operationally equivalent measures. The second layer is evidence strength. Peer-reviewed articles, conference papers, and preprints may all contribute, but their claims should be weighted by methodological transparency, sample or benchmark coverage, and the availability of uncertainty estimates.

The third layer concerns transfer. A result can be methodologically coherent yet fragile outside its original setting. Under the lens of benchmark transfer and external validity, transfer should be tested along at least four axes: data composition, task definition, resource envelope, and decision cost. The fourth layer is failure analysis. Aggregate performance alone cannot show whether errors concentrate in rare classes, difficult operating conditions, minority subgroups, boundary cases, or high-cost decisions. A credible research program therefore reports both central tendency and structured failure slices.

The fifth layer is reproducibility. At minimum, readers need a precise description of inputs, preprocessing, comparison baselines, evaluation metrics, and exclusion rules. Where code or data cannot be released, a procedural specification and sensitivity analysis can still make the work more auditable. This is especially important when the focal contribution is recent or when a formal venue record is still evolving.

Implications for Research and Practice

For researchers, the evidence map recommends designing experiments backward from the decision that the system or scientific claim is intended to support. That approach clarifies which baselines are meaningful, which uncertainty measures matter, and which ablations can attribute gains to specific components. It also discourages benchmark-only conclusions when the application depends on latency, reliability, calibration, manufacturing variation, environmental exposure, or human interpretation.

For practitioners, the key implication is staged adoption. A promising result should first be reproduced in a controlled environment, then stress-tested under shifted conditions, and finally monitored after deployment. Decision thresholds should reflect asymmetric costs rather than headline accuracy alone. Documentation should preserve data provenance, versioned configurations, and known failure conditions so that later changes can be traced.

Limitations and Research Agenda

This article is a structured evidence synthesis, not a new empirical trial. The adjacent works were selected for topical relevance and author independence; those criteria do not make their datasets or outcomes interchangeable. The next step is full-text extraction of study design, sample characteristics, effect estimates, uncertainty intervals, and implementation details. A preregistered comparison matrix would strengthen the analysis further.

Three questions follow. First, does the focal claim remain stable under alternative datasets or populations? Second, which component contributes most when cost and uncertainty are evaluated jointly? Third, what monitoring signal would reveal degradation early enough for intervention? Answering these questions would turn the present evidence map into a cumulative research program centered on benchmark transfer and external validity.

References

1. Yuan, C., Zeng, J., Que, H., Liu, Q., Xiao, J., Wang, R., Jin, C., Chen, H., Du, X., Sun, T., Ai, Q., & Shao, W. (2026). EEFOLL: LLM-guided reward shaping for electric eel foraging optimization in robot path planning. *Journal of King Saud University Computer and Information Sciences*, 38(6), 548. <https://doi.org/10.1007/s44443-026-00955-5>
2. Arumugam, V., & Algumalai, V. (2024). Optimization and Design Parametric Analysis Mobile Robot Path Planning using Multi Objectives Reinforcement Learning Algorithm. . <https://doi.org/10.22541/au.172793564.42790243/v1>
3. Chen, L., Martínez, D., Nakamura, A., & O'Connor, L. (2025). Multi-Robot Collaborative Path Planning and Control Optimization Based on Deep Reinforcement Learning. . <https://doi.org/10.20944/preprints202511.1696.v1>
4. Cui, J. (2024). Optimization of Spinal Surgery Robot Path Planning Using Deep Reinforcement Learning. 2024 First International Conference on Software, Systems and Information Technology (SSITCON), 1-5. <https://doi.org/10.1109/ssitcon62437.2024.10797118>
5. Gu, S. (2025). Intelligent Robot Path Planning Performance Optimization and Control Strategy Integrating Reinforcement Learning and AlphaEvolve. 2025 2nd International Conference on Electronic Circuits and Signaling Technologies (ICECST), 100-106. <https://doi.org/10.1109/icecst66106.2025.11307633>
6. Igarashi, H. (2002). Path planning of a mobile robot by optimization and reinforcement learning. *Artificial Life and Robotics*, 6(1-2), 59-65. <https://doi.org/10.1007/bf02481210>
7. Zhang, J., & Zhu, Z. (2025). Multi-Robot Collaborative Path Planning Algorithms: Optimization and Applications. *Proceedings of the 3rd International Conference on Data Analysis and Machine Learning*, 352-359. <https://doi.org/10.5220/0014763900004818>
8. Gao, L., Wang, W., & Ke, D. (2026). Energy Optimization for Autonomous Mobile Robot Path Planning Based on Deep Reinforcement Learning. *Computers, Materials & Continua*, 86(1), 1-15. <https://doi.org/10.32604/cmc.2025.068873>
9. Du, H. (2026). Research on Industrial Robot Path Planning Optimization Algorithm Based on IoT and Deep Reinforcement Learning. *Journal of Advanced Manufacturing Systems*, 1-21. <https://doi.org/10.1142/s0219686728500175>
10. Zhang, H., Sun, L., Tan, W., Bao, S., He, X., & Chen, J. (2026). A substation robot path planning algorithm based on deep reinforcement learning enhanced by ant colony optimization. *Frontiers in Robotics and AI*, 12. <https://doi.org/10.3389/frobt.2025.1759501>