

Code Lineages as Curricula for Language-Model Reasoning

Kathryn Perry, Gloria Powell, Teresa Long^{*}

^{*} Corresponding author: Teresa Long

Review Article | Journal of Algorithmic Discovery and Applied AI | ISCSR Research Publishing | Published 23 September 2026

ABSTRACT

Can revision history teach reasoning strategies that isolated examples cannot? This review answers by treating lineage-aware learning as a property of a sociotechnical workflow rather than a feature that can be read from average accuracy. The focal setting is software tasks represented as evolving graphs, where speed and fluency can conceal semantic loss, correlated self-evaluation errors, or domain-specific failure. Evidence from the assigned publications is synthesized with foundational studies of calibration, distribution shift, causal structure, and responsible deployment. Four requirements follow: preserve the lineage of source texts, alternative drafts, reward signals, confidence estimates, domain corpora, and human judgments; measure stability across relevant perturbations; connect confidence to a specific action; and maintain a route for human challenge and correction. The framework distinguishes descriptive performance from decision utility and separates uncertainty about the world from uncertainty created by the model and its evaluator. It also shows why faster inference or richer reasoning is valuable only when it improves a defined decision under a transparent resource budget. The article is a literature review and research agenda, not a report of a newly completed trial.

KEYWORDS

code lineages; language-model reasoning; reasoning; confidence; human; decision; uncertainty

Introduction

Can revision history teach reasoning strategies that isolated examples cannot? The central difficulty is that operational use turns a prediction into one component of an operational system. The visible result sits at the end of choices about measurement, encoding, comparison, computation, and intervention. In software tasks represented as evolving graphs, 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. The analysis therefore approaches lineage-aware learning as an evidence problem. The review asks which observations and counterfactual comparisons can support a defensible assurance claim. In software tasks represented as evolving graphs, this distinction should be tested before it changes an operational decision.

The argument proceeds from a pipeline perspective. Its unit is the trajectory from tokenization and draft generation through scoring, revision, compression, and release. 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. 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 documented - a reward signal, a graph relation, a physical boundary condition, a confidence gate, or a revision trigger. The synthesis therefore asks whether that documented component remains meaningful when source texts, alternative drafts, reward signals, confidence estimates, domain corpora, and human judgments change, and whether the proposed safeguards lead to selective revision, distillation, abstention, expert review, and continuous assessment in place of merely producing another metric. This point narrows the research task for software tasks represented as evolving graphs: compare alternatives under the same information and resource constraints.

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 deployed pipeline 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 evaluation 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 in place of presenting confidence as a decorative number. For the present focus, lineage-aware learning therefore becomes a criterion for deciding which evidence deserves further scrutiny.

The conceptual framework begins by separating four 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 selective revision, distillation, abstention, expert review, and continuous assessment. 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 lineage-aware learning testable because each claim can be assigned to a component and a measurable transition. The implication is especially consequential in software tasks represented as evolving graphs, where a small modeling choice can redirect attention and resources.

What the Core Studies Establish

SAINF: Intrinsic Self-Correction for Robust Machine Translation with Large Language Models asks how a language model can identify and repair fragile translations without an external quality-estimation model. It uses a conditional workflow that produces a chain-of-dictionary translation, self-scores it, and triggers targeted term interpretation, validation, and revision below a threshold. The relevant evidence is that experiments on ChallengeWMT across language pairs report stronger quality and robustness than prompting and interpretation baselines. The method makes self-correction selective, limiting expensive revision to cases the model judges uncertain. In the present review, this work matters because lineage-aware learning cannot be evaluated as an abstract virtue; it must change how evidence is collected and how an action is withheld. Self-assessment can share the generator's blind spots, and threshold behavior may shift across languages, domains, and model families. (Zhang et al. 2026).

A useful test case is Evo-CuRL: Curriculum-Aware Reinforcement Learning over Code Lineage Graphs for Software Engineering Reasoning. Rather than treating its output as an isolated score, the study frames how software-engineering reasoning can learn from the evolving ancestry of code states rather than isolated prompts and patches. Its central mechanism is curriculum-aware reinforcement learning organized over code-lineage graphs, so training can expose increasingly difficult relationships among revisions and outcomes, and its evidential basis is that the conference record establishes a graph-and-curriculum formulation for software reasoning, with evaluation reported in the software-engineering setting. The lineage representation offers a natural way to retain failed branches, successful repairs, and prerequisite relations during learning. For software tasks represented as evolving graphs, 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 value of a curriculum depends on graph construction, difficulty ordering, and whether benchmark lineages resemble real collaborative repositories. (Tan et al. 2026).

Evidence for lineage-aware learning becomes concrete in One-Step Generative Distillation. The paper addresses how generative feature distillation can avoid the cost and information loss of many denoising steps through MeanFlow Distillation uses an average-velocity-field identity for a direct teacher-to-student transformation and a generative mask loss for adaptive feature weighting. Its evaluation is informative because classification and semantic-segmentation experiments report competitive accuracy with a single generative step. The work recasts distillation as a transport problem that can be both generative and operationally efficient. 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, the evidence is task-specific, and one-step matching may conceal mode loss or teacher bias that aggregate accuracy does not expose. The appropriate use of the result is therefore bounded, comparative, and explicit about unresolved deployment conditions (Chen et al. 2026).

A Traceable System Architecture

Operational design in software tasks represented as evolving graphs 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. For the present focus, lineage-aware learning therefore becomes a criterion for deciding which evidence deserves further scrutiny.

Computational effort should vary with evidential need. 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. This point narrows the research task for software tasks represented as evolving graphs: compare alternatives under the same information and resource constraints.

Testing Claims Rather Than Scores

Evaluation begins by writing each claim in testable form. 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, lineage-aware learning therefore becomes a criterion for deciding which evidence deserves further scrutiny.

Measurement is meaningful only when connected to the decision rule. 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 lineage-aware learning, the minimum report should include overall performance, slice-level results, uncertainty intervals, coverage-risk curves, stability across runs, and failure examples linked back to source texts, alternative drafts, reward signals, confidence estimates, domain corpora, and human judgments. If the deployed pipeline updates incrementally, the assessment 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. In software tasks represented as evolving graphs, this distinction should be tested before it changes an operational decision.

Relevant Foundations

Several established literatures clarify the design space. Domain-adapted language modeling demonstrates why financial vocabulary cannot be treated as generic prose (Araci 2019). Financial text analysis shows that common-language sentiment resources can misread domain-specific terms (Loughran and McDonald 2011). Continued pretraining can improve domain fit but may also narrow coverage if deployment domains remain heterogeneous (Gururangan et al. 2020). Transformer attention provides the architectural basis for long-range contextual modeling in modern language systems (Vaswani et al. 2017). Together these findings imply that lineage-aware learning should be evaluated against explicit alternatives and failure costs. In Code Lineages as Curricula for Language-Model Reasoning, this literature supplies a specific test of lineage-aware learning within software tasks represented as evolving graphs.

The evidence base offers complementary tests rather than one universal metric. Subword modeling addresses rare forms while also making tokenization a consequential design choice (Sennrich, Haddow, and Birch 2016). BLEU made corpus-level comparison scalable but cannot represent every semantic or cultural error (Papineni et al. 2002). Learned translation metrics improve correlation with human judgments while introducing model and domain dependencies of their own (Rei et al. 2020). They also caution against interpreting one benchmark, one split, or one explanatory artifact as a complete validation argument. In *Code Lineages as Curricula for Language-Model Reasoning*, this literature supplies a specific test of lineage-aware learning within software tasks represented as evolving graphs.

A credible assurance case can be built from findings that operate at different levels. Self-consistency uses diverse reasoning samples as evidence, but agreement can still reflect shared bias (Wang et al. 2023). Human-feedback training demonstrates the power and specification sensitivity of learned reward models (Ouyang et al. 2022). Direct preference optimization simplifies alignment training while retaining dependence on preference-data quality (Rafailov et al. 2023). Their combined value lies in making assumptions visible enough to challenge before deployment and to revise after failure. In *Code Lineages as Curricula for Language-Model Reasoning*, this literature supplies a specific test of lineage-aware learning within software tasks represented as evolving graphs.

Experiments the Field Still Needs

The first research priority is failure-mechanism sampling instead of another convenient random split. 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 lineage-aware learning makes that boundary observable and, in principle, auditable.

Beyond individual benchmarks, research must address how findings are retained and updated. 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 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 pipeline. 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. For the present focus, lineage-aware learning therefore becomes a criterion for deciding which evidence deserves further scrutiny.

Deployment Controls

A governance layer translates validation results into permissions and constraints. A field use specification should name allowed uses, prohibited uses, escalation thresholds, data-retention rules, and the person or role that can halt the deployed workflow. 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 predictive component and the surrounding interface. Versioned models without versioned prompts, retrieval indexes, rules, and datasets do not support reliable reconstruction. For the present focus, lineage-aware learning therefore becomes a criterion for deciding which evidence deserves further scrutiny.

Human review is an engineered pipeline, not an automatic safeguard. 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 software tasks represented as evolving graphs, 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 learned component. In software tasks represented as evolving graphs, this distinction should be tested before it changes an operational decision.

Conclusion

Lineage-aware learning is credible only when it is attached to an clearly stated evidence chain and decision policy. Useful mechanisms recur across the literature: structured exploration, typed graphs, conditional revision, adaptive computation, physical constraints, and repeated testing. Their limitations make clear that benchmark scope and implementation scope are different. For software tasks represented as evolving graphs, the practical standard should be a traceable pipeline that measures shift and instability, supports selective revision, distillation, abstention, expert review, and continuous testing, 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 clearly stated deferral and independent verifiability. Seen through lineage-aware learning, the issue is not merely technical; it determines which claims remain open to challenge.

References

1. Zhang, Yin, et al. "SAINF: Intrinsic Self-Correction for Robust Machine Translation with Large Language Models." *Frontiers of Computer Science* (2026).
2. Tan, Wei, et al. "Evo-CuRL: Curriculum-Aware Reinforcement Learning over Code Lineage Graphs for Software Engineering Reasoning." *Proceedings of the 2026 International Conference on Multimedia Retrieval* (2026): 1327-1335.
3. Chen, Yiwei, et al. "One-Step Generative Distillation." *ICASSP 2026 - 2026 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (2026).
4. Araci, Dogu. "FinBERT: Financial Sentiment Analysis with Pre-Trained Language Models." *arXiv preprint arXiv:1908.10063*, 2019.
5. Loughran, Tim, and Bill McDonald. "When Is a Liability Not a Liability? Textual Analysis, Dictionaries, and 10-Ks." *Journal of Finance*, vol. 66, no. 1, 2011, pp. 35-65.
6. Gururangan, Suchin, et al. "Don't Stop Pretraining: Adapt Language Models to Domains and Tasks." *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, 2020, pp. 8342-8360.
7. Vaswani, Ashish, et al. "Attention Is All You Need." *Advances in Neural Information Processing Systems*, vol. 30, 2017.
8. Sennrich, Rico, Barry Haddow, and Alexandra Birch. "Neural Machine Translation of Rare Words with Subword Units." *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics*, 2016, pp. 1715-1725.
9. Papineni, Kishore, et al. "BLEU: A Method for Automatic Evaluation of Machine Translation." *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics*, 2002, pp. 311-318.
10. Rei, Ricardo, et al. "COMET: A Neural Framework for MT Evaluation." *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing*, 2020, pp. 2685-2702.
11. Wang, Xuezhi, et al. "Self-Consistency Improves Chain of Thought Reasoning in Language Models." *International Conference on Learning Representations*, 2023.
12. Ouyang, Long, et al. "Training Language Models to Follow Instructions with Human Feedback." *Advances in Neural Information Processing Systems*, vol. 35, 2022, pp. 27730-27744.
13. Rafailov, Rafael, et al. "Direct Preference Optimization: Your Language Model Is Secretly a Reward Model." *Advances in Neural Information Processing Systems*, vol. 36, 2023.