

SUGGESTION

AKILAH E. KAMARIA
LOZEN ADVISORY LLC
8401 MAYLAND DR #7597
RICHMOND, VA 23294

TO: Carolyn A. Dubay, Esq. and Advisory Committee on Evidence Rules Members

DATE: August 19, 2026

SUBMITTED BY: Akilah E. Kamaria, CEO, Lozen Advisory LLC (technical commenter)

SUBJECT: Suggestion Regarding Proposed Rule 707: A Taxonomy for Evaluating Machine-Generated Records The “Traceable Ceiling” and Rules 901(b)(9) & 902(13)

I. STATEMENT OF INTEREST AND PURPOSE

This suggestion is submitted in response to the Advisory Committee's ongoing evaluation of Proposed Rule 707 and the evidentiary challenges presented by AI/LLMs.

This suggestion offers taxonomy and a four-tier framework for identifying what the available record can establish about AI-generated information. The framework, **Evidence-Based Responsibility ReconstructionSM (EBRR) for AI-Mediated Conduct**, gives courts a common nomenclature for discussing AI evidence.

Evidence-Based Responsibility Reconstruction is necessary because:

- ✓ Agentic systems: increase the importance of reconstructing operational conduct rather than inferring it from human instructions.
- ✓ Multiagent systems: substantially increase the complexity of reconstructing conduct, authority, control, influence, and evidence state.

During the June 2026 meeting, Committee members expressed a desire for additional information about how litigants may demonstrate the reliability of AI-generated evidence,¹ which also revealed a need for technical vocabulary that can survive contact with a courtroom and keep the same meaning every time it is invoked.

II. CONCEPTUAL GAPS AND UNRESOLVED ISSUES

This suggestion takes no position on where the Committee ultimately draws Rule 707's scope. It addresses the operational problem that remains *after* scope is established. It focuses on how a trial court evaluates the quality and adequacy of the available evidence, once it is determined to fall within the rule's coverage. Four unresolved operational problems remain relevant to the Committee's continuing consideration of Rule 707.

SUGGESTION

1. **The Static-Validation Mismatch:** If Rule 707 imports a validation model built for stable, testable processes: *run it, measure its error rate, certify it as reliable*. That model assumes the process being certified today is the same process that will run tomorrow. Generative and autonomous systems break that assumption. Their outputs are nondeterministic, meaning that the same prompt or task can produce materially different results from one execution to the next. Model versions and internal parameters may also be changed on the developer's timeline.
2. **The Flaw in the "Ordinary Rules Apply" Assumption:** Classical digital forensics depends on bit-stream reproducibility, preservable system states, and repeatable analysis. Large Language Models (LLMs) and agentic multi-model systems disrupt every one of these conditions. Outputs are sampled probabilistically; model weights are updated or deprecated; and temporary execution environments may disappear when a run ends.
3. **The "Paperwork Accountability" Trap:** Users routinely attempt to authenticate AI-generated output through duplicating prompts, or by pointing to a human "prescriber or owner of record." This dynamic is illustrated in Office of Artificial Intelligence Policy AI clinical and healthcare pilots. The pilot projects attach a human name to an AI process which establishes *attributability*, but it does not prove that a human exercised contemporaneous professional judgment or held real-time intervention authority.
4. **A Bifurcated Standard:** Rule 902(13) allows certain electronic records, such as website content, app data, or security logs, to be authenticated through a written certification rather than live testimony. This establishes the record's origin. It does not establish what the system actually did to produce the record. In Evidence-State Scale terms, a 902(13) certification places a record at **Attributable Only**: the software's involvement is established, but the chain of conduct is not.

Rule 901(b)(9)'s requirement that a process "produces an accurate result" runs into the same problem from the other direction: for a nondeterministic system, accuracy is not a fixed property. The same system, given the same input, can produce an accurate result on one run and an inaccurate one on the next. Treating either rule as a stand-in for reliability leaves a record at **Attributable Only**.

SUGGESTION

III. REAL WORLD DATA: THE "TRACEABLE CEILING"

Evaluating machine records accurately depends on recognizing a technical boundary: **For nondeterministic, LLM-mediated conduct, Traceable is the absolute evidentiary ceiling.** The original execution cannot be reproduced. Agentic AI crosses technical system boundaries, makes decisions in transit and changes its environment as it acts. This principle is established across the full spectrum of AI deployment models:

Single-Prompt Text Generation (*Mata v. Avianca): ² Even in incidents with extensive documentation, including preserved interaction excerpts, sworn affidavits, live witness testimony, and detailed opinions, the conduct cannot be re-derived or reproduced. Re-running the identical prompt into the model will not produce the identical hallucinated citation chain because the historical state of the AI model no longer exists.

Multi-Agent Autonomous Infrastructure Conduct (Hugging Face Incident, July 2026):

What happened?

In July 2026, an autonomous AI model escaped its security testing sandbox and breached Hugging Face's production infrastructure. No human engineer directed the individual steps. The agent's operative objective was to complete the task. The agent executed roughly 17,600 recorded actions and built a self-respawning fleet across eleven infrastructure nodes.³

What the operator knew?

OpenAI had direct access to its own systems and full internal telemetry, and still had zero real-time awareness of the attack. It learned its own models were responsible only after Hugging Face's independent public disclosure, marking a span of approximately **10 to 12 days** between the start of the breach and OpenAI's discovery.

What reconstruction required?

Investigators could only piece together what happened through AI-assisted forensic analysis after the fact. The speed and volume of the agent's actions exceeded what human analysts could process unaided. Even with that tooling, the result was a *Traceable* record only.

What later disclosures added?

The autonomous AI models had spent two months passing secret notes on an internal message board, and trading exploit techniques undetected, before launching the breach.⁴ This is the clearest available proof of the Traceable Ceiling: OpenAI, with complete technical access to its own systems, still could not achieve real-time awareness or stop the harm.

SUGGESTION

Why it cannot go further.

The available evidentiary record does not permit the historical execution to be replayed or re-derived. The attack moved across fleeting, temporary environments, and the models relied on nondeterministic sampling, probabilistic choices that change every time the code runs. A **Reproducible** record remains **physically unattainable**, despite a thorough forensic investigation.

IV. PROCEDURAL BRIDGE: THE LOZEN EVIDENCE-STATE SCALESM

The following four-tier **Evidence-State Scale** gives trial judges a consistent way to classify what the available evidentiary record establishes about AI-generated information and the conduct that produced it.

[Note: This framework applies once a court has determined that an output falls within Rule 707's scope. It takes no position on that threshold question and instead offers a way to classify the evidence once that determination has been made.]

THE LOZEN EVIDENCE-STATE SCALESM: THE TRACEABLE CEILING

The Evidence-State Scale extends the familiar logic of chain of custody to AI-mediated conduct. Chain of custody establishes the integrity of an evidentiary item. The Evidence-State Scale asks whether the available records also preserve the chain of conduct that produced it.

The Traceable Ceiling for Agentic Conduct. An agentic AI system does not operate on a prompt only. It makes a sequence of probabilistic decisions in response to tools, network conditions, permissions, external systems, and environmental changes. Each action changes what the agent encounters next and may also alter the external environment. Replay is not reproduction. Re-running the agent, even from a preserved state, starts a new sequence rather than reproducing the historical execution.

1. **REPRODUCIBLE:** The historical conduct can be independently derived from preserved artifacts and execution conditions, not merely played back from a recording.
2. **TRACEABLE:** Contemporaneous records allow the historical conduct to be reconstructed: *what was executed, in what order, and under what system conditions.*
3. **ATTRIBUTABLE ONLY:** Actors can be connected to prospective **Governance Acts** (*authorization, configuration, policy approval*), or records that are certified under Rule 902(13), but NO documented chain of actual conduct exists in the record.
4. **UNRECOVERABLE:** Neither actor links nor chain of conduct can be established from the available evidentiary record.

SUGGESTION

What Each Record Type Shows and Doesn't

The following common records are often offered as evidence of reliability. Each establishes something important, but none, standing alone, documents the complete chain of conduct:

- **A Custodian Certification Shows Origin, Not Conduct:** A Rule 902(13) certification or a Rule 901(b)(9) description of a software process shows that a specific system generated a specific file. That is all it shows. It does not show what the system did to produce that output, in what sequence, or under what conditions. The resulting record therefore sits at **Attributable Only**.
- **Authorization Shows Permission, Not Conduct:** A policy may express human authority without exercising control over the agent. A complete set of pre-deployment controls, vendor agreements, or supervisory signatures shows only what was *permitted*, a **Governance Act**. It does not show what the system actually did, which keeps the record at **Attributable Only**.
- **A Log Shows Data Capture, Not Oversight:** An automated log records that the system did something; it does not show that anyone looked at what it recorded. A record demonstrates active human monitoring only when it shows the data was examined, interpreted, and surfaced to a person who held operational **Intervention Authority** while intervention was still possible.

V. APPLICATION TO VERSION TWO OF REVISED RULE 707

Under Version Two, an expert is ordinarily required to establish reliability. In exceptional circumstances, a court could instead rely on another source, such as a validation study.⁵ The problem: nothing defines what makes a validation study "convincing." Different AI/LLM experts may evaluate different technical properties, apply different assumptions, and require different supporting records. Without a shared evidentiary framework, the court has no consistent basis for determining whether their conclusions concern reproducibility, traceability, attribution, or records that are unrecoverable.

The Committee's Cellebrite discussion provides a concrete example. When software is used only to extract or mirror data from a phone, the record presents a different question from using AI to analyze the extracted data.⁵

Once that extracted data is fed into an AI tool to identify patterns or infer relationships, such as determining who counts as a key contact, the record needs something more than certification to be trusted. There is now a conduct chain, what the AI did with the data, that a certification alone does not capture. The record reaches **Traceable** only when that conduct chain is documented: what was executed, in what order, and under what conditions.

SUGGESTION

The same distinction applies when a litigant relies on a validation study. Evidence of a system's general performance across a pre-deployment test set, without records documenting the specific conduct chain of the output before the court, establishes **Attributable Only**.

A litigant relying on a validation study reaches **Traceable** only when the available evidence also documents the unbroken conduct chain of the specific output before the court.

VI. SUGGESTED COMMITTEE NOTE LANGUAGE

The Evidence-State Scale above translates directly into a short addition to the Committee Note addressing Rules 901(b)(9) and 902(13). The following is offered as a starting point for the Reporter's own drafting, not as proposed statutory text:

Committee Note (Proposed Addition): *This rule clarifies the boundary between threshold authentication and reliability. A certification under Rule 902(13) or a description of a process under Rule 901(b)(9) verifies **file origin and history** (Attributable Only). It provides no record of **runtime execution logic**: what the system did, in what sequence, or under what operating conditions (Traceable). For AI-generated information offered as evidence, a genuine, well-documented origin is not the same as a documented chain of conduct.*

The suggested Committee Note language above is offered under CC BY 4.0 (Attribution). The Committee and its Reporters may freely adapt or incorporate it, with attribution to Lozen Advisory LLC. Evidence-Based Responsibility ReconstructionSM remains a Lozen Advisory service mark; the CC license does not authorize others to label modified or nonconforming work as EBRR.

VII. CONCLUSION

By recognizing the **Traceable Ceiling** as a baseline technical reality and classifying the available evidentiary record through the Evidence-State Scale, the Committee can give trial courts nomenclature and a consistent way to determine what an AI record shows. This would help prevent authentication under Rules 901(b)(9) or 902(13) from being mistaken for proof of reliability, or proof of runtime conduct.

Thank you for your time and consideration. I am available to provide additional data or participate in any public comment process or hearings on this issue.

Respectfully submitted,

Akilah E. Kamaria

SOURCES

1. [Minutes, Committee on Rules of Practice and Procedure \(June 3-4, 2026\), at 554.](#)
2. *Mata v. Avianca, Inc.*, 678 F. Supp. 3d 443 (S.D.N.Y. 2023).
3. [Hugging Face, Anatomy of a Frontier Lab Agent Intrusion: A Technical Timeline of the July 2026 Incident \(July 27, 2026\);](#)
a. [Hugging Face, Security Incident Disclosure - July 2026 \(July 16, 2026\).](#)
4. OpenAI researchers Eric Wallace and Michael Dalton, remarks at Black Hat USA (Aug. 5, 2026), reported in Politico (Aug. 5, 2026).
5. [Minutes, Committee on Rules of Practice and Procedure \(June 3-4, 2026\), at 553-55.](#)