

EMERGING DIGITAL FORENSICS • CONCEPTUAL PAPER

Against the First Story

A Falsification-First Investigative Methodology
for AI-Augmented Digital Forensics

James Allister Odd, PhD

j.a.odd@ntrlnk.net

Affiliation: NTRLNK, Oregon, USA

ORCID: 0009-0008-2976-6313

DOCUMENT TYPE: Research Paper

PUBLICATION DATE: June 2026

VERSION: 1.0

LICENSE: Creative Commons Attribution 4.0 International (CC BY 4.0)

NTRLNK RELATIONSHIP

This paper was authored by or in association with NTRLNK. The author is responsible for the research methods, analysis, findings, and conclusions presented.

Copyright (c) 2026 James Allister Odd Except where otherwise noted, this work is licensed under the Creative Commons Attribution 4.0 International License. Third-party materials and protected case information are excluded.

Abstract

Digital forensic methods are strong at preserving, extracting, and documenting artifacts, yet the inferential path from artifacts to an account of events remains vulnerable to confirmation bias, premature narrative closure, dependent evidence, and opaque automation. These risks intensify when generative artificial intelligence is used to summarize records, suggest leads, or draft findings: a fluent output can appear corroborative even when it is derived from the same artifacts as the examiner's initial theory. This conceptual paper proposes Falsification-First Digital Investigation (FFDI), an investigative methodology that organizes examination around attempts to disconfirm competing explanations. FFDI adds four linked controls to established acquisition and examination practice: a competing-hypothesis register; an Evidence Dependency Graph that groups observations by causal and transformational origin; a precommitted falsification matrix; and an AI-lead quarantine that prevents machine-generated propositions from being counted as evidence. The methodology uses a staged loop—scope, preserve, model, challenge, discriminate, reconstruct, review, and report—with explicit stop and redirection criteria. It does not replace chain of custody, tool validation, or examiner judgment. Instead, it makes the reasoning between preserved data and reported conclusions more visible, contestable, and reproducible. An evaluation programme is proposed using synthetic cases, seeded ground truth, controlled dependency structures, and comparisons with conventional workflows. Primary outcomes include error rate, calibration, time to correction, dependency-aware corroboration, and inter-examiner reproducibility. FFDI offers a practical new direction for digital forensics in which the quality of an investigation is judged not by how much evidence supports its first story, but by how seriously credible alternatives were tested.

Keywords: digital forensics; investigative methodology; falsification; competing hypotheses; evidence provenance; confirmation bias; generative AI; reproducibility

1. Introduction

Digital forensic investigations rarely fail because no data exist. They more often struggle because data are abundant, distributed, duplicated, transformed by tools, and open to more than one explanation. A login record may indicate intentional access, automated token refresh, credential theft, clock drift, or parser error. A deleted file may reflect concealment, routine retention, synchronization behavior, or an application update. The central investigative problem is therefore not only recovery. It is disciplined inference from incomplete traces to claims about past events.

Established guidance appropriately emphasizes legal authority, preservation, validated methods, working copies, documentation, and reporting [1–4]. Process models have also organized investigations into preparation, collection, examination, analysis, and presentation phases [5–7]. These contributions provide necessary procedural order. However, a phase model does not by itself ensure that an examiner searches fairly for explanations that contradict an early theory. Nor does conventional chain of custody reveal

whether several apparently corroborating findings all originated from one underlying event, one database row, or one automated interpretation.

The arrival of generative AI sharpens this methodological gap. An AI assistant can cluster messages, translate log entries, propose timelines, and identify potential inconsistencies. It can also hallucinate, omit qualifying details, inherit errors from prompts, and restate an examiner's assumption in more persuasive language. If the output is copied into notes and later compared with a tool report derived from the same source, circular agreement may masquerade as independent corroboration. NIST's Generative AI Profile emphasizes governance, measurement, provenance, human oversight, and the management of confabulation and information-integrity risks [8]. Digital forensics needs corresponding case-level controls.

This paper proposes Falsification-First Digital Investigation (FFDI). The method begins with multiple plausible hypotheses and gives priority to observations most capable of showing that a leading hypothesis is wrong. It records how evidence products depend on source artifacts and transformations. It treats AI output as a lead or analytic aid, never as an independent evidential source. Its purpose is not to import a rigid philosophy-of-science formula into casework. Its purpose is operational: to make investigative choices, dependencies, revisions, and uncertainty inspectable.

The contribution is fourfold. First, the paper defines a falsification-first investigative cycle that can sit inside existing legal and laboratory controls. Second, it introduces an Evidence Dependency Graph (EDG) for distinguishing numerical abundance from genuine independence. Third, it defines an AI-lead quarantine for AI-augmented examinations. Fourth, it presents a testable evaluation design without claiming empirical results.

2. The Methodological Gap

2.1 From artifact handling to event inference

ISO/IEC 27037 addresses identification, collection, acquisition, and preservation of potential digital evidence [3]. NIST SP 800-86 describes a practical sequence of collection, examination, analysis, and reporting [1]. SWGDE guidance similarly stresses authorization, integrity, systematic examination, documentation, and communication of conflicting or exculpatory information [2,4]. These controls protect reliability at critical stages, but an additional layer is needed to govern how findings change belief about competing accounts.

Carrier's hypothesis-based approach described digital investigation as testing claims about a system's history through indirect observations [5]. Later process-model research sought standardization and harmonization [6,7]. FFDI extends this lineage by operationalizing three issues together: active disconfirmation, evidence dependence, and AI-mediated reasoning. The extension matters because two

technically accurate extractions can still support a poor conclusion if they are interpreted through an unchallenged theory or incorrectly treated as independent.

2.2 Narrative lock-in

Investigators must form provisional theories to decide where to look. The risk arises when a useful starting theory becomes the organizing story before alternatives have been articulated. Search terms, tool selections, time windows, and interview questions then preferentially recover compatible material. Contradictory artifacts appear peripheral; ambiguous artifacts acquire a one-sided interpretation. This is narrative lock-in.

Contextual bias has been demonstrated in forensic decision-making beyond digital evidence, showing that expectation and task-irrelevant information can affect judgment [9]. Digital investigations add a feedback mechanism: a theory controls the query, the query controls the returned data, and the returned data appears to validate the theory. FFDI does not assume that examiners can become context-free. It instead creates procedural friction against unnoticed commitment.

2.3 False corroboration through dependency

Corroboration is often described informally as agreement among artifacts or tools. Yet independence is a causal property, not a count. Consider a cloud audit event that is exported to JSON, parsed into a timeline, summarized by an AI assistant, and quoted in an examiner's note. Four objects now exist, but all descend from one event record. Agreement among them may verify transformation consistency; it does not constitute four independent observations of the event.

Provenance research has shown the value of recording the tools and transformations that connect acquired data to derived products [10]. FFDI turns that record into an inferential control. It asks which findings share a source, which share a clock, identity provider, parser, or collection path, and which can fail together. Evidence products within the same dependency family are not discarded. They are simply prevented from inflating the apparent weight of corroboration.

2.4 Automation authority

Traditional forensic tools can produce errors through incomplete parsers, unsupported versions, configuration choices, or misunderstood semantics. Generative AI adds variable output, uncertain source attribution, and persuasive language. A model response is a transformation of prompts, retrieved content, system instructions, and model behavior. It may identify a promising lead, but it did not witness the event. Treating the response as evidence confuses a proposition generator with an evidential source.

FFDI therefore separates evidential objects, analytic products, and leads. Evidential objects are acquired or preserved representations whose origin is documented. Analytic products are reproducible transformations such as parsed tables or validated calculations. Leads are propositions that direct further

testing. AI output remains in the third category unless and until a human examiner verifies its factual components against preserved sources using documented methods.

3. Design Principles

FFDI is governed by six principles.

1. **Plurality before preference.** Record multiple reasonable explanations before intensive analysis. Include a benign or process-error alternative where plausible and retain “insufficient information” as a legitimate outcome.
2. **Disconfirmation before accumulation.** Prioritize tests that could distinguish or defeat leading hypotheses, not merely add compatible details.
3. **Independence before counting.** Evaluate common origin and common failure modes before describing findings as corroborative.
4. **Traceability before fluency.** A concise or persuasive analytic output has no special status. Every material claim must trace to preserved data and documented transformations.
5. **Revision without penalty.** Changing a hypothesis in response to evidence is a sign of a functioning method, not examiner failure. Revisions are logged rather than hidden.
6. **Proportionality throughout.** Examination remains bounded by legal authority, necessity, privacy, volatility, cost, and potential harm.

These principles are compatible with conventional preservation and reporting requirements. FFDI begins only after authority and urgent preservation needs are addressed, and it never authorizes access outside the approved scope.

4. Core Investigative Records

4.1 Competing-Hypothesis Register

The Competing-Hypothesis Register (CHR) is opened when the investigative question is defined. A hypothesis must describe a potentially distinguishable state or event, not a person's general character or guilt. For a suspicious account access, suitable entries might be: H1, the account owner intentionally accessed the repository; H2, another person used stolen credentials; H3, an authorized automated service generated the event; H4, the record is a logging or interpretation artifact; and H5, available evidence cannot discriminate among these explanations.

Each entry records scope, assumptions, predicted observations, potential disconfirmers, and revision history. Hypotheses need not be mutually exhaustive, and hybrid explanations may emerge. The register nevertheless prevents the leading explanation from being the only explanation written down.

4.2 Evidence Dependency Graph

The Evidence Dependency Graph (EDG) is a directed graph whose nodes represent source objects, acquisition events, transformations, analytic products, and reported observations. Edges state relationships such as “acquired from,” “parsed from,” “summarized from,” “shares clock with,” “shares identity provider with,” or “manually transcribed from.” Each reported observation must connect backward to a preserved source or be marked as an unverified lead.

The graph supports three practical judgments. First, it identifies evidence families that share a causal source. Second, it reveals shared failure modes—for example, two logs generated by different services but synchronized to the same compromised identity platform. Third, it distinguishes transformation verification from event corroboration. Two parsers agreeing on one database can increase confidence that the database was decoded correctly; it does not independently establish that the recorded real-world event occurred.

4.3 Falsification Matrix

The Falsification Matrix crosses hypotheses with candidate observations. Cells record whether an observation is predicted, compatible, inconsistent, or non-discriminating. The matrix is completed provisionally before high-volume searching whenever practical. This precommitment makes hindsight reinterpretation more visible.

Candidate observation	H1 intentional owner action	H2 credential theft	H3 authorized automation	H4 logging artifact
Interactive session with owner-held device key	Predicted	Usually inconsistent	Inconsistent	Possible if record corrupt
Token used from new infrastructure without device proof	Compatible	Predicted	Compatible	Compatible
Service principal identifier and scheduled-job trace	Inconsistent	Inconsistent	Predicted	Compatible
Independent endpoint telemetry contradicts cloud time	Inconsistent	Compatible	Compatible	Predicted

The matrix does not mechanically decide the case. It helps select high-discrimination tests and exposes observations that are merely compatible with every hypothesis. Where empirical likelihood ratios are unavailable, FFDI uses transparent qualitative judgments rather than invented numerical precision.

4.4 Decision and Revision Log

The Decision and Revision Log records why a source was collected, a test was prioritized, a hypothesis changed, a time window expanded, or an examination stopped. Entries include timestamp, examiner, authority basis, inputs considered, decision, and anticipated consequence. Material reversals remain visible. This creates an audit trail for reasoning rather than only a chronology of tool actions.

4.5 AI-Lead Quarantine

Any proposition generated, ranked, translated, clustered, or summarized by generative AI enters an AI-lead register with the model and interface used, date, relevant settings where available, prompt or task description, source materials supplied, output identifier, and verification status. The proposition cannot populate the “observation” column of the Falsification Matrix until verified against preserved evidence.

The quarantine has five rules:

7. AI output is never an independent corroborating source.
8. A citation supplied by a model is not accepted until the cited material is retrieved and checked.
9. Sensitive evidence is not submitted to an external model without authority, an approved data-handling arrangement, and documented necessity.
10. Material model output is preserved when it influenced investigative action, including enough context to understand the interaction.
11. Final wording must distinguish examiner findings from machine-assisted organization or drafting.

5. The FFDI Workflow

5.1 Stage 1: Scope the question

Translate the request into bounded investigative questions. Record legal authority, relevant systems, custodians, time periods, prohibited areas, urgency, and decision thresholds. Replace broad prompts such as “find evidence of wrongdoing” with event-focused questions such as “which principal initiated the repository export between 08:00 and 10:00 UTC, and what evidence distinguishes interactive action from automated execution?”

5.2 Stage 2: Preserve volatile and foundational sources

Preserve data according to applicable standards and organizational procedures [1–4]. Hashes, acquisition metadata, errors, time sources, tool versions, and chain-of-custody events are recorded. FFDI adds an early dependency assessment: which sources are authoritative, which are replicas, and which may overwrite, synchronize, or derive from others? Volatile independent sources receive priority where lawful and proportionate.

5.3 Stage 3: Model competing explanations

Construct the initial CHR with input from case facts but before detailed exposure to potentially biasing narratives where operationally feasible. At least one reviewer may be given a reduced-context task to evaluate a pivotal artifact. The objective is not artificial blindness to relevant information; it is staged disclosure, so the examiner first interprets the artifact's technical meaning and then integrates wider context.

5.4 Stage 4: Map dependencies

Populate the EDG as sources are acquired and products generated. Assign each observation an evidence-family identifier. A family can be defined by common origin, common transformation, or common failure mode. Cross-family agreement is stronger than within-family repetition, but independence is assessed case by case rather than presumed from different filenames or vendors.

5.5 Stage 5: Precommit discriminating tests

Build the Falsification Matrix and select tests by expected discriminatory value, volatility, intrusiveness, cost, and reversibility. A useful test is one for which plausible outcomes would change the relative standing of hypotheses. Searching a thousand additional messages for the same keyword may have lower value than checking one independently administered authentication factor.

For each selected test, record the predicted consequence for each hypothesis before observing the result. If precommitment is impossible because the artifact was already seen, record that fact. FFDI values candor over a fictional clean sequence.

5.6 Stage 6: Execute, verify, and challenge

Perform examination on protected working copies using validated or tested methods. Confirm pivotal results by an alternative method when feasible. Alternative methods may test a transformation, but the EDG must still show whether they rely on the same source. Actively search for artifacts that would be expected if the leading hypothesis were false. Record contradictory, exculpatory, and null results with the same care as supporting results.

5.7 Stage 7: Reconstruct with uncertainty

Construct the event account from verified observations, maintaining separate fields for source time, normalized time, uncertainty interval, actor attribution, and dependency family. Explain alternative sequences that remain compatible with the evidence. Avoid converting absence of an artifact into evidence of absence unless collection completeness, retention, logging state, and expected artifact generation have been established.

5.8 Stage 8: Review and report

A reviewer tests the strongest conclusion by selecting at least one plausible rival hypothesis and asking what evidence most threatens the reported account. The final report states the question, authority and scope, items examined, methods and limitations, material observations, dependency relationships, alternative hypotheses considered, disconfirming tests, unresolved contradictions, and conclusion strength. SWGDE's reporting requirements provide a baseline for minimum report elements [4]; FFDI adds a concise inferential audit trail.

6. Stop, Redirect, and Escalation Criteria

Investigations can expand indefinitely unless stopping is principled. FFDI defines four stopping conditions. A sufficiency stop occurs when the decision-maker's question is answered to the required standard and remaining tests are unlikely to alter the conclusion. An insufficiency stop occurs when accessible evidence cannot discriminate among material hypotheses. A proportionality stop occurs when the likely value of further collection does not justify privacy intrusion, cost, delay, or operational harm. An authority stop occurs when the next useful test falls outside legal or policy authorization.

Redirection is required when a pivotal observation is invalidated, a new hypothesis explains more of the evidence with fewer unsupported assumptions, or the EDG reveals that supposed corroboration collapses into one dependency family. Escalation is required for scope expansion, destructive testing, privileged or highly sensitive data, uncertain legal authority, or deployment of an unvalidated technique on a dispositive issue.

7. Worked Example: Suspicious Repository Export

Assume an organization discovers a large source-code export attributed to an employee account shortly before resignation. A conventional narrative may quickly become “the employee stole the code.” FFDI first frames the event question and preserves identity-provider logs, repository audit records, endpoint telemetry, automation records, network data, and relevant administrative configuration.

The CHR records intentional employee action, stolen credentials, authorized automation, administrator action misattributed to the employee, logging error, and insufficiency. The EDG shows that the security dashboard, CSV export, and AI-generated incident summary all derive from the same repository API event. They are one evidence family. Identity-provider logs are nominally separate, but both systems may rely on the same account identifier and time synchronization service; those shared dependencies are recorded.

The matrix identifies high-value disconfirmers: hardware-backed device authentication, an interactive endpoint process tree, a service-principal token, scheduler traces, repository API semantics, and evidence that the account mapping changed. The examiner predicts how each outcome would affect each hypothesis before retrieving it. An AI assistant suggests that the export volume resembles a known exfiltration pattern. The claim remains quarantined until the underlying calculation, comparison set, and artifact selection are independently verified.

Suppose endpoint telemetry shows no interactive repository client, while automation logs reveal a scheduled backup using a token labeled with the employee's account due to a legacy configuration. Repository and identity logs are accurate, yet the first story is wrong. FFDI records the revision, verifies the automation path, tests whether the backup destination was authorized, and reports the remaining

uncertainty. The example illustrates the method's aim: accurate artifacts do not guarantee accurate attribution.

8. Evaluation Programme

FFDI should be evaluated before claims of improved performance are made. A suitable study would compare trained examiners using a conventional laboratory workflow with examiners using FFDI on matched synthetic cases. Cases should contain seeded ground truth, plausible alternative explanations, realistic missingness, contradictory timestamps, duplicated derivative products, parser error, and AI-generated leads of mixed quality.

8.1 Research questions

12. Does FFDI reduce false attribution and premature closure without materially increasing missed true events?
13. Do examiners distinguish independent corroboration from derivative repetition more accurately?
14. Does AI-lead quarantine reduce the incorporation of unsupported model claims?
15. Does FFDI improve calibration between expressed confidence and ground-truth accuracy?
16. Can another examiner reproduce the reported reasoning and identify why hypotheses changed?

8.2 Measures

Primary measures should include conclusion accuracy, false-attribution rate, sensitivity to exculpatory artifacts, time to first correct revision, and calibration scores. Process measures should include number and diversity of hypotheses, proportion of pivotal observations assigned to a dependency family, proportion of AI leads verified before use, and completeness of decision logs. Reproducibility should be measured by whether an independent reviewer can reconstruct material decisions from the case record, not merely reproduce a tool's output.

8.3 Experimental safeguards

Scenarios, answer keys, and scoring rules should be preregistered. Examiners should be stratified by experience and randomly assigned where feasible. The study should separate training effects from methodology effects and report workload, time, and usability. Because documentation can improve simply through increased effort, the control workflow should receive an equivalent time budget. Human-participant research requires institutional ethics review, informed consent, appropriate handling of performance data, and protection against employment consequences.

8.4 Success and failure criteria

FFDI would be promising if it reduces serious attribution errors, improves calibration and dependency recognition, and preserves acceptable timeliness. It would require revision if paperwork grows without measurable inferential benefit, if examiners generate cosmetic alternatives that do not affect testing, or if the method suppresses timely action in volatile incidents. Negative results should be published because they identify which controls are burdensome or ineffective.

9. Implementation and Governance

Organizations can pilot FFDI without replacing their forensic platforms. The four records may initially be implemented as controlled case forms linked to evidence identifiers. Tool integrations should favor open export formats, stable identifiers, versioned transformations, and append-only revision history. Automated assistance may propose dependency edges or candidate disconfirmers, but a human examiner must approve material classifications.

Training should use cases in which the obvious story is sometimes correct and sometimes wrong. If every exercise rewards reversal, examiners will learn a new bias toward contrarian conclusions. Assessment should reward appropriate uncertainty, accurate dependency mapping, proportionate testing, and documented revision—not the number of hypotheses or the complexity of the graph.

Quality assurance should sample closed cases for four questions: Were reasonable alternatives recorded? Were pivotal tests genuinely discriminating? Was independence assessed causally? Did any AI-derived proposition enter the report without source verification? Audit findings should improve procedures and training rather than become quotas that encourage superficial compliance.

FFDI records may contain sensitive hypotheses about individuals who are later excluded. Access controls and retention schedules are therefore essential. Reports should disclose relevant alternatives and limitations without unnecessarily circulating speculative personal information. The CHR is an investigative control, not a license to memorialize unsupported allegations indefinitely.

10. Discussion

FFDI shifts the unit of quality from artifact volume to inferential resilience. A conclusion is stronger when it survives tests designed to defeat it, when its observations come from appropriately independent sources, and when the route from source to statement can be reproduced. This orientation also changes how AI is used. The model can help generate questions, organize large corpora, or identify apparent contradictions, but the methodology prevents its fluency from becoming evidential authority.

The method does not eliminate judgment. Classifying dependencies, selecting reasonable hypotheses, and deciding sufficiency require expertise. FFDI makes those judgments explicit enough to challenge. It also

avoids a simplistic demand that every hypothesis be disproved. In open systems, complete falsification may be impossible. The practical goal is to seek the most informative available tests and report what remains unresolved.

An EDG may also improve cross-disciplinary communication. Lawyers can see which conclusions depend on contested sources; incident responders can identify shared telemetry failure modes; laboratories can locate transformations requiring validation; and researchers can study how evidence structures affect error. Standardized schemas for the CHR, EDG, and AI-lead register could support tool interchange, but premature standardization should be avoided until field evaluations reveal which fields are consistently useful.

11. Limitations

This is a conceptual proposal. No empirical comparison has yet established that FFDI improves outcomes. The method may impose disproportionate overhead on simple examinations and may be difficult under severe time pressure. Hypothesis generation can still omit the correct explanation. Dependency graphs can become complex, and causal independence is sometimes uncertain. Blinded or staged review may be infeasible in small teams. Qualitative matrix ratings can conceal disagreement unless reasons are recorded.

The worked example simplifies identity, logging, and organizational context. Legal disclosure obligations and admissibility rules vary by jurisdiction. FFDI must therefore be adapted through competent legal advice, laboratory policy, and validation. Its AI controls address inferential use, not every cybersecurity, confidentiality, intellectual-property, or vendor risk associated with model deployment.

12. Conclusion

Emerging digital forensics requires more than faster extraction and larger analytic models. It requires methods that remain reliable when evidence is abundant, derivative, ambiguous, and mediated by automation. Falsification-First Digital Investigation offers a practical response: state credible alternatives, map evidence dependencies, precommit discriminating tests, quarantine AI-generated leads, log revisions, and stop with explicit uncertainty when the data cannot decide.

The method's central claim is deliberately testable. Investigations structured to challenge their first story should produce fewer confident errors and clearer audit trails than investigations structured mainly to accumulate support. Controlled evaluation, practitioner piloting, and adversarial review are the necessary next steps. Until those results exist, FFDI should be regarded as a research agenda and pilot methodology—not a validated replacement for established forensic standards.

References

17. Kent, K., Chevalier, S., Grance, T., & Dang, H. (2006). *Guide to Integrating Forensic Techniques into Incident Response* (NIST SP 800-86). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.SP.800-86>
18. Scientific Working Group on Digital Evidence. (2025). *Best Practices for Computer Forensic Examinations* (SWGDE 18-F-001-2.0). <https://www.swgde.org/documents/published-complete-listing/18-f-001-swgde-best-practices-for-computer-forensic-examination/>
19. International Organization for Standardization. (2012). *ISO/IEC 27037:2012: Guidelines for identification, collection, acquisition and preservation of digital evidence*. <https://www.iso.org/standard/44381.html>
20. Scientific Working Group on Digital Evidence. (2018). *Requirements for Report Writing in Digital and Multimedia Forensics* (SWGDE 18-Q-002). <https://www.swgde.org/documents/published-complete-listing/18-q-002-swgde-requirements-for-report-writing-in-digital-and-multimedia-forensics/>
21. Carrier, B. D. (2006). *A Hypothesis-Based Approach to Digital Forensic Investigations* [Doctoral dissertation, Purdue University]. https://www.cerias.purdue.edu/apps/reports_and_papers/view/2920
22. Kohn, M. D., Eloff, M. M., & Eloff, J. H. P. (2013). Integrated digital forensic process model. *Computers & Security*, 38, 103–115. <https://doi.org/10.1016/j.cose.2013.05.001>
23. Valjarevic, A., & Venter, H. S. (2015). A comprehensive and harmonized digital forensic investigation process model. *Journal of Forensic Sciences*, 60(6), 1467–1483. <https://doi.org/10.1111/1556-4029.12823>
24. Autio, C., et al. (2024). *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile* (NIST AI 600-1). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.600-1>
25. Dror, I. E., Charlton, D., & Péron, A. E. (2006). Contextual information renders experts vulnerable to making erroneous identifications. *Forensic Science International*, 156(1), 74–78. <https://doi.org/10.1016/j.forsciint.2005.10.017>
26. Garfinkel, S. L., Nelson, A. J., White, D., & Roussev, V. (2009). DEX: Digital evidence provenance supporting reproducibility and comparison. *Digital Investigation*, 6(Supplement), S48–S56. <https://doi.org/10.1016/j.diin.2009.06.011>
27. Heuer, R. J., Jr. (1999). *Psychology of Intelligence Analysis*. Center for the Study of Intelligence, Central Intelligence Agency. <https://www.cia.gov/resources/csi/static/Psychology-of-Intelligence-Analysis.pdf>
28. Montasari, R. (2017). A standardised data acquisition process model for digital forensic investigations. *International Journal of Information and Computer Security*, 9(3), 229–249. <https://doi.org/10.1504/IJICS.2017.085139>

Declarations

Funding: This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Conflicts of interest: The author declares no conflict of interest.

Data availability: No empirical data were generated for this conceptual paper. The proposed evaluation materials should be released with any subsequent experimental study, subject to ethical, legal, and security constraints.

Ethics statement: No human participants or personal data were used in preparing this conceptual paper. The proposed human-participant evaluation requires appropriate institutional ethics review before recruitment.

Use of generative AI: The author used generative artificial intelligence as an assistive writing tool to help address difficulties associated with dyslexia, including support with spelling, grammar, sentence structure, organization, and clarity. The author developed the paper's arguments, reviewed and revised the generated text, checked the cited sources, and takes full responsibility for the accuracy, originality, analysis, and final content. Generative AI was not used to generate or analyze empirical data. This disclosure should be adapted to the target journal's policy.