

Scientific Knowledge Graphs Should Preserve Negative Evidence

Douglas Ward, Peter Torres, Kyle Peterson*

* Corresponding author: Kyle Peterson

Review Article | Frontiers in Integrative Science | ISCSR Research Publishing | Published 23 September 2026

ABSTRACT

Reliable deployment in experimental campaigns under costly pressure-temperature conditions depends on more than obtaining a strong benchmark result. This conceptual analysis uses negative and null evidence to study how claims travel from data to model output and then to action. Its central thesis is that the unit of assurance must be the linked history of sample preparation, pressure-temperature path, instrument output, code version, fit, and scientific claim. The reviewed evidence shows recurring risks from hidden distribution change, correlated evaluation error, missing provenance, and optimization objectives that omit downstream costs. In response, the article proposes a layered evaluation program combining controlled perturbations, subgroup and scenario analysis, repeated runs, calibration or selective prediction, and monitoring after release. It also asks who can inspect, override, and learn from failures. By integrating the assigned target papers with established scholarship, the synthesis clarifies which findings transfer across domains and which remain local to a benchmark, dataset, or experimental apparatus. The goal is a testable research program for bounded, traceable, and revisable systems.

KEYWORDS

scientific knowledge graphs should preserve negative evidence; scientific; negative; experimental; pressure-temperature; benchmark; output

Introduction

How should failed syntheses and ambiguous spectra guide later inference? The problem resists a learned component-only answer; once deployed, a predictor becomes part of a wider decision process. The visible result sits at the end of choices about measurement, encoding, comparison, computation, and intervention. In experimental campaigns under costly pressure-temperature conditions, these choices interact with institutions and physical processes that the model only partially represents. A high score can coexist with fragile inputs, an evaluator that rewards the wrong surface feature, or an action rule that turns modest uncertainty into a costly decision. For this reason, the synthesis examines negative and null evidence as an evidence problem. What matters is an clearly stated account of the evidence needed to justify assurance. The article's emphasis on negative and null evidence makes that boundary observable and, in principle, auditable.

The argument proceeds from a workflow perspective. Its unit is the linked history of sample preparation, pressure-temperature path, instrument output, code version, fit, and scientific claim. The benefit is that it avoids attributing every problem to the predictor: attributing every failure to model architecture while ignoring data capture, labeling, retrieval, validation, interface design, or organizational incentives. The same view makes differences among the focal studies easier to interpret. Although their application settings differ, they each make some part of an inference chain more explicit - a reward signal, a graph relation, a physical boundary condition, a confidence gate, or a revision trigger. The practical question is whether that explicit component remains meaningful when instrument traces, diffraction and spectroscopy products, simulation outputs, laboratory records, and analysis repositories change, and whether the proposed safeguards lead to replication, reanalysis, anomaly review, and explicit preservation of competing hypotheses rather than merely producing another metric. This requirement gives practical content to the article's central question: How should failed syntheses and ambiguous spectra guide later inference?

Conceptual Foundations

Reliability analysis must not collapse aleatory variability, epistemic uncertainty, and strategic response into one category. Aleatory variation concerns irreducible differences among cases. Epistemic uncertainty concerns missing knowledge, limited samples, or unsupported regions of the input space. Strategic adaptation occurs when users, adversaries, markets, or institutions respond to the system itself. These sources demand different controls. Repeated measurement can characterize variation; new evidence and conservative extrapolation can reduce epistemic uncertainty; red teaming and feedback-aware assessment address adaptation. Combining them into one confidence score hides which intervention is appropriate. A defensible system therefore reports uncertainty in a form tied to an available response rather than presenting confidence as a decorative number. This requirement gives practical content to the article's central question: How should failed syntheses and ambiguous spectra guide later inference?

Four analytically distinct layers organize the problem. The evidence layer records observations and their provenance. The representation layer turns those observations into features, graphs, tokens, or simulated states. The inference layer produces predictions, explanations, translations, anomaly scores, or candidate actions. The decision layer maps those outputs to replication, reanalysis, anomaly review, and documented preservation of competing hypotheses. Reliability can fail at any boundary. Better inference cannot repair evidence that omitted the relevant condition, and a calibrated probability cannot rescue an action rule whose costs were never specified. Conversely, a conservative decision policy may reduce harm even when the underlying model remains imperfect. This layered view makes negative and null evidence testable because each claim can be assigned to a component and a measurable transition. The implication is especially consequential in experimental campaigns under costly pressure-temperature conditions, where a small modeling choice can redirect attention and resources.

What the Core Studies Establish

MuSK: Multi-Scale Knowledge Learning for Provenance-Graph Anomaly Detection asks how stealthy multi-stage attacks can be detected in heterogeneous system-provenance graphs without abundant attack labels. It uses self-supervised recovery of masked node attributes, meta-path edges, and positional structure, followed by meta-path-guided subgraph verification. The relevant evidence is that evaluation spans nine provenance settings, including DARPA Transparent Computing hosts, and reports high precision, complete recall on those hosts, and efficiency gains from cached incremental expansion. The design joins local attributes, heterogeneous causal paths, and global position while moving decisions from isolated nodes to suspicious subgraphs. In the present review, this work matters because negative and null evidence cannot be evaluated as an abstract virtue; it must change how evidence is collected and how an action is withheld. Performance may depend on audit completeness, benign workload diversity, attack coverage, and the stability of hand-selected meta-path semantics. (Ma et al. 2026).

A useful test case is BoostAPR: Boosting Automated Program Repair via Execution-Grounded Reinforcement Learning with Dual Reward Models. Rather than treating its output as an isolated score, the study frames how reinforcement learning for program repair can overcome sparse execution feedback and coarse sequence-level credit. Its central mechanism is a three-stage pipeline combining execution-verified supervised traces, sequence- and line-level reward models, and PPO with reward redistribution to critical edit regions, and its evidential basis is that evaluation spans SWE-bench Verified, Defects4J transfer, HumanEval-Java, and QuixBugs, with reported gains in both repository repair and cross-language generalization. Line-level credit assignment makes test execution informative at the granularity where patches actually change behavior. For experimental campaigns under costly pressure-temperature conditions, the transferable lesson is methodological: a system should preserve the path from signal to decision and expose the conditions under which that path may fail. The boundary is equally important. Benchmark success does not by itself establish maintainability, specification fidelity, or robustness to incomplete and flaky tests. (Li et al. 2026).

Evidence for negative and null evidence becomes concrete in Synthesis and Stability of High-Energy-Density Niobium Nitrides under High-Pressure Conditions. The paper addresses which nitrogen-rich niobium phases can be synthesized and stabilized under extreme pressure through high-pressure synthesis and structural stability analysis of niobium nitrides motivated by dense nitrogen bonding and recoverable energy storage. Its evaluation is informative because the publication establishes phase synthesis and stability as the relevant evidence chain rather than inferring performance from composition alone. It places niobium nitrides in the broader search for high-energy-density polynitrogen solids. This does not make the method a universal component; it identifies a design hypothesis that must be re-tested in the article's focal setting. In particular, recovery, kinetic persistence, scale-up, and safe energy release remain distinct questions from phase

stability in a high-pressure cell. The appropriate use of the result is therefore bounded, comparative, and explicit about unresolved deployment conditions (Chen et al. 2025).

Relevant Foundations

Several established literatures clarify the design space. Backtracking work established causal system history as a practical resource for intrusion investigation (King and Chen 2003). HOLMES demonstrates the value of correlating low-level events into higher-level attack stages (Milajerdi et al. 2019). UNICORN uses provenance structure to learn long-running system behavior and surface persistent threats (Han et al. 2020). Whole-system provenance research exposes the engineering tension among capture completeness, query utility, and runtime overhead (Pasquier et al. 2017). Together these findings imply that negative and null evidence should be evaluated against explicit alternatives and failure costs. In *Scientific Knowledge Graphs Should Preserve Negative Evidence*, this literature supplies a specific test of negative and null evidence within experimental campaigns under costly pressure-temperature conditions.

The evidence base offers complementary tests rather than one universal metric. Relational graph convolution provides a foundation for treating typed edges as distinct evidence channels (Schlichtkrull et al. 2018). Inductive graph representation learning is essential when new entities arrive after training (Hamilton, Ying, and Leskovec 2017). Graph attention makes neighborhood weighting learnable but does not by itself guarantee causal relevance (Velickovic et al. 2018). They also caution against interpreting one benchmark, one split, or one explanatory artifact as a complete validation argument. In *Scientific Knowledge Graphs Should Preserve Negative Evidence*, this literature supplies a specific test of negative and null evidence within experimental campaigns under costly pressure-temperature conditions.

A credible assurance case can be built from findings that operate at different levels. Anomaly-detection research distinguishes point, contextual, and collective anomalies, a distinction that maps naturally to provenance subgraphs (Chandola, Banerjee, and Kumar 2009). Open-world security analysis warns that laboratory labels and stationary traffic assumptions rarely survive deployment (Sommer and Paxson 2010). Automated threat-intelligence extraction shows both the reach and the noise of public indicator sources (Liao et al. 2016). Their combined value lies in making assumptions visible enough to challenge before deployment and to revise after failure. In *Scientific Knowledge Graphs Should Preserve Negative Evidence*, this literature supplies a specific test of negative and null evidence within experimental campaigns under costly pressure-temperature conditions.

A Traceable System Architecture

The evidence architecture for experimental campaigns under costly pressure-temperature conditions begins with typed evidence. Each record should identify its source, time, collection conditions, transformations, and relation to other records. Graphs are useful when relations carry meaning, but graph construction is itself a hypothesis. An edge may denote temporal adjacency, physical causation, code dependency, shared infrastructure, or analyst assertion; treating these as interchangeable creates false certainty. The architecture should therefore preserve edge types and confidence, retain alternative links when evidence is ambiguous, and version both the data schema and the rules that construct the graph. A model may aggregate this structure, but the original evidence must remain recoverable for audit. The implication is especially consequential in experimental campaigns under costly pressure-temperature conditions, where a small modeling choice can redirect attention and resources.

The next design choice concerns where additional computation is worth its cost. Easy, well-supported cases can take a fast path. Shifted, ambiguous, or high-cost cases should activate additional precision, retrieval, alternative drafts, simulation, or expert review. This conditional design is preferable to applying maximum computation everywhere because resource cost is part of deployment quality. It is also preferable to an unconditional fast path because efficiency can amplify a systematic error at scale. The trigger must be evaluated as carefully as the expensive model: coverage, false reassurance, subgroup effects, latency, and the consequences of abstention all matter. The output is not simply a prediction but a package containing evidence links, uncertainty, the action taken, and the conditions that caused escalation. The article's emphasis on negative and null evidence makes that boundary observable and, in principle, auditable.

Testing Claims Rather Than Scores

The first testing artifact should list the claims the operational tool is expected to support. Each intended claim - such as improved robustness, safer abstention, faster updating, or better structural localization - needs a target population, comparator, outcome, and boundary condition. Random splits are insufficient when related samples share a patient, repository, instrument run, season, organization, or simulated design family. Grouped and temporal splits better approximate the novelty encountered after field use. Stress tests should vary one factor at a time when attribution is important, then combine factors to measure interaction. Repeated runs are necessary whenever decoding, optimization, retrieval, or numerical kernels can vary. Reporting only the best seed or a pooled mean makes operational instability disappear. In experimental campaigns under costly pressure-temperature conditions, this distinction should be tested before it changes an operational decision.

Metric choice must be derived from the decision. Discrimination measures whether cases are ranked correctly; calibration asks whether confidence corresponds to observed frequency; selective risk asks what error remains after deferral; robustness measures performance across defined perturbations; and utility assigns costs to actions. These are not interchangeable. For negative and null evidence, the minimum report should include overall performance, slice-level results, uncertainty intervals, coverage-risk curves, stability across runs, and failure examples linked back to instrument traces, diffraction and spectroscopy products, simulation outputs, laboratory records, and analysis repositories. If the system updates incrementally, the validation must also test forgetting, stale evidence, and rollback. If an automatic judge supplies labels, a sample needs independent human review and content-preserving perturbations that reveal whether the judge is grading substance. This point narrows the research task for experimental campaigns under costly pressure-temperature conditions: compare alternatives under the same information and resource constraints.

Deployment Controls

Operational rules are the governance expression of the evidence. A field use specification should name allowed uses, prohibited uses, escalation thresholds, data-retention rules, and the person or role that can halt the operational tool. Monitoring should track input shift, missingness, abstention, disagreement, latency, overrides, and downstream outcomes. A rising override rate may indicate model drift, a changed operational sequence, or new user behavior; treating it only as operator noncompliance wastes evidence. Incident review should reconstruct the entire chain, including the learned component and the surrounding interface. Versioned models without versioned prompts, retrieval indexes, rules, and datasets do not support reliable reconstruction. In experimental campaigns under costly pressure-temperature conditions, this distinction should be tested before it changes an operational decision.

A credible human-review process must be specified and tested. Reviewers need enough context to challenge the output, but not so much generated rationale that weak evidence is obscured by volume. Escalation queues require prioritization and workload limits. A reject option that sends nearly everything to experts is safe only on paper; a reject option that rarely fires may be equally ineffective. For experimental campaigns under costly pressure-temperature conditions, oversight should therefore be evaluated with realistic staffing, time pressure, and disagreement. The system should record the final action and its reason without turning that record into surveillance unrelated to the stated purpose. Contestability, privacy, and security are properties of this institutional process as much as of the estimator. This point narrows the research task for experimental campaigns under costly pressure-temperature conditions: compare alternatives under the same information and resource constraints.

Experiments the Field Still Needs

Research should begin with cases selected for what they can reveal about failure, not merely for availability. Cases should represent unsupported regions, ambiguous evidence, conflicting modalities, rare but costly conditions, and realistic adversarial transformations. Each case needs provenance and a reason for inclusion. The second priority is intervention testing: when uncertainty is detected, compare additional computation, retrieval, alternative reasoning paths, model ensembles, and human escalation under the same resource budget. This reveals whether a proposed safeguard actually changes the decision or only changes the explanation. The third priority is transfer. A method should be re-evaluated across sites, devices, languages, organizations, or physical regimes before broad claims are made. The implication is especially consequential in experimental campaigns under costly pressure-temperature conditions, where a small modeling choice can redirect attention and resources.

Long-term progress depends on cumulative records of both successful and failed approaches. Benchmarks should preserve negative results, failed branches, and changed definitions so later work does not learn only from successful demonstrations. Shared reporting should include data lineage, versioned assessment code, confidence intervals, and enough error material to reproduce major conclusions. Prospective studies should predefine monitoring and stopping rules, then measure how people adapt to the system. Finally, researchers should compare simple baselines with complex architectures under equivalent information. Complexity is justified when it adds a measurable capability - for example, structural localization, calibrated deferral, or incremental updating - not merely when it creates a richer narrative around the same errors. This requirement gives practical content to the article's central question: How should failed syntheses and ambiguous spectra guide later inference?

Conclusion

Negative and null evidence is credible only when it is attached to an explicit evidence chain and decision policy. The review identifies several concrete mechanisms: structured exploration, typed graphs, conditional revision, adaptive computation, physical constraints, and repeated assessment. At the same time, success on the originating benchmark cannot define a safe deployment boundary. For experimental campaigns under costly pressure-temperature conditions, the practical standard should be a traceable operational sequence that measures shift and instability, supports replication, reanalysis, anomaly review, and explicit preservation of competing hypotheses, and preserves enough evidence for an independent reviewer to reconstruct failure. Future work should compare safeguards under realistic resource and staffing limits, publish negative and subgroup results, and treat monitors and automated judges as components requiring their own validation. Reliability is demonstrated through warranted answers, consequential uncertainty, and a record that another observer can inspect. Seen through negative and null evidence, the issue is not merely technical; it determines which claims remain open to challenge.

References

1. Ma, Mingjun, et al. "MuSK: Multi-Scale Knowledge Learning for Provenance-Graph Anomaly Detection." *Computer Networks* 289 (2026): 112728.
2. Li, Yuanhao, et al. "BoostAPR: Boosting Automated Program Repair via Execution-Grounded Reinforcement Learning with Dual Reward Models." *arXiv preprint arXiv:2605.09134* (2026).
3. Chen, Huawei, et al. "Synthesis and Stability of High-Energy-Density Niobium Nitrides under High-Pressure Conditions." *Inorganic Chemistry* 64.1 (2025): 692-700.
4. King, Samuel T., and Peter M. Chen. "Backtracking Intrusions." *Proceedings of the Nineteenth ACM Symposium on Operating Systems Principles*, 2003, pp. 223-236.
5. Milajerdi, Sadegh M., et al. "HOLMES: Real-Time APT Detection through Correlation of Suspicious Information Flows." *2019 IEEE Symposium on Security and Privacy*, 2019, pp. 1137-1152.
6. Han, Xueyuan, et al. "UNICORN: Runtime Provenance-Based Detector for Advanced Persistent Threats." *Network and Distributed System Security Symposium*, 2020.
7. Pasquier, Thomas, et al. "Runtime Analysis of Whole-System Provenance." *Proceedings of the 2017 ACM SIGSAC Conference on Computer and Communications Security*, 2017, pp. 1601-1614.
8. Schlichtkrull, Michael, et al. "Modeling Relational Data with Graph Convolutional Networks." *The Semantic Web*, Springer, 2018, pp. 593-607.
9. Hamilton, William L., Rex Ying, and Jure Leskovec. "Inductive Representation Learning on Large Graphs." *Advances in Neural Information Processing Systems*, vol. 30, 2017.
10. Velickovic, Petar, et al. "Graph Attention Networks." *International Conference on Learning Representations*, 2018.
11. Chandola, Varun, Arindam Banerjee, and Vipin Kumar. "Anomaly Detection: A Survey." *ACM Computing Surveys*, vol. 41, no. 3, 2009, article 15.
12. Sommer, Robin, and Vern Paxson. "Outside the Closed World: On Using Machine Learning for Network Intrusion Detection." *2010 IEEE Symposium on Security and Privacy*, 2010, pp. 305-316.
13. Liao, Xiaojing, et al. "Acing the IOC Game: Toward Automatic Discovery and Analysis of Open-Source Cyber Threat Intelligence." *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*, 2016, pp. 755-766.