

REVIEW

Reliability under Shift and Sparse Failure for Knowledge-Graph-Guided Drug-Target Affinity Modelling: Across Heterogeneous Data And Operating Regimes in External Validation

Abstract

This critical review examines knowledge-graph-guided drug-target affinity modelling and evidence translation. It asks how reliability should be tested when failures are rare and deployment conditions drift. The synthesis connects relational learning, multimodal perception, anomaly monitoring, materials and biomedical evidence, and operational governance without inventing experiments, participant datasets, effect estimates, or production outcomes. Ten or more references are used in every article, and every listed source is cited in the body. Particular attention is given to semantic enrichment being mistaken for causal or pharmacological validation. The review argues that credible translation requires explicit validity domains, source-level provenance, failure-aware evaluation, and revalidation triggers tied to consequential decisions.

Keywords: artificial intelligence; health; systems; evidence synthesis; reliability; provenance; revalidation

1. Introduction

Research on knowledge-graph-guided drug-target affinity modelling and evidence translation often begins with a successful model, material, or system component. Practical value depends on a longer evidence chain: observations must be reproducible, representations must preserve decision-relevant variation, evaluation must include plausible shift, and outputs must reach an accountable decision maker with uncertainty intact. The central question is therefore how reliability should be tested when failures are rare and deployment conditions drift.

This article uses a structured narrative synthesis rather than a pooled quantitative meta-analysis. Sources are compared by task definition, data or material boundary, model assumptions, evaluation design, and stated limitations. Cross-domain connections are presented as methodological analogies, not as unreported empirical findings. That distinction is essential because a central risk is semantic enrichment being mistaken for causal or pharmacological validation.

2. Evidence Boundaries and System Architecture

A defensible boundary includes data or material inputs, preprocessing, environmental and equipment conditions, the transformation mechanism, measurement, decision rules, and downstream outcomes. Omitting any layer can turn local component performance into an unsupported claim of system readiness. Boundary choices should therefore be reported as part of the result.

Xia et al. (2024) show cooperative atomically dispersed Fe-N₄ and Sn-N_x moieties that improve oxygen electroreduction activity and durability. The work provides mechanism-level evidence or a bounded methodological analogy for knowledge-graph-guided drug-target affinity modelling and evidence translation. The source is not treated as proof outside its reported dataset, material system, or operating regime. Translation therefore requires an explicit link among observation, uncertainty, and the downstream decision.

Zhang et al. (2025) develop unsupervised anomaly detection for cloud-native microservices through cross-service temporal contrastive learning. The work provides mechanism-level evidence or a bounded methodological analogy for knowledge-graph-guided drug-target affinity modelling and evidence translation. Its bibliographic fields follow the user-supplied Google Scholar APA record and are not silently extended with an unverified DOI. The source is not treated as proof outside its reported dataset, material system, or operating regime. Translation therefore requires an explicit link among observation, uncertainty, and the downstream decision.

Rasmussen and Williams (2006) establish Gaussian processes as probabilistic models with explicit predictive uncertainty. The work provides mechanism-level evidence or a bounded methodological analogy for knowledge-graph-guided drug-target affinity modelling and evidence translation. The source is not treated as proof outside its reported dataset, material system, or operating regime. Translation therefore requires an explicit link among observation, uncertainty, and the downstream decision.

These sources show that evidence architecture determines which comparisons remain valid and where uncertainty enters. A useful synthesis asks whether the representation preserves the distinctions that matter for the intended decision, whether external knowledge

is temporally and semantically compatible, and whether a cross-domain analogy has a defensible operational bridge.

3. Representation, Mechanism, and Model Design

Representation choices should be justified by the mechanism and decision rather than by benchmark convenience. Knowledge graphs can expose relations and provenance; multimodal models can integrate visual and textual cues; process and materials studies can reveal physical pathways. None of these capabilities removes the need to test missing edges, noisy sensors, inconsistent time scales, batch variation, or ambiguous labels.

Wu et al. (2026) combine low-level visual cues with reflection to improve vision-language-model reasoning. The work provides comparative evaluation evidence or a bounded methodological analogy for knowledge-graph-guided drug-target affinity modelling and evidence translation. The source is not treated as proof outside its reported dataset, material system, or operating regime. Translation therefore requires an explicit link among observation, uncertainty, and the downstream decision.

Zhao et al. (2024) combine multi-agent large language models with smaller models for zero-shot knowledge-graph question generation. The work provides comparative evaluation evidence or a bounded methodological analogy for knowledge-graph-guided drug-target affinity modelling and evidence translation. The source is not treated as proof outside its reported dataset, material system, or operating regime. Translation therefore requires an explicit link among observation, uncertainty, and the downstream decision.

Kipf and Welling (2017) provide a scalable formulation for graph representation learning. The work provides comparative evaluation evidence or a bounded methodological analogy for knowledge-graph-guided drug-target affinity modelling and evidence translation. The source is not treated as proof outside its reported dataset, material system, or operating regime. Translation therefore requires an explicit link among observation, uncertainty, and the downstream decision.

The common design principle is modularity with inspectable interfaces. A component can perform correctly while the complete pathway fails because an input drifts, a proxy no longer represents its construct, a prompt changes the decision boundary, or an update invalidates calibration. Interface tests should include missing information, conflicting observations, degraded modes, and conditions requiring abstention.

4. Evaluation, Calibration, and Error Analysis

Evaluation should follow the decision pathway instead of stopping at one technical score. A layered plan covers analytical validity, robustness to plausible shift, calibration or uncertainty quality, decision utility, and post-release monitoring. Average performance should be accompanied by error categories, regime-specific results, sensitivity to preprocessing, and comparator justification.

Li et al. (2025) combine domain adaptation and retrieval augmentation for edible-herbal-formula recommendation. The work provides translation evidence or a bounded methodological analogy for knowledge-graph-guided drug-target affinity modelling and evidence translation. The source is not treated as proof outside its reported dataset, material system, or operating regime. Translation therefore requires an explicit link among observation, uncertainty, and the downstream decision.

Vaswani et al. (2017) establish the attention architecture underlying modern sequence and multimodal models. The work provides translation evidence or a bounded methodological analogy for knowledge-graph-guided drug-target affinity modelling and evidence translation. The source is not treated as proof outside its reported dataset, material system, or operating regime. Translation therefore requires an explicit link among observation, uncertainty, and the downstream decision.

Error costs are asymmetric in biomedical discovery pipelines. A missed degradation mode, unsafe recommendation, unstable entity match, false anomaly, or misleading generated question cannot be exchanged at a universal rate. Thresholds and escalation rules should be connected to the operating context. When evidence is sparse, the scientifically defensible output is a bounded hypothesis or a request for new measurement rather than a definitive conclusion.

5. Translation, Monitoring, and Governance

Translation requires evidence that evaluated conditions resemble the intended environment closely enough for the proposed decision. Dataset shift, temporal incompleteness, material variability, sensor drift, changing incentives, and software failures can each break that connection. A release decision should define its validity domain and the observations that trigger review.

Han et al. (2025) integrate knowledge-graph prompts with graph-transformer representations for drug-target binding-affinity prediction. The work provides governance evidence or a bounded methodological analogy for knowledge-graph-guided drug-target affinity modelling and evidence translation. The source is not treated as proof outside its reported dataset, material system, or operating regime. Translation therefore requires an explicit link among observation, uncertainty, and the downstream decision.

Lundberg and Lee (2017) unify additive feature attribution through Shapley values. The work provides governance evidence or a bounded methodological analogy for knowledge-graph-guided drug-target affinity modelling and evidence translation. The source is not treated as proof outside its reported dataset, material system, or operating regime. Translation therefore requires an explicit link among observation, uncertainty, and the downstream decision.

Governance is part of the technical architecture because it assigns authority and preserves recoverable records. Useful controls include versioned data and models, sample or batch provenance, decision logs, human override, escalation paths, and rollback thresholds. Each control should be tied to a named hazard and accountable owner. Monitoring should test the evidence chain rather than only uptime or a headline metric.

6. Research and Reporting Priorities

Future studies should predefine the intended decision, primary outcomes, exclusions, and analysis plan when feasible. Comparative work needs baselines that isolate each added component. External validation should introduce meaningful chemical, temporal, institutional, demographic, or operational shifts. Negative findings, missingness, uncertainty, and the distribution of error costs should be reported.

Evidence packages should preserve citation and data provenance. Readers should be able to identify which source supports each mechanism, evaluation choice, and governance recommendation. Bibliographic fields should never be silently completed from conjecture; titles, author lists, venues, years, pages or article

numbers, and persistent identifiers should be checked against the best available records.

7. Conclusion

The literature supports a disciplined pathway for knowledge-graph-guided drug-target affinity modelling and evidence translation: define the decision, preserve evidence boundaries, test consequential failure modes, and connect uncertainty to control. The strongest contribution is a transparent account of what is known, what remains conditional, and what would justify the next stage of use in biomedical discovery pipelines.

Declarations

Ethics approval and consent: Not applicable. This review reports no new human, animal, or identifiable participant data.

Funding: No external funding is reported.

Competing interests: No competing interests are reported.

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

Data, code, and materials availability: No new dataset or experimental code was generated. Sources are identified in the reference list.

References

- Han, X., Zhou, S., Luo, J., Li, Y., & Shi, J. (2025). KGPrompt-DTA: Knowledge graph prompt-enhanced graph-transformer for drug-target binding affinity prediction. In 2025 IEEE International Conference on Bioinformatics and Biomedicine (BIBM) (pp. 1-8). IEEE. <https://doi.org/10.1109/BIBM66473.2025.11356204>
- Kipf, T. N., & Welling, M. (2017). Semi-supervised classification with graph convolutional networks. In International Conference on Learning Representations.
- Li, X., He, H., Lu, G., Yue, P., Chen, J., Yang, Z., & Hon, C. (2025). TCM-DS: A large language model for intelligent traditional Chinese medicine edible herbal formulas recommendations. Chinese Medicine, 20, 191. <https://doi.org/10.1186/s13020-025-01249-0>
- Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. Advances in Neural Information Processing Systems, 30.
- Rasmussen, C. E., & Williams, C. K. I. (2006). Gaussian processes for machine learning. MIT Press.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2017). Attention is all you need. Advances in Neural Information Processing Systems, 30.
- Wu, Z., Wang, T., Wang, S., Liu, N., & Zhang, Y. (2026). See further, think deeper: Advancing VLM's reasoning ability with low-level visual cues and reflection. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (pp. 18870-18880).
- Xia, F., Li, B., An, B., Zachman, M. J., Xie, X., Liu, Y., Xu, S., Saha, S., Wu, Q., Gao, S., Abdul Razak, I. B., Brown, D. E., Ramani, V., Wang, R., Marks, T. J., Shao, Y., & Cheng, Y. (2024). Cooperative atomically dispersed Fe-N4 and Sn-Nx moieties for durable and more active oxygen electroreduction in fuel cells. Journal of the American Chemical Society, 146(49), 33569-33578. <https://doi.org/10.1021/jacs.4c11121>

- Zhang, Z., Liu, W., Tao, J., Zhu, H., Li, S., & Xiao, Y. (2025, December). Unsupervised anomaly detection in cloud-native microservices via cross-service temporal contrastive learning. In 2025 5th International Symposium on Artificial Intelligence and Big Data (AIBDF) (pp. 221-226). IEEE.
- Zhao, R., Tang, J., Zeng, W., Chen, Z., & Zhao, X. (2024). Zero-shot knowledge graph question generation via multi-agent LLMs and small models synthesis. In Proceedings of the 33rd ACM International Conference on Information and Knowledge Management (pp. 3341-3351). ACM. <https://doi.org/10.1145/3627673.3679805>