

# Foundations and Positioning of the Zemi Method

---

**Companion to:** The Zemi Method, Version 1.1

**Version:** 1.3

**Status:** Published. Prepared to accompany the Version 1.1 method and worked example. Principal sources from Version 1.2 remain verified; Version 1.3 records pressure-test-driven changes adopted in Version 1.1.

**Versioning note:** This companion is versioned independently from the method because it records changes in positioning, literature review, acknowledgements, and pressure-test history.

**First published:** July 17, 2026

**This version published:** July 25, 2026

**Canonical home:** <https://zemimethod.com>

**Copyright:** © 2026 Kevin V. Watson. All rights reserved.

**Author:** Kevin V. Watson

**Web:** <https://kevinvwatson.com>

**Contact:** [kevinvwatson@gmail.com](mailto:kevinvwatson@gmail.com)

**Material changes from v1.2:** Companion aligned to the Version 1.1 method. Added a section documenting the Version 1.1 pressure-test cycle and the controls it produced: Authority and Scope, disclosure authority, transfer integrity and retained-custody boundaries, source completeness, prior-work treatment, suspect/source narrowing, layered-infrastructure mapping, known-finding support basis, standards and benchmark mapping, pending and exhausted undetermined findings, AI role classification, tool-derived measurement limits, derivative-presentation controls, and artifact-conformance requirements. The prior Version 1.0 authority-and-scope limitation is updated to reflect its resolution in Version 1.1.

## Contents

|                                                                   |    |
|-------------------------------------------------------------------|----|
| Purpose of this document .....                                    | 4  |
| Version 1.1 pressure-test cycle .....                             | 4  |
| The traditions the method relies on .....                         | 6  |
| Contemporaneous and parallel work .....                           | 7  |
| Recent digital forensic evaluation and reporting literature ..... | 9  |
| What the method inherits .....                                    | 11 |
| What the method contributes .....                                 | 11 |
| What is claimed, and what is not .....                            | 12 |
| The gaps the method is designed to address .....                  | 12 |
| References .....                                                  | 13 |

## Purpose of this document

The Zemi Method is a compact, normative doctrine. It states what the practice requires and stops there. This companion does the work the doctrine deliberately does not: it identifies the traditions the method relies on, acknowledges its predecessors, and states precisely what the method claims and what it does not.

The method was not invented from nothing, and this document does not pretend otherwise. Most durable methodologies are built through synthesis, extension, and formalization of ideas that earned their standing elsewhere. The honest position, and the defensible one, is to name those ideas, show where they came from, and be exact about what the synthesis contributes.

The narrow claim is that the Zemi Method integrates claim-testing, bounded classification, human adjudication, automation control, and AI assurance into one practitioner-facing evidentiary discipline for decisions. It does not claim to originate the underlying reasoning traditions or replace specialized forensic, attribution, evaluative-reporting, or legal frameworks.

## Version 1.1 pressure-test cycle

Version 1.1 is a pressure-test revision rather than a change in the method's premise. The premise remains that every claim relevant to a decision is a hypothesis until independent evidence supports it. What changed is the set of controls required before a finding may be treated as defensible.

The revision was developed by applying the method against public incident reports, a public legal complaint, and a published video-forensic case study. Those tests were not validation studies and are not claimed as empirical proof. They were structured stress tests: each case was used to ask whether a competent practitioner could do less work, hide a material uncertainty, or overstate a claim while still appearing to comply with the method.

The tests exposed several repeatable gaps. First, Version 1.0 checked claims across independent sources but did not expressly require the practitioner to test whether the extraction from a single source was complete and faithful for the question being asked. Second, authority was implicit in engagement practice but not a named phase, and authority to perform work was not separated from authority to disclose findings or evidence. Third, cross-border and remote-acquisition scenarios showed that authority to release evidence is not the same as proof that transferred evidence remained intact, securely packaged, and received by the authorized recipient. Fourth, onsite-only and retained-custody scenarios showed the inverse problem: where the evidentiary claim is that devices, data, temporary media, or workspaces never left a defined boundary, non-transfer must be recorded rather than merely asserted. Fifth, second-opinion and litigation-review scenarios showed that prior reports, prior examiner findings, and investigative summaries need to be treated as source material rather than inherited as established findings. Sixth, anonymous-account and pattern-analysis scenarios showed that suspect, actor, account, device, system, or source narrowing can be mistaken for attribution unless the report states the remaining candidate pool and the evidence still required for stronger attribution. Seventh, CDN, cloud, domain, and financial-record scenarios showed that

technical delivery, account ownership, operational control, legal responsibility, and financial benefit can sit at different layers and must not be collapsed into a single responsibility finding. Eighth, standards-based defence work showed that a statement of compliance, alignment, or deviation can overreach unless the report maps the finding to a specific standard, version, requirement, applicability basis, evidence, and interpretive limit. Ninth, undetermined findings did not distinguish matters still being pursued from matters exhausted within the available authority and proportionate scope. Tenth, known findings identified epistemic status but did not require the report to state the support basis that made the finding known. Eleventh, the AI boundary addressed automation as assistance but did not require role classification when an AI system was material to the engagement. Twelfth, public video-forensic material showed that tool-derived measurements, enhancements, comparisons, calculations, and transformations require their own stated processing basis, validation checks, and uncertainty limits. Thirteenth, CCTV and demonstrative-evidence scenarios showed that excerpts, compilations, edited videos, timelines, transcripts, visualizations, and other derivative presentations can be mistaken for source evidence unless their source links, selection criteria, excluded context, and non-substitution limits are recorded.

Version 1.1 responds by adding an Authority and Scope phase before Frame; separating work authority from disclosure authority; requiring transfer integrity records where evidence or forensic copies are packaged, encrypted, transferred, or delivered; requiring retained-custody records where evidence, devices, temporary media, workspaces, or data must remain within a defined boundary; requiring source-completeness checks for material extractions; requiring prior reports, prior examiner findings, and investigative summaries to be tested against their underlying evidence before being relied on; requiring suspect, actor, account, device, system, or source narrowing to be distinguished from attribution where unique responsibility is not established; requiring layered infrastructure and financial records to distinguish technical delivery, account ownership, operational control, legal responsibility, and financial benefit where material; requiring standards, frameworks, policies, and benchmarks to be mapped to the specific requirement and evidence where a finding depends on them; requiring known findings to state their support basis; requiring undetermined findings to be marked pending or exhausted; requiring AI role classification where AI is material; requiring tool-derived measurement or transformation limits to be stated when a finding depends on them; and requiring derivative presentations to identify the source evidence, selection or editing criteria, source locators, material excluded context, and non-substitution limits.

These additions do not broaden the method's originality claim. They make the doctrine more demanding of its own practitioners. A Version 1.1-conformant report must now show not only that a claim was tested against independent evidence, but that the practitioner was authorized to do and release the work, that material extracts did not silently omit relevant records, that automation and AI did not cross the adjudication line, that tool-derived outputs were not treated as self-explaining evidence, and that derivative presentations were not treated as substitutes for the source record.

The same consequence applies to the method's supporting materials. Materials prepared under Version 1.0 do not automatically satisfy Version 1.1. A worked example, report skeleton,

checklist, visual summary, or template should not be described as Version 1.1-conformant unless it includes the Authority and Scope phase, identifies the mode of work, justifies proportionality and scope limits, records source-completeness limits, separates work authority from disclosure authority, records transfer integrity where evidence or forensic copies are transferred, records retained-custody boundaries where material is deliberately kept in place, treats prior work as source material where material, distinguishes narrowing from attribution where material, separates layered infrastructure responsibilities where material, maps standards or benchmarks where material, states support basis for known findings, states pending or exhausted status for undetermined findings, records AI role or tool-derived measurement limits where material, and records derivative-presentation controls where material.

## The traditions the method relies on

**The epistemology of testing.** The method's premise, that every claim relevant to a decision is a hypothesis until independent evidence supports it, descends from the falsificationist account of scientific reasoning associated with Karl Popper. The requirement in the Frame phase that a working theory state its own breaking condition is that account applied to casework. The method makes no claim to this epistemology. It claims only its consistent operational use.

**The anatomy of argument.** Stephen Toulmin's model of argumentation separates a claim from its supporting data, from the warrant connecting them, and from the qualifiers bounding them. The method's requirement that reports state the reasoning connecting evidence to conclusion, and its classification of findings as known, assumed, or undetermined, are working translations of that anatomy into report structure.

**The mapping of evidence.** John Henry Wigmore's chart method, developed for legal evidence a century ago and carried forward in the modern evidence scholarship of Anderson, Schum, and Twining, established that every inferential step between evidence and conclusion can and should be made explicit and inspectable. The method's traceability requirement and its defensibility check are a lightweight, practice-ready expression of that tradition.

**The competition of hypotheses.** Richards Heuer's Analysis of Competing Hypotheses, developed for intelligence analysis, requires that rival explanations be tested against the same evidence together, with disconfirmation privileged over confirmation. The Frame phase adopts this discipline directly: where the evidence admits more than one explanation, the surviving explanation is the one the evidence failed to break, not the one considered first.

**The calculus of uncertainty.** Bayesian evaluative reporting, reflected in European forensic practice through likelihood-ratio approaches and the ENFSI guideline and primer on evaluative reporting, represents the most developed probabilistic treatment of evidential strength. The method takes a different route for its reporting layer, categorical classification rather than probabilistic calculus, because categorical boundaries are auditable and communicable to decision-makers and triers of fact. Nothing in the method forbids probabilistic evaluation within

analysis where the evidence type supports it. The choice concerns how conclusions are communicated and owned, not how strength of evidence may be assessed.

**The precedent for graded vocabularies.** The GRADE framework in evidence-based medicine demonstrated that a spare, categorical grading language, published, versioned, and consistently applied, can become shared infrastructure for an entire field. The method's version discipline is modeled on that precedent deliberately.

**The forensic process standards.** The established process models and consensus documents of the field, including NIST SP 800-86, ISO/IEC 27037, ISO/IEC 27043, and SWGDE best-practice guidance for digital evidence collection, govern the handling, acquisition, and examination of digital evidence. The method does not replace or compete with them. It governs the reasoning above them, and it expects engagements to comply with the applicable standards beneath it.

**The tradition of investigative distrust.** The nearest published predecessor to the method's premise is Zero Trust Digital Forensics, proposed by Neale, Kennedy, Price, Yu, and Nuseibeh in 2022, which defined the strategy as one where each aspect of an investigation is treated as unreliable until verified, and introduced multifaceted verification of digital artefacts as its operating principle. The kinship is substantial and is acknowledged without reservation: both positions reject unexamined trust in investigative components, and both respond with verification. That paper also modeled the honest posture this document adopts, noting openly that the practices implementing its strategy were not all fundamentally new. The authors presented their strategy as theoretical, with its practical application deliberately left to future work, and the research program built on it has remained centered on the detection of artefact tampering. The Zemi Method stands in this tradition and extends it in kind as well as scope: from a verification strategy at the artifact level to an operating doctrine at the decision level, with a lifecycle, an adjudication gate, an automation boundary, and an assurance application the earlier work did not contemplate.

**The emerging consensus on human accountability for automated assistance.** Across fields, a common position is forming: automated systems may assist evidence work, humans remain accountable for judgments, and consequential automated contributions must be disclosed. In evidence synthesis, the 2025 joint position of Cochrane, the Campbell Collaboration, JBI, and the Collaboration for Environmental Evidence requires human oversight and transparent reporting wherever automation makes or suggests judgments. In forensic practice, current vendor guidance holds that AI changes the review method while the evidentiary process and its human validation remain unchanged. The method's fourth principle, AI-assisted, never AI-decided, belongs to this emerging consensus and claims no priority over it. What the method adds is procedural form: a named adjudication phase that automation is barred from crossing, and a reporting requirement that keeps the automated contribution separable from human judgment.

## Contemporaneous and parallel work

The traditions above are sources and antecedents. This section records work that is parallel rather than developmental: published before this companion revision, developed independently,

and arriving at overlapping positions. It appears separately because the method does not derive from it, and describing it as a source for the method's development would be inaccurate.

**The FACT Attribution Framework.** In December 2025, Brett Shavers published the FACT Attribution Framework, a jurisprudence-centered attribution methodology for digital forensics and incident response, structured around four stages: Forensic Authority and Compliance, Analyze Evidence, Correlate and Sequence, and Testify and Transfer. FACT was released following pre-release review by a panel of practitioners drawn from law enforcement, corporate practice, consulting, academia, and the legal profession.

FACT was published before the Zemi Method, Version 1.0, but the Zemi Method was developed and published without knowledge of FACT. The author first encountered FACT in July 2026, after publishing Version 1.0 on July 17, 2026, and records it here at the first opportunity. The convergence between the two is substantial and is stated plainly rather than minimized.

Both works are published, versioned, DOI-anchored practitioner doctrines that position themselves above the established process standards of the field rather than in competition with them. Both require competing explanations to be framed and tested rather than assumed away. Both treat falsification as a duty rather than an option. Both hold that evidence drawn from a shared mechanism does not constitute independent corroboration, a point FACT expresses as independence being a matter of distinct causal origins rather than of format or location. Both favor categorical expression of conclusions over formal probabilistic likelihood ratios in ordinary practice, though they explain that choice from different emphases: FACT from the absence of robust empirical error rates across most artifact types, and the Zemi Method from auditability, communicability, and human ownership of conclusions. Both require that automated contributions be disclosed and independently validated against source evidence rather than relied upon directly. And both treat the refusal to conclude as a legitimate and sometimes required outcome, expressed in FACT as a non-attribution finding and in the Zemi Method as classification of a matter as undetermined.

That two practitioners working separately arrived at these positions within months of one another is best explained by common descent rather than by influence in either direction. Independence of sources, falsification, competing hypotheses, and categorical expression of uncertainty are longstanding elements of forensic and scientific reasoning, and both works cite that ancestry openly. The convergence is evidence that the underlying tradition is sound and that the questions it now faces, particularly the question of automation, are being asked across the field at the same time.

The two methods address different subjects. FACT answers who acted: it is an attribution overlay, bounded to person-level conclusions, built on identity layers that separate artifact from device from account from session from person, and grounded in a stage devoted to lawful authority, scope, and proportionality. The Zemi Method addresses whether a claim relevant to a decision can be trusted, of which attribution is one instance among several. Its second subject, the examination of an AI system's own claims under the same evidentiary standard, is not within FACT's scope, which treats automation as a tool to be disclosed and as an alternative actor to be considered rather than as a matter under examination.

The relationship is therefore complementary. An engagement may reasonably apply the Zemi Method to establish whether the underlying claims hold and FACT to govern whether those established findings support a person-level attribution. Practitioners working on attribution questions are directed to FACT on its own terms.

## Recent digital forensic evaluation and reporting literature

The FACT comparison prompted a wider pass through recent digital forensic scholarship. That pass exposed a weakness in the earlier companion: it treated classical reasoning traditions more fully than recent work in Forensic Science International: Digital Investigation, Science & Justice, DFRWS, and adjacent digital forensic evaluation literature. This section records that correction. It is not exhaustive, and it does not turn this companion into a literature review. Its purpose is to identify the works closest to the method's reasoning, reporting, automation, and assurance boundaries.

**Decision support and reliability of findings.** Horsman's Digital Evidence Reporting and Decision Support framework, published in 2019, is a close predecessor to the method's concern with when a practitioner may safely report an inference, assumption, or conclusion. DERDS formalizes investigative decision-making and asks whether identified digital data is reliable enough to disclose. The Zemi Method does not claim priority over that concern. Its contribution is different in form: a broader claim-to-finding lifecycle that begins with framing the question, requires independent corroboration where available, classifies findings at a human adjudication gate, and applies the same discipline beyond criminal case reporting to assurance questions.

**Digital Evidence Certainty Descriptors and their critique.** Horsman's Digital Evidence Certainty Descriptors proposed a harmonized verbal approach to expressing uncertainty in digital forensic findings, arguing that fine-grained quantification of uncertainty in digital evidence may not yet be achievable. That work is directly adjacent to the method's use of known, assumed, and undetermined. It must therefore be acknowledged plainly.

The critique by Biedermann and Kotsoglou is equally important. They argue that digital evidence should not be treated as exceptional in a way that bypasses lessons already learned in forensic science, and they criticize verbal descriptor systems for conceptual ambiguity and legal interpretive risk. The Zemi Method accepts the force of that warning. Its three classifications are not intended as a verbal scale of evidential strength, nor as substitutes for probability or likelihood ratios. Known, assumed, and undetermined classify the epistemic status of a proposition in the report: supported within stated limits, relied upon without direct proof for a stated purpose, or unresolved on the available evidence. They are categories of kind, not degrees of confidence. A known finding may still be narrow, weakly supported, or contingent on stated limits. The classification says that support exists and has been bounded; it does not say how strong the evidence is in probabilistic terms.

**Preliminary evaluative opinions and the investigative/evaluative distinction.** Casey's work on standardizing preliminary evaluative opinions on digital evidence is a significant predecessor in the expression of digital forensic conclusions. It emphasizes competing hypotheses, strength

of evidence, avoidance of transposed conditional reasoning, and clarity for decision-makers. The Zemi Method shares the concern for balanced reasoning and understandable expression, but it does not present itself as a likelihood-ratio or strength-of-evidence framework. Version 1.1 responds to this literature by requiring the report to identify the mode of work as investigative, preliminary evaluative, or fully evaluative, and to justify that mode by the purpose of the engagement. That is a boundary control, not a claim that the method supplies a full evaluative-reporting calculus where one is required.

**Knowledge bases, weaknesses, and mitigations.** SOLVE-IT, proposed by Hargreaves, van Beek, and Casey in 2025, is a digital forensic knowledge base inspired by MITRE ATT&CK. It catalogs techniques, objectives, weaknesses, mitigations, examples, and references. SOLVE-IT is not a competing methodology to Zemi. It operates at the technique and knowledge-base layer, whereas the Zemi Method operates at the reasoning and reporting layer. In practice, SOLVE-IT could strengthen the Corroborate, Analyze, and Defensibility Check phases by helping practitioners identify technique-level weaknesses and mitigations that should be accounted for before a finding is released.

**Cognitive bias, reliability, and reporting practice.** Sunde and Dror's work on reliability and biasability in digital forensics gives empirical support to the need for structured framing, controlled exposure to contextual information, and explicit reasoning. In their study, 53 digital forensic examiners analyzed the same evidence file. The results showed low reliability across observations of traces, interpretations of observed traces, and conclusions, and showed that observations were affected by biasing contextual information. Sunde's later comparative study of digital forensic reporting practice found inconsistent conclusion types, varied certainty expressions, and insufficient reference to established frameworks. These studies do not merely support the method; they also challenge it. A classification system is only useful if different practitioners can apply it consistently, and if readers understand what the classifications mean. Inter-rater reliability and reader comprehension therefore remain validation questions for the Zemi Method.

**Peer review and scrutiny.** Horsman and Sunde's work on peer review in digital forensics is directly relevant to the method's publication posture. They describe peer review as a primary quality assurance mechanism in digital forensic casework, while also warning that the field lacks consistent, standardized approaches to what peer review should examine and how deeply it should operate. That literature frames the method's own review gap more precisely. Public release on Zenodo and GitHub made the method dated, citable, and open to scrutiny. It did not make the method peer reviewed. Formal peer review, structured expert review, and documented responses to critique remain separate validation tasks.

**Automation and AI in forensic reasoning.** Bollé, Casey, and Jacquet's work on automated systems supporting forensic analysis is the closest identified predecessor to the method's automation and assurance boundary. Their evaluation hierarchy of performance, understandability, and forensic evaluation anticipates several concerns in the method's fourth principle: automated outputs must be critically evaluated, explainable enough to support a forensic conclusion, and fit for the specific question. The 2026 survey of large language models

in digital forensics reinforces the same point in the specific context of generative AI: probabilistic model outputs create tension with deterministic forensic requirements, and explainability, reproducibility, validation, and admissibility remain central concerns. The Zemi Method's AI-assisted, never AI-decided rule stands within this developing literature, not outside it. Version 1.1 adds a narrower reporting control: where an AI system is material, the report classifies the AI role as AI present, AI-assisted, AI-orchestrated, or AI-autonomous claimed, based on evidence rather than labels, prompts, logs, or vendor descriptions alone.

**Attribution and identity-layer literature.** FACT's own literature review surveys the attribution field and points to a body of work on source, activity, account, device, actor, and legal implication. The Zemi Method does not reproduce that survey here, because attribution is not this method's subject. Zemi may be used in attribution work only as a claim-testing discipline. It does not replace attribution-specific frameworks, legal analysis, or identity-layer models. Practitioners working on attribution questions should read FACT and the attribution literature it identifies on their own terms.

The effect of this literature pass is narrowing, not abandonment. The method no longer stands as though recent digital forensic scholarship were silent on decision support, uncertainty, reporting, attribution, automation, or AI. It stands as a synthesis that must be read beside that work.

## What the method inherits

From these traditions the method inherits nearly all of its raw material: hypothesis testing, stated breaking conditions, explicit inference, competing explanations, bounded conclusions, verification before trust, chain of custody, and human ownership of judgment. None of these is claimed as original, and a reader who recognizes their sources is reading the method correctly.

## What the method contributes

The contribution is integration, structure, and extension.

Integration: the elements above are combined into a single doctrine with one epistemology, rather than existing as separate practices a skilled examiner may or may not apply.

Structure: the doctrine is procedural. It has a lifecycle with a mandatory human adjudication gate, a categorical classification applied at that gate, a standing defensibility check, and a versioned, dated public form that an engagement can be held to. In Version 1.1, that structure becomes stricter: authority and scope are checked before framing, source completeness is tested before extraction results are trusted, disclosure authority is separated from work authority, known findings state their support basis, and undetermined findings state whether they are pending or exhausted.

Extension: the same evidentiary standard is applied to a second subject. On the assurance side, the AI system itself becomes the matter under examination, its claims corroborated, its controls evidenced, its accountability located. The symmetry between investigating with AI assistance

and investigating AI systems under one claim-to-finding discipline remains the method's most distinctive extension, but that claim is now stated narrowly. Prior work has addressed automated-system evaluation, AI use in digital forensics, and forensic knowledge bases.

## **What is claimed, and what is not**

**Claimed:** the Zemi Method integrates established principles of forensic verification, evidentiary reasoning, bounded conclusions, and human accountability into a single, versioned methodology governing both digital investigation and AI assurance. The integration and structure are not claimed as unique. Contemporaneous work has arrived at comparable forms, and the elements combined are, individually, the common property of the field. What is claimed is the particular combination, the decision-level rather than attribution-level framing, and the extension of the same claim-to-finding discipline to AI assurance.

**Not claimed:** that the method invented independent verification, human accountability for automated assistance, falsifiable framing, or any of its component principles. Those claims would be false, and the method's own third principle forbids making them.

**Resolved in Version 1.1:** the published Version 1.0 doctrine did not contain an authority and scope stage equivalent to FACT's first stage. The Frame phase defined the question to be answered, but did not itself require a separate determination that the examiner was legally, contractually, ethically, and professionally permitted to ask that question, obtain the evidence needed to answer it, or release the resulting findings. Version 1.1 addresses that limitation by adding Authority and Scope as Phase 0 and by separating authority to perform the work from authority to disclose evidence, findings, or reports. The change is not cosmetic. If authority to perform the work is absent, the engagement stops and no Zemi-governed finding is issued. If authority exists but scope, access, mandate, or disclosure authority is limited, the limitation must be justified, stated in the report, and reflected in the classification of any affected finding.

## **The gaps the method is designed to address**

The traditions above leave a practitioner in a specific position: the reasoning tools exist in scholarship, the handling standards exist in consensus documents, and the accountability principle is emerging in position statements, but they are not bound into an operating discipline that a client can read before an engagement and hold the practice to afterward.

Contemporaneous and predecessor work has closed more of that gap than the first companion version recognized. DERDS addressed decision support for reliable reporting. DECDs addressed uncertainty descriptors. Casey addressed preliminary evaluative opinions. SOLVE-IT addressed technique knowledge, weaknesses, and mitigations. FACT addressed person-level attribution. Bollé, Casey, and Jacquet addressed automated-system evaluation. Horsman and Sunde addressed peer review. Sunde and Dror addressed bias and reliability. These are not footnotes around the method. They are the immediate neighborhood in which the method now has to stand.

The narrower statement is this. The reasoning discipline required for claim-testing generally, as distinct from person-level attribution, technique cataloguing, or strength-of-evidence expression, has not been found in a published practitioner doctrine in this form. Neither has the extension of one evidentiary standard across two subjects, so that the same discipline governs an investigation assisted by an AI system and an examination of the AI system itself. Version 1.1 narrows the practical gap further by requiring authority, source completeness, disclosure authority, transfer integrity, retained-custody boundaries, suspect/source narrowing, layered-infrastructure mapping, standards or benchmark mapping, AI role, and tool-derived measurement limits to be made visible in the same report structure as the finding itself. Those are the gaps the Zemi Method attempts to occupy. Whether it occupies them well is a question its casework, review history, and future validation will answer, which is as it should be.

## References

- Anderson, T., Schum, D., & Twining, W. (2005). *Analysis of Evidence* (2nd ed.). Cambridge University Press.
- Biedermann, A., & Kotsoglou, K. (2020). Digital evidence exceptionalism? A review and discussion of conceptual hurdles in digital evidence transformation. *Forensic Science International: Synergy*, 2, 262-274. <https://doi.org/10.1016/j.fsisyn.2020.08.004>
- Bollé, T., Casey, E., & Jacquet, M. (2020). The role of evaluations in reaching decisions using automated systems supporting forensic analysis. *Forensic Science International: Digital Investigation*, 34, 301016. <https://doi.org/10.1016/j.fsidi.2020.301016>
- Casey, E. (2020). Standardization of forming and expressing preliminary evaluative opinions on digital evidence. *Forensic Science International: Digital Investigation*, 32, 200888. <https://doi.org/10.1016/j.fsidi.2019.200888>
- Champhod, C., Biedermann, A., Vuille, J., Willis, S., & De Kinder, J. (2016). *ENFSI Guideline for Evaluative Reporting in Forensic Science: A Primer for Legal Practitioners*.
- Chernyshev, M., Baig, Z., Syed, N., Doss, R., & Shore, M. (2026). Large language models in digital forensics: Capabilities, challenges and future directions. *Forensic Science International: Digital Investigation*, 56, 302043. <https://doi.org/10.1016/j.fsidi.2025.302043>
- Flemyng, E., Noel-Storr, A., Macura, B., Gartlehner, G., Thomas, J., Meerpohl, J. J., Jordan, Z., Minx, J., Eisele-Metzger, A., Hamel, C., Jemio, P., Porritt, K., & Grainger, M. (2025). Position statement on artificial intelligence (AI) use in evidence synthesis across Cochrane, the Campbell Collaboration, JBI and the Collaboration for Environmental Evidence 2025. *Environmental Evidence*, 14, 20. <https://doi.org/10.1186/s13750-025-00374-5>
- Guyatt, G. H., Oxman, A. D., Vist, G. E., Kunz, R., Falck-Ytter, Y., Alonso-Coello, P., & Schunemann, H. J. (2008). GRADE: An emerging consensus on rating quality of evidence and strength of recommendations. *BMJ*, 336(7650), 924-926.

- Hargreaves, C., van Beek, H., & Casey, E. (2025). SOLVE-IT: A proposed digital forensic knowledge base inspired by MITRE ATT&CK. *Forensic Science International: Digital Investigation*, 52, 301864. <https://doi.org/10.1016/j.fsidi.2025.301864>
- Heuer, R. J. (1999). *Psychology of Intelligence Analysis*. Center for the Study of Intelligence, Central Intelligence Agency.
- Horsman, G. (2019). Formalising investigative decision making in digital forensics: Proposing the Digital Evidence Reporting and Decision Support (DERDS) framework. *Digital Investigation*, 28, 146-151. <https://doi.org/10.1016/j.diin.2019.01.007>
- Horsman, G. (2020). Digital Evidence Certainty Descriptors (DECDs). *Forensic Science International: Digital Investigation*, 32, 200896. <https://doi.org/10.1016/j.fsidi.2019.200896>
- Horsman, G., & Sunde, N. (2020). Part 1: The need for peer review in digital forensics. *Forensic Science International: Digital Investigation*, 35, 301062. <https://doi.org/10.1016/j.fsidi.2020.301062>
- International Organization for Standardization. (2012). *ISO/IEC 27037:2012: Information technology - Security techniques - Guidelines for identification, collection, acquisition and preservation of digital evidence*.
- International Organization for Standardization. (2015). *ISO/IEC 27043:2015: Information technology - Security techniques - Incident investigation principles and processes*.
- Kent, K., Chevalier, S., Grance, T., & Dang, H. (2006). *Guide to integrating forensic techniques into incident response* (NIST Special Publication 800-86). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.SP.800-86>
- Neale, C. (2023). Fool me once: A systematic review of techniques to authenticate digital artefacts. *Forensic Science International: Digital Investigation*, 45, 301516. <https://doi.org/10.1016/j.fsidi.2023.301516>
- Neale, C., Kennedy, I., Price, B., Yu, Y., & Nuseibeh, B. (2022). The case for Zero Trust Digital Forensics. *Forensic Science International: Digital Investigation*, 40, 301352. <https://doi.org/10.1016/j.fsidi.2022.301352>
- Popper, K. R. (1959). *The Logic of Scientific Discovery*. Hutchinson.
- Shavers, B. (2025). *The FACT Attribution Framework* (v1.1). Zenodo. <https://doi.org/10.5281/zenodo.18005597>
- Scientific Working Group on Digital Evidence. (2025). *SWGDE best practices for digital evidence collection* (Version 2.0). <https://www.swgde.org>
- Sunde, N. (2021). What does a digital forensics opinion look like? A comparative study of digital forensics and forensic science reporting practices. *Science & Justice*, 61(5), 586-596. <https://doi.org/10.1016/j.scijus.2021.06.010>

- Sunde, N., & Dror, I. E. (2021). A hierarchy of expert performance (HEP) applied to digital forensics: Reliability and biasability in digital forensics decision making. *Forensic Science International: Digital Investigation*, 37, 301175.  
<https://doi.org/10.1016/j.fsidi.2021.301175>
- Toulmin, S. E. (1958). *The Uses of Argument*. Cambridge University Press.
- Wigmore, J. H. (1913). *The Principles of Judicial Proof*. Little, Brown and Company.
- Willis, S. M., McKenna, L., McDermott, S., O'Donnell, G., Barrett, A., Rasmusson, B., Hoglund, T., Nordgaard, A., Berger, C. E. H., Sjerps, M. J., Molina, J. J. L., Zadora, G., Aitken, C. G. G., Lovelock, T., Lunt, L., Champod, C., Biedermann, A., Hicks, T. N., & Taroni, F. (2015). *ENFSI Guideline for Evaluative Reporting in Forensic Science*. European Network of Forensic Science Institutes.

---

*This companion is maintained alongside the method and revised as the surrounding field develops. Its purpose is not to argue that the method is unique. It is to make sure the method's debts are on the record before its contributions are.*