

Provenance Graphs for High-Pressure Experimental Evidence

Jacob Nelson, Gary Hill, Nicholas Campbell*

* Corresponding author: Nicholas Campbell

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

ABSTRACT

Reliable deployment in diamond-anvil experiments with multi-instrument observations depends on more than obtaining a strong benchmark result. This conceptual analysis uses scientific provenance 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

provenance graphs; high-pressure experimental evidence; provenance; experimental; benchmark; scientific; output

Introduction

How can causal data lineage help distinguish phase evidence from preparation artifacts? The difficulty begins with a basic observation: prediction is only one stage of a deployed process. Before an output appears, designers have already decided what counts as evidence, which representations are admissible, how alternatives are compared, and how uncertainty changes action. In diamond-anvil experiments with multi-instrument observations, these choices interact with institutions and physical processes that the predictive component 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. This makes it useful to treat scientific provenance as an evidence problem. Its purpose is to specify the observations and comparisons needed before reliability can be claimed. Seen through scientific provenance, the issue is not merely technical; it determines which claims remain open to challenge.

A pipeline view organizes the comparison. Its unit is the linked history of sample preparation, pressure-temperature path, instrument output, code version, fit, and scientific claim. This scope makes it harder to commit a familiar error: attributing every failure to model architecture while ignoring data capture, labeling, retrieval, validation, interface design, or organizational incentives. This framing places the focal studies on comparable evidential terms. Although their application settings differ, they each make some part of an inference chain more clearly stated - a reward signal, a graph relation, a physical boundary condition, a confidence gate, or a revision trigger. The synthesis therefore asks whether that clearly stated 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 clearly stated preservation of competing hypotheses instead of merely producing another metric. In diamond-anvil experiments with multi-instrument observations, this distinction should be tested before it changes an operational decision.

Methods and Boundaries in the Assigned Studies

The strongest contribution of MuSK: Multi-Scale Knowledge Learning for Provenance-Graph Anomaly Detection is not a generic claim that learning improves performance. It is the more specific proposition that how stealthy multi-stage attacks can be detected in heterogeneous system-provenance graphs without abundant attack labels. The proposed route is self-supervised recovery of masked node attributes, meta-path edges, and positional structure, followed by meta-path-guided subgraph verification. Support comes from the fact 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. Read through the lens of scientific provenance, the study also illustrates why a reported average must remain connected to the workflow that produced it. Performance may depend on audit completeness, benign workload diversity, attack coverage, and the stability of hand-selected meta-path semantics. (Ma et al. 2026).

BoostAPR: Boosting Automated Program Repair via Execution-Grounded Reinforcement Learning with Dual Reward Models asks how reinforcement learning for program repair can overcome sparse execution feedback and coarse sequence-level credit. It uses a three-stage pipeline combining execution-verified supervised traces, sequence- and line-level reward models, and PPO with reward redistribution to critical edit regions. The relevant evidence 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. In the present review, this work matters because scientific provenance cannot be evaluated as an abstract virtue; it must change how evidence is collected and how an action is withheld. Benchmark success does not by itself establish maintainability, specification fidelity, or robustness to incomplete and flaky tests. (Li et al. 2026).

A useful test case is Possible H₂O Storage in the Crystal Structure of CaSiO₃ Perovskite. Rather than treating its output as an isolated score, the study frames whether CaSiO₃ perovskite can host hydrogen at lower-mantle pressures and temperatures. Its central mechanism is laser-heated diamond-anvil synthesis from starting materials with different water contents, combined with in-situ diffraction and infrared observations, and its evidential basis is that non-cubic peak splitting, a possible OH vibration, and systematically smaller unit-cell volumes in hydrous syntheses provide convergent but indirect evidence of water storage. The work expands the lower-mantle water budget beyond bridgmanite and ferropericlase. For diamond-anvil experiments with multi-instrument observations, 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. The authors emphasize that exact water quantification and exclusion of signals from coexisting hydrous phases remain difficult under pressure. (Chen et al. 2020).

Evidence for scientific provenance becomes concrete in LARGER: Lexically Anchored Repository Graph Exploration and Retrieval. The paper addresses how coding agents can convert lexical matches into structurally meaningful repository context through lexical anchors aligned to a repository graph, followed by confidence-filtered expansion of local structural neighborhoods inside an existing command-line search loop. Its evaluation is informative because evaluation covers localization, test generation, and codebase understanding, with reported file-level retrieval gains on multiple benchmarks. The method integrates graph evidence without requiring a separate graph database or a disconnected traversal interface. 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, confidence thresholds and graph freshness can create silent omissions, especially in generated code, reflection-heavy systems, or repositories with incomplete build metadata. The appropriate use of the result is therefore bounded, comparative, and explicit about unresolved deployment conditions (Hu et al. 2026).

Separating Evidence, Inference, and Action

The analysis distinguishes evidence, representation, inference, and decision as separate layers. 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 clearly stated 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 scientific provenance testable because each claim can be assigned to a component and a measurable transition. The article's emphasis on scientific provenance makes that boundary observable and, in principle, auditable.

Three kinds of uncertainty require different responses: aleatory variation, epistemic limits, and strategic adaptation. 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 deployed process 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 validation 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. For the present focus, scientific provenance therefore becomes a criterion for deciding which evidence deserves further scrutiny.

Architecture for Bounded Automation

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 operational use 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. Seen through scientific provenance, the issue is not merely technical; it determines which claims remain open to challenge.

A traceable operational use in diamond-anvil experiments with multi-instrument observations 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 retain 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. In diamond-anvil experiments with multi-instrument observations, this distinction should be tested before it changes an operational decision.

Established Tests and Design Principles

Several established literatures clarify the design space. Automated threat-intelligence extraction shows both the reach and the noise of public indicator sources (Liao et al. 2016). ATT&CK provides a shared behavioral vocabulary, but its categories remain an evolving analyst model rather than ground truth (Strom et al. 2018). FAIR principles make provenance, identifiers, and reusable metadata part of scientific evidence quality (Wilkinson et al. 2016). Backtracking work established causal system history as a practical resource for intrusion investigation (King and Chen 2003). Together these findings imply that scientific provenance should be evaluated against explicit alternatives and failure costs. In Provenance Graphs for High-Pressure Experimental Evidence, this literature supplies a specific test of scientific provenance within diamond-anvil experiments with multi-instrument observations.

The evidence base offers complementary tests rather than one universal metric. 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). They also caution against interpreting one benchmark, one split, or one explanatory artifact as a complete validation argument. In Provenance Graphs for High-Pressure Experimental Evidence, this literature supplies a specific test of scientific provenance within diamond-anvil experiments with multi-instrument observations.

A credible assurance case can be built from findings that operate at different levels. 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). Their combined value lies in making assumptions visible enough to challenge before deployment and to revise after failure. In Provenance Graphs for High-Pressure Experimental Evidence, this literature supplies a specific test of scientific provenance within diamond-anvil experiments with multi-instrument observations.

Measurement Under Shift and Repetition

The action and its costs should determine the reported metrics. 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 scientific provenance, 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 operational tool 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 diamond-anvil experiments with multi-instrument observations: compare alternatives under the same information and resource constraints.

The first testing artifact should list the claims the application 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. For the present focus, scientific provenance therefore becomes a criterion for deciding which evidence deserves further scrutiny.

Priorities for Empirical Work

A second priority is to maintain the history of evidence rather than only final successes. Benchmarks should maintain negative results, failed branches, and changed definitions so later work does not learn only from successful demonstrations. Shared reporting should include data lineage, versioned evaluation 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 deployed operational sequence. 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. In diamond-anvil experiments with multi-instrument observations, this distinction should be tested before it changes an operational decision.

Future evaluation sets should be organized around plausible mechanisms of failure. 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 article's emphasis on scientific provenance makes

that boundary observable and, in principle, auditable.

Limits, Monitoring, and Contestability

Placing a person in the loop does not by itself create effective control. 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 diamond-anvil experiments with multi-instrument observations, 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 predictive component. This point narrows the research task for diamond-anvil experiments with multi-instrument observations: compare alternatives under the same information and resource constraints.

Evidence must eventually become a set of enforceable operating conditions. A field use specification should name allowed uses, prohibited uses, escalation thresholds, data-retention rules, and the person or role that can halt the application. Monitoring should track input shift, missingness, abstention, disagreement, latency, overrides, and downstream outcomes. A rising override rate may indicate model drift, a changed workflow, or new user behavior; treating it only as operator noncompliance wastes evidence. Incident review should reconstruct the entire chain, including the model and the surrounding interface. Versioned models without versioned prompts, retrieval indexes, rules, and datasets do not support reliable reconstruction. For the present focus, scientific provenance therefore becomes a criterion for deciding which evidence deserves further scrutiny.

Conclusion

Scientific provenance is credible only when it is attached to an explicit evidence chain and decision policy. Across the reviewed studies, several mechanisms can strengthen that chain: structured exploration, typed graphs, conditional revision, adaptive computation, physical constraints, and repeated evaluation. None of these results makes the original benchmark a complete map of safe use. For diamond-anvil experiments with multi-instrument observations, 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. The final standard is not uninterrupted automation but evidence-linked action with explicit deferral and independent verifiability. In diamond-anvil experiments with multi-instrument observations, this distinction should be tested before it changes an operational decision.

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. "Possible H₂O Storage in the Crystal Structure of CaSiO₃ Perovskite." *Physics of the Earth and Planetary Interiors* 299 (2020): 106412.
4. Hu, Yuntong, et al. "LARGER: Lexically Anchored Repository Graph Exploration and Retrieval." *arXiv preprint arXiv:2605.16352* (2026).
5. 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.
6. Strom, Blake E., et al. *MITRE ATT&CK: Design and Philosophy*. MITRE Corporation, 2018.
7. Wilkinson, Mark D., et al. "The FAIR Guiding Principles for Scientific Data Management and Stewardship." *Scientific Data*, vol. 3, 2016, article 160018.
8. King, Samuel T., and Peter M. Chen. "Backtracking Intrusions." *Proceedings of the Nineteenth ACM Symposium on Operating Systems Principles*, 2003, pp. 223-236.

9. 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.
10. Han, Xueyuan, et al. "UNICORN: Runtime Provenance-Based Detector for Advanced Persistent Threats." Network and Distributed System Security Symposium, 2020.
11. 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.
12. Schlichtkrull, Michael, et al. "Modeling Relational Data with Graph Convolutional Networks." The Semantic Web, Springer, 2018, pp. 593-607.
13. Hamilton, William L., Rex Ying, and Jure Leskovec. "Inductive Representation Learning on Large Graphs." Advances in Neural Information Processing Systems, vol. 30, 2017.
14. Velickovic, Petar, et al. "Graph Attention Networks." International Conference on Learning Representations, 2018.