

The Zemi Method

An evidentiary methodology for investigation and assurance

Version 1.2

Kevin V. Watson

August 10, 2026

Final publication revision August 25, 2026

© 2026 Kevin V. Watson. All rights reserved.

Document information

Version: 1.2

Status: Final publication version. Author revision; not peer reviewed or empirically validated.

Drafted: August 10, 2026

Final publication revision: August 25, 2026

Zemi Method Version 1.2 DOI: <https://doi.org/10.5281/zenodo.22084060>

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

Material changes from v1.1

Defines a finding and the evidence-to-finding pathway; states directly that evidence is not a finding; strengthens the definition of independent corroboration; clarifies that the method supports but does not make organizational, legal, operational, or governance decisions; states the method's integrative contribution; connects the method to Artifact-to-Finding Promotion; adds AI Action Provenance and a proportionate Control Provenance requirement for consequential agentic activity; distinguishes execution, control, causation, authorization, intent, autonomy, and accountability; adds representation completeness where human-visible and machine-readable content may differ; and reorganizes the Defensibility Check by theme without removing Version 1.1 controls.

Release basis. Version 1.2 develops the method from the Version 1.1 pressure-test record, post-publication clarity review, and pressure testing against reported agentic-AI security incidents. These activities are structured author review, not empirical validation. External review, engagement-based validation, inter-rater reliability work, reader-comprehension testing, and peer-reviewed publication remain future validation paths.

Contents

Premise	4
Scope and contribution.....	4
Principles.....	5
From source material to finding	6
Finding classifications	7
The investigative lifecycle	7
0. Authority and Scope.....	7
1. Frame	8
2. Preserve	8
3. Corroborate.....	9
4. Analyze.....	9
5. Adjudicate	10
6. Report	10
The assurance practice	11
AI role classification	11
Where automation sits, and where it does not.....	12
Peer review and validation status.....	12
Transition note for examples and templates.....	13
The defensibility check.....	13
Authority and Scope	13
Evidence and Source Completeness	13
Preservation, Transfer, and Custody.....	13
Analysis and Classification	14
Automation and AI.....	14
Reporting and Disclosure	15
Release, Correction, and Version Discipline	15
Version discipline	15

Premise

Every important organizational decision is an evidentiary decision.

Courts, incident response, internal investigations, AI governance, board decisions, compliance, fraud, digital forensics, and AI forensics all eventually ask the same question:

Can this claim be trusted, and what decision can be defended because of it?

Most failures of evidence-based decision-making do not begin with a lack of tools. They begin because an assumption went untested. A system was trusted to record what happened. A log was read as complete. A control was believed to be working because it was documented as working. A summary was treated as the record instead of a path back to the record.

The Zemi Method starts from the opposite position. Every claim that matters to a decision, including claims produced by an organization's own systems, is a hypothesis until proportionate testing establishes what the available evidence supports.

This is not caution for its own sake. An audit record can appear authoritative and still be incomplete in a way that only shows under corroboration. An application-layer log that captures most access may present as the full account while a database-layer record shows access paths the application never recorded. The application record is not necessarily false. It is partial, and partial in a direction no one flagged. The method exists to catch that failure before it becomes the basis for a decision.

Scope and contribution

The Zemi Method applies wherever a decision depends on digital or machine-generated evidence. Its first applications are digital forensic examination, incident reconstruction, expert evidence, litigation support, AI assurance, AI governance evidence, AI incident examination, AI forensics, and synthetic-media assessment. Future applications may change as evidence changes, but the reasoning discipline remains the same.

The method does not replace engagement-specific procedures, legal advice, lawful authority, forensic acquisition or examination standards, tool validation, chain-of-custody requirements, peer review, AI governance frameworks, or professional judgment. It governs the reasoning above them: how claims are framed, how sources are tested, how automation is controlled, how conclusions are bounded, and how findings are reported.

Its contribution is not a new acquisition procedure or a replacement for forensic standards. It integrates claim-as-hypothesis reasoning, source completeness, independent corroboration, human adjudication, bounded finding classifications, AI-role control, and defensibility checking into one repeatable reporting discipline.

The method supports decisions; it does not make them. It produces defensible findings that authorized decision-makers may rely on, challenge, reject, or supplement according to their role and decision

context. It does not make organizational, legal, operational, or governance decisions, and it does not transfer accountability for those decisions to the practitioner or to an automated system.

The depth of verification is proportionate to the stakes of the decision it serves. Where time, cost, authority, access, technical feasibility, or retention limits restrict verification, the restriction is justified, stated as a limitation, and reflected in the classification of any affected finding.

The method may be used in three modes. The mode is determined by the engagement's purpose, decision context, and required form of opinion, not by the examiner's preference. The selected mode and the reason for selecting it are stated in the report.

- **Investigative mode**, where the purpose is to determine what the available evidence can establish about an event, system, claim, or sequence.
- **Preliminary evaluative mode**, where competing propositions are considered and the practitioner expresses what the available evidence appears to support, subject to stated limits and without claiming full evaluative finality.
- **Fully evaluative mode**, where the practitioner evaluates evidence against formal competing propositions for a legal, regulatory, or adjudicative decision. Where a formal likelihood-ratio, probability, or discipline-specific evaluative framework is required, that framework governs the expression of evidential strength. The Zemi Method may govern the framing, preservation, corroboration, automation control, classification, and reporting discipline around that work, but it does not replace the required evaluative framework.

The public method states the philosophy and control structure. Implementation detail, including tools, queries, templates, checklists, and internal review procedures, remains engagement-specific and proprietary.

The Zemi Method is developed and maintained by Kevin V. Watson, carrying on business as Zemi North Advisory. Publication permits the method to be read, cited, reviewed, and tested. It does not place the text in the public domain, waive the stated copyright, or make engagement-specific implementation materials public.

For Zemi North, the method is not a service line. It is the governing discipline from which service lines, reports, assurance frameworks, software, training, and future research should derive.

Principles

Five principles govern every engagement, on both the investigation side and the assurance side.

1. Evidence over assertion. A conclusion stands on what can be shown, not on what a system or a person states. A stakeholder statement, vendor statement, model output, log entry, file artifact, prior report, or tool result is source material or evidence to be tested, not a finding to be recorded by itself.

2. Corroboration and completeness. A single untested assertion is a lead. A source-derived proposition becomes a finding only when the source, extraction, and failure mode have been tested to the proportionate limit of the engagement. Independence means that a corroborating source is not merely

a restatement, derivative, replica, or downstream output of the same evidentiary pipeline. A source is independent only to the extent that it was produced through a materially separate process, control point, system layer, actor, or observation path. A source is also tested for whether the extraction taken from it is complete for the question, because an incomplete extraction can misrepresent the source as surely as the wrong source would. Where independent corroboration is available, it is required. Where it is unavailable, the limitation is reported rather than hidden.

3. Bounded claims. State what the evidence supports and no more. Where the evidence runs out, say where. Every material conclusion is separated into what is known, what is assumed, and what remains undetermined. These classifications describe the epistemic status of a proposition, not its strength on a probability scale. A narrower finding that holds is more defensible than a broader finding with an unsupported edge.

4. AI-assisted, never AI-decided. Automated systems may retrieve, correlate, summarize, propose patterns, and draft. They do not adjudicate. A person owns every conclusion and can trace it to source. When an AI system is itself under examination, its outputs and accounts remain evidence to be tested; the human adjudication boundary does not change.

5. Defensibility by construction. Build the case so a reviewer, opposing expert, regulator, or court can follow every step from source material to conclusion without taking any material step on trust. Defensibility is designed in from the authority and scope gate onward, not added at the end.

From source material to finding

A finding is a human-adjudicated statement that expresses what the available evidence supports, assumes, or cannot resolve within the stated authority, scope, and limits of the engagement.

Evidence is not itself a finding. A log entry, file artifact, AI output, vendor statement, prior report, detector score, or tool result may support, contradict, or qualify a finding, but it does not become a conclusion merely because it was recorded or generated.

The method's pathway is:

Source material → Observation → Claim → Tested claim → Human-adjudicated finding → Decision support

An observation records what was perceived or extracted. A claim gives the observation proposed meaning. Testing examines that claim against source completeness, failure modes, competing explanations, and independent evidence. Human adjudication classifies the resulting finding and states its support and limits. The finding then supports a decision without making the decision itself.

Artifact-to-Finding Promotion, Working Paper Version 1.1, is the aligned application of this discipline to AI-generated and tool-derived outputs. It examines the conditions that should be satisfied before an artifact is elevated into a human finding or relied on for a consequential decision. The Zemi Method Version 1.2 provides the broader reasoning and reporting framework within which that promotion occurs. The working paper applies the same evidence-to-finding pathway, evidentiary-lineage test for

independence, representation-completeness requirement, decision-support boundary, and control-provenance limit used here.

Finding classifications

The method uses three classifications at the human adjudication gate. They are categories of kind, not degrees of confidence.

Known means the proposition is supported by evidence within the stated scope and limits of the engagement. It does not mean absolute certainty or necessarily strong probabilistic support. For each known finding, the report states the support basis: direct evidence, indirect evidence, single independent corroboration, multiple independent corroboration, inference from a bounded record, or a tool-derived measurement with stated processing and uncertainty limits. The support basis explains how the finding is supported without converting the classification into a probability or verbal strength scale.

Assumed means the proposition is relied on for a stated and limited purpose without being directly established by the available evidence. An assumption is identified, bounded, and kept separate from what is known.

Undetermined means the proposition cannot be resolved on the available evidence. Evidence may be unavailable, conflicting, incomplete, outside authority or scope, or insufficiently independent. Undetermined is not a failure to decide; it is a finding that the record does not support a stronger conclusion. An undetermined matter is reported as **pending** where identified evidence or work remains outstanding and resolution is still being pursued, or **exhausted** where available and obtainable evidence has been pursued to the proportionate limit of the engagement.

The investigative lifecycle

Seven phases carry an engagement from authority to finding. The role of automation is stated at each phase, and it narrows deliberately as the work approaches a conclusion.

0. Authority and Scope

Before the question is framed, determine whether the examiner is legally, contractually, ethically, and professionally permitted to ask the question, obtain the evidence, and perform the work. Identify the client, decision-maker, mandate, jurisdictional or policy limits, conflicts, excluded systems, access constraints, privileged material, required standards, and limits on release.

Authority to perform the work is separate from authority to disclose evidence, findings, or reports. 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 engagement proceeds only if the limitation is justified, reported, and reflected in any affected finding.

The method does not create authority. It requires authority to be identified before evidence work is treated as defensible and disclosure authority to be identified before findings or supporting material leave the practice.

1. Frame

Define the question. Identify the mode of work. Surface the assumptions the question carries. Write down what would have to be true for the working theory to hold and what evidence would confirm or break it.

Where the evidence admits more than one explanation, competing explanations are framed and tested against the same evidence together, not sequentially. A theory with no stated breaking condition is not yet an investigation.

Where AI activity is material, avoid framing that presumes autonomy, intent, escape, deception, causation, or responsibility. Each is a separate proposition to be tested.

2. Preserve

Acquire and protect evidence with integrity intact. Maintain chain of custody, use hashing and verified copies where applicable, and record the condition and source of each artifact.

Before an extracted result, query output, parsed artifact, export, or summary is relied on, assess whether the extraction appears complete for the source and question. Prior reports, examiner findings, and investigative summaries are preserved as source material, not inherited as established findings, unless their underlying evidence and reasoning are independently inspected or re-performed.

Identify material filters, date ranges, permissions, pagination, parser limitations, excluded fields, failed records, timezone handling, and known retention limits where they may affect the result.

Where a material source may present different content to human and machine readers, preserve and compare the relevant visible, machine-readable, and transformed representations, and document any material difference. These representations remain views or transformations of the same source unless they arise through materially independent evidentiary lineages.

Where evidence or forensic copies are packaged, encrypted, transferred, or delivered, record the packaging or encryption method, transfer path, authorized recipient, integrity verification, and receipt confirmation. Where evidence, devices, temporary media, workspaces, or data must remain within a defined location or custody boundary, record the boundary, who observed or controlled it, what remained inside it, any temporary media created, and how retained custody or non-transfer was verified.

Where automated systems contribute to the engagement, preserve the inputs provided to them, the outputs returned, and the tool and version records needed to test their contribution later.

Where an AI system can initiate, select, coordinate, or execute consequential actions, preserve the material action trajectory and execution context where available and proportionate. This **AI Action Provenance** includes the assigned objective; model and orchestration components and versions;

prompts and material context; tools, permissions, credentials, and network access; memory or execution state where recorded; configuration and control states; tool invocations and responses; API and network activity; human instructions, approvals, interventions, and overrides; timestamps; environmental responses; and resulting state changes. Hidden chain-of-thought is not required. The evidentiary requirement concerns observable or recorded actions, conditions, and state transitions.

Where attribution is material, also preserve proportionate **Control Provenance**: the available evidence of who or what created, configured, authorized, deployed, operated, modified, supervised, or intervened in the system or its material capabilities. Relevant records may concern accounts, credentials, sessions, devices, administrative activity, deployments, configurations, orchestration, approvals, and communications. A control artifact does not by itself establish human identity, authorship, authorization, intent, or responsibility.

Nothing downstream is defensible if this phase is weak, so it is not rushed.

3. Corroborate

Test each claim against independent sources wherever available and report the limitation where they are not. Corroboration includes testing whether a result is a complete and faithful extraction of its source for the question; an incomplete or unvalidated extraction is a limitation, not the source itself.

A prior report may identify a lead, claimed finding, or gap, but it does not corroborate itself. Test the underlying evidence, reasoning, and failure mode rather than relying on the prior conclusion as evidence of itself.

Where independent sources conflict, treat the conflict as evidence and investigate it. Do not average it or resolve it toward the working theory. A conflict that cannot be resolved is classified as undetermined with both sources shown.

AI-generated accounts, agent logs, orchestration traces, and model explanations do not alone establish that an action occurred or why it occurred. Where material, compare them with independent host, network, authentication, cloud, application, provider, external-system, or human records. Two displays or exports derived from the same orchestration pipeline are not independent merely because they look different.

Do not promote account ownership, credential use, device association, infrastructure ownership, or administrative access directly into a finding about human control or responsibility. Test the material attribution link against independent evidence and plausible alternatives such as delegation, sharing, compromise, inherited access, service operation, or automation.

Automation earns its place here by aligning records across systems. The person decides what agreement means, whether sources are independent, and whether each extraction can be trusted.

4. Analyze

Reconstruct the sequence of events, evaluate competing explanations, and identify what the evidence supports, contradicts, leaves open, or requires as an assumption.

When AI activity is material, distinguish at least the following where relevant: the action technically executed; the component that selected, transmitted, or performed it; the conditions that enabled it; its causal contribution to the outcome; whether it was within authority and scope; whether an objective was assigned or independently selected; whether intent or awareness can be supported; the role of human configuration or intervention; and the location of organizational, legal, or operational accountability. Technical execution does not by itself establish autonomous escape, independent objective formation, intent, or accountability.

Reconstruct action provenance and control provenance separately where attribution is material. Evidence that establishes technical execution does not necessarily establish who exercised material control; evidence of control does not necessarily establish authorization, knowledge, intent, or responsibility. Each transition is a distinct claim requiring proportionate support.

Automation may propose patterns and sequences. The investigator confirms or rejects each against the underlying evidence rather than against the summary.

5. Adjudicate

This is the human decision gate. A person sets conclusions, bounded to what the evidence supports. Findings are classified as known, assumed, or undetermined before they are relied on.

Known findings state their support basis. Undetermined findings state whether resolution is pending or exhausted. Classification remains separate from any statement about evidential strength.

No automated system crosses this line. This is principle four made procedural.

6. Report

State each finding, the evidence behind it, and its limits. Report material contrary evidence alongside the finding. State the reasoning connecting evidence to conclusion rather than leaving the reader to reconstruct it.

Reports identify the authority and scope relied on; the mode of work and why it was selected; why the verification depth was proportionate; sources relied on; completeness limits of material extractions; support basis for known findings; resolution status of undetermined findings; assumptions; material limitations; automation used in the engagement; and the AI role classification where AI is material.

Where evidence narrows a suspect, actor, account, device, system, or source population without uniquely identifying responsibility, state the remaining candidate pool, narrowing criteria, supporting evidence, level of attribution supported, and evidence still required for stronger attribution.

Where a finding depends on layered infrastructure or intermediaries, distinguish technical delivery, account ownership, operational control, legal responsibility, and financial benefit where material, and state which layer each item of evidence supports.

Where a finding depends on compliance with, alignment to, or deviation from a standard, framework, policy, or benchmark, identify the standard and version where material, the requirement relied on, why it applies, the evidence mapped to it, and interpretive limits.

Where a finding depends on a tool-derived measurement, enhancement, comparison, calculation, transformation, or AI-derived output, state the tool and version where material, processing basis, material settings or steps, validation checks, known error or uncertainty limits, and the role the output played in the finding.

Where a finding relies on a derivative presentation, preserve or identify the underlying source, state the selection or editing criteria, link each excerpt to a source locator, identify material excluded context, and state that the presentation is not a substitute for source evidence.

Where consequential AI activity is material, report the action-provenance and, where attribution is material, control-provenance sources and limits. Keep execution, control attribution, causation, authorization, intent, autonomy, and responsibility or accountability analytically distinct. Labels such as "rogue," "escaped," or "autonomous" are not substituted for findings.

Before release, identify authorized recipients; privilege, confidentiality, statutory, contractual, or policy limits; transfer protections and verification where material is transferred; and retained-custody controls where material deliberately remains within a boundary.

The assurance practice

The same discipline applies to a different subject. On the investigation side the method tests claims about what happened. On the assurance side it tests claims about a system itself: what it does, where it fails, under what conditions it operates, what controls govern it, and who remains accountable.

A model's own account of its behavior is a claim to be corroborated. A vendor's stated accuracy is an assertion until testing supports it. A control described in documentation is a hypothesis until evidence shows it operating. The subject changes; the evidentiary standard does not.

For assurance work, the method asks whether evidence supports the claims made about a system, including its capabilities, limits, failure modes, control operation, human oversight, and reliance in decisions. It does not certify governance merely because a policy exists, and it does not replace the governance framework against which a claim is assessed.

Authority and scope matter with particular force in assurance work. Determine whether testing may involve live systems, synthetic datasets, production records, protected classes, model outputs, prompt records, audit logs, configuration files, agent trajectories, external connectivity, or vendor-controlled evidence. Revisit authority and scope when access or system boundaries change. Where a conclusion depends on inaccessible evidence, report that limitation rather than absorb it into the finding.

AI role classification

Where an AI system is material, classify its role before drawing conclusions. The classification describes what the evidence shows the system did. It does not assign legal responsibility and does not accept labels, logs, prompts, model explanations, or vendor statements as proof without corroboration.

- **AI present:** an AI system exists in the environment, artifact set, or workflow, but the evidence does not show that it materially affected the activity under examination.
- **AI-assisted:** an AI system supports a human actor through drafting, troubleshooting, summarization, coding, search, or analysis, but the evidence does not show it independently selecting objectives or executing material actions.
- **AI-orchestrated:** an AI system coordinates or triggers material workflow steps under human-configured objectives, rules, prompts, agents, tools, or integrations, while accountability remains with the human or organization controlling the system.
- **AI-autonomous claimed:** a claim that an AI system independently selected objectives, made decisions, or executed material actions without human control. This classification requires independent evidence and is not accepted from labels or system-generated records alone.

Report the classification with its evidence and limits. Evidence of autonomous execution does not necessarily establish autonomous objective selection, escape, intent, or accountability. Where autonomy is claimed but not independently supported, the claim remains undetermined.

Where automation sits, and where it does not

Automation may support retrieval, correlation, pattern proposal, drafting, and summarization. It is barred from conclusion, adjudication, and sign-off.

The reason is accountability. A conclusion must be owned by someone who can be questioned, trace the finding to source, disclose material limits, and be held accountable for the judgment.

Preserve and document automated contributions so they remain separable from human judgment. When an automated system is the subject of the investigation rather than a practitioner tool, preserve its material action provenance and test its records against independent evidence.

Peer review and validation status

Publication makes a method dated, citable, and open to scrutiny. It does not make it peer reviewed, empirically validated, or accepted by the field.

The method is both a governing discipline for practice and a research object to be tested. Known validation questions include whether practitioners classify findings consistently, whether readers understand the classifications, whether the authority and scope gate improves defensibility, whether AI Action Provenance and proportionate Control Provenance are practicable across different agent architectures, and whether the method improves reporting quality without disproportionate burden.

Structured expert review, case-based testing, inter-rater reliability work, reader-comprehension testing, engagement-based validation, and eventual peer-reviewed publication remain future paths. Until that work exists, the method's status is bounded: it is a practitioner methodology under author review, not a validated scientific technique.

Transition note for examples and templates

Materials prepared under Version 1.1 do not automatically satisfy Version 1.2. Before being described as Version 1.2-conformant, examples, report skeletons, templates, checklists, and visual summaries must incorporate the finding definition and evidence-to-finding pathway, the strengthened independence test, the decision-support boundary, representation completeness where material, AI Action Provenance where consequential AI activity is material, and Control Provenance where attribution is material. This requirement does not mean every artifact needs an AI, attribution, or cross-representation section where the issue is irrelevant.

The defensibility check

Before a finding leaves the practice, it passes a standing test.

Authority and Scope

- Is lawful, contractual, ethical, and professional authority identified?
- Is the scope of the question, evidence access, and excluded material stated?
- Is the mode identified and justified by the engagement's purpose?
- Are proportionality, access, technical-feasibility, or scope limits justified rather than merely noted?
- Is disclosure authority identified before evidence, findings, or reports are released?

Evidence and Source Completeness

- Have material queries, exports, parses, extracted results, and summaries been assessed for completeness against the source and question?
- Where a material source may differ between human-visible, machine-readable, or transformed representations, have those representations been preserved, compared, and any material difference documented?
- Have prior reports, findings, and investigative summaries been treated as source material and tested against underlying evidence rather than inherited as findings?
- Can every finding be traced to inspectable source material and stated reasoning?
- Does any claimed corroboration arise through a materially separate process, control point, system layer, actor, or observation path rather than the same evidentiary pipeline?

Preservation, Transfer, and Custody

- Is the condition, source, integrity, and custody history of material evidence recorded?
- Where evidence or forensic copies are transferred, are packaging or encryption, transfer path, authorized recipient, integrity verification, and receipt confirmation recorded?

- Where material must remain within a custody boundary, are the boundary, observer or controller, retained material, temporary-media handling, and verification of non-transfer recorded?
- Where consequential AI activity is material, are the available action trajectory, execution context, control states, human interventions, and resulting state changes preserved proportionately?
- Where attribution is material, is the available evidence of material control preserved proportionately and separately from the action trajectory?

Analysis and Classification

- Are observations, claims, tested claims, and findings kept distinct?
- Are claims bounded to what the evidence supports?
- Are known facts, assumptions, and undetermined matters separated?
- Do known findings state their support basis without converting classification into a probability or verbal strength scale?
- Do undetermined findings state whether resolution is pending or exhausted?
- Is each classification consistent with the evidence record?
- Have competing explanations and material conflicting or contrary evidence been tested and reported?
- Where evidence narrows a candidate population without uniquely identifying responsibility, are the remaining pool, narrowing criteria, support, attribution level, and evidence required for stronger attribution stated?
- Where layered infrastructure is material, are technical delivery, account ownership, operational control, legal responsibility, and financial benefit distinguished, without treating association as proof of control or responsibility?
- Would the finding hold if an independent examiner repeated the work from the preserved evidence?

Automation and AI

- Is the contribution of automation documented, preserved, and separable from human judgment?
- Where AI is material, is its role classified and bounded to what the evidence supports?
- Are AI or orchestration records corroborated where they are relied on to establish consequential actions?
- Are technical execution, control attribution, causation, authorization, objective selection, intent, autonomy, human intervention, and responsibility or accountability distinguished where material?
- Where a finding depends on a tool-derived or AI-derived output, are processing basis, validation checks, uncertainty limits, and the output's role stated?
- Was every finding human-adjudicated and signed off by an accountable person?

Reporting and Disclosure

- Does each finding state its evidence, reasoning, classification, support basis where known, contrary evidence, assumptions, and material limits?
- Where a finding depends on a standard, framework, policy, or benchmark, are the requirement, applicability basis, evidence mapping, and interpretive limits stated?
- Where a finding relies on a derivative presentation, are the source, selection criteria, source locator, excluded context, and non-substitution limit stated?
- Are material limitations clear enough for a decision-maker to account for them?
- Does the report support the decision without purporting to make it?

Release, Correction, and Version Discipline

- Are the method version and any material artifact-conformance limits identified?
- Is release limited to authorized recipients under the applicable disclosure constraints?
- Is there a record sufficient to correct a finding later if it proves unsupported, incomplete, or materially misstated?

A finding that fails any applicable item returns to the phase where the gap opened. Applicability is stated rather than assumed away. Nothing reaches a client on the strength of the method's reputation alone; it reaches them on the strength of the record behind it.

Version discipline

The Zemi Method is maintained in dated, versioned form. Each public version records its publication date and material changes. Engagement reports identify the version relied on.

Versioning matters because the method is not written after the fact to fit a conclusion. It exists before the engagement, guides the work, and constrains the finding.

Version 1.2 is an intentionally approved final publication version. It does not alter the historical Version 1.1 DOI record or retroactively make any Version 1.1-conformant supporting artifact conformant with Version 1.2.

If a finding issued under the method is later shown to be unsupported, incomplete, or materially misstated against the preserved evidence, issue a correction identifying the finding, reason, and evidence with the same visibility as the original.

The Zemi Method is maintained in dated, versioned form at its canonical home, <https://zemimethod.com>. Comments and scholarly correspondence are welcome and inform future versions.
