

SAMPLE BRIEF FOR ADMISSIBILITY CHALLENGE

PREPARED BY ARGENT FORENSICS

This document is for demonstration purposes only and does not represent legal advice or clinical guidance. Practitioners should apply their judgement as to the application of these concepts to their real-world practice decisions. No particular outcome is guaranteed.

May, 2026

IN THE UNITED STATES DISTRICT COURT
STATE OF MINNESOTA

[PLAINTIFF/STATE/PETITIONER],

vs.

[DEFENDANT/RESPONDENT].

Case no. _____

MEMORANDUM OF LAW IN SUPPORT
OF THE ADMISSIBILITY OF EXPERT
TESTIMONY UTILIZING AI-ENHANCED
COMPUTER-ASSISTED RECORD
REVIEW (ARGENT)

MEMORANDUM IN SUPPORT OF EXPERT ADMISSIBILITY

I. Introduction

This memorandum is submitted in support of the admissibility of the expert psychological testimony of Dr. [Expert]. A motion has been filed seeking to exclude or limit Dr. [Expert]'s testimony under the premise that the use of **Argent Forensics**, a specialized, computer-assisted record review platform, compromises the reliability of the underlying clinical methodology.

This mischaracterizes the nature of modern computer-assisted record review and obscures the fundamental distinction between an autonomous decision-making algorithm and a clinical instrument.

Argent does not form diagnostic opinions, draw forensic conclusions, or substitute for the cognitive processes of the testifying expert. Rather, it functions as a highly sophisticated digital highlighter and index clerk. It utilizes a secure **Retrieval-Augmented Generation (RAG)** architecture to locate, organize, and summarize voluminous collateral records. Crucially, every factual extraction generated by the platform is anchored by a non-negotiable, page-level citation back to the original source document, which the expert manually verifies.

Because the expert maintains sole intellectual agency – manually cross-checking every cited fact against the raw record before incorporating it into the clinical formulation – the methodology meets and exceeds the reliability standards required by *Daubert v. Merrell Dow Pharmaceuticals, Inc.* 509 US 579 (1993) and *Kumho Tire Co. v. Carmichael*, 526 US 137 (1999).

II. Evidentiary Framework Under *Daubert* and *Kumho Tire*

Under Federal Rule of Evidence 702 (and corresponding state standards), the trial court serves as a “gatekeeper” to ensure that an expert’s testimony rests on a reliable foundation and is relevant to the task at hand. The Supreme Court further clarified, in *Kumho*, that this gatekeeping obligation applies not only to “scientific” knowledge, but to all “technical” or “other specialized” knowledge.

Importantly, Federal Rule of Evidence 702 was amended to clarify that the proponent of expert testimony must demonstrate that it is more likely than not that the admissibility requirements are satisfied. This shifted the burden of proof on admissibility more explicitly to the party offering the expert. This development underscores the importance of clearly articulating the reliability of machine-assisted methodologies, which this memorandum does in detail below.

The Crucial Distinction Between Open-Domain Generative AI and Closed-Domain RAG

There may be efforts to conflate Argent’s Retrieval-Augmented Generation (RAG) architecture with the “open-domain” generative artificial intelligence models that have recently drawn judicial scrutiny. See, e.g. *Mata v. Avianca*, 678 F. Supp 3d 443 (SDNY 2023) (addressing the dangers of ungrounded, unverified consumer chatbots hallucinating legal citations). This comparison misrepresents the underlying computer science.

Ungrounded consumer chatbots operate “closed book,” relying on mathematical probability to reconstruct plausible-sounding answers from their general training data – a process highly vulnerable to “hallucinations.” In stark contrast, Argent’s RAG architecture operates “open book” in a closed-domain system restricted entirely to the specific, uploaded collateral records of the case. The Large Language Model (LLM) within Argent does not serve as an independent source of information; rather, it serves as a linguistic formatting engine.

The primary, data-finding phase of Argent’s pipeline is functionally equivalent to the standard Technology Assisted Review (TAR) and e-discovery indexing tools that federal

and state courts have approved for many years. See *Da Silva Moore v. Publicis Groupe*, 287 FRD 182 (SDNY 2012). Argent applies a highly sophisticated translation layer to those retrieved, cited segments. Because the platform hardcodes page-level citations to every output, and because the expert’s protocol requires manual verification of those cited original pages, the risk of unverified “hallucination” in the final work product is substantially reduced.

The Supreme Court has repeatedly emphasized that the *Daubert* inquiry is a flexible one. The five identified factors: 1) Testability; 2) Peer Review and Publication; 3) Rate of Error; 4) Existence and Maintenance of Standards and Controls; and 5) General Acceptance, do not constitute a rigid, mandatory checklist. Instead, the focus must remain on the expert’s principles and methodology, not on the conclusions they generate.

Furthermore, courts have consistently distinguished between “black-box” systems that attempt to automate clinical or legal outcomes and technology-assisted review tools that merely make human review more comprehensive and efficient. See *Da Silva Moore v. Publicis Group*, 287 FRD 182 (SDNY 2012), holding that computer-assisted document review (Technology Assisted Review, or TAR) and predictive coding are acceptable and highly reliable methods for analyzing large volumes of electronically stored information. Additionally, see in *United States v. O’Keefe*, 537 F Supp. 2d 14 (DDC 2008), recognizing computer-assisted review as a legitimate and efficient methodology. As demonstrated below, Argent falls squarely into this category, satisfying each of the *Daubert* reliability factors.

III. Factor-By-Factor *Daubert* Analysis

FACTOR 1: TESTABILITY & VERIFIABILITY

Whether the expert’s technique or theory can be (and has been) tested and assessed for reliability.

The first *Daubert* factor assesses whether a technique can be disproven, tested, or refuted through replicable results. Dr. [Expert]’s record review methodology using Argent is testable in the most direct, empirical sense: **every factual claim is accompanied by an audit trail.**

Unlike consumer-grade artificial intelligence models that generate answers from a closed-book, pre-trained “memory” which cannot be independently verified and is often protected by proprietary trade secret status, Argent operates on a Retrieval-Augmented Generation (RAG) architecture. See *Lewis, P. et al., Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks*, arXiv:2005.11401 (2020), establishing RAG as the

foundational peer-reviewed architecture for constrained, document-specific AI retrieval. Argent is consistent with this in that when the expert queries the case file, the system is constrained to search only the specific, uploaded document set. It retrieves the relevant text segments and attaches a page-level citation to the summary output.

Harmonizing this with traditional practice, there is very little distinction between this and manual record review methods where an expert can be challenged on proving the foundation of a cited fact by pointing to its location in the collateral record. The fact that AI helped illuminate the fact is not relevant to the factor of testability as long as the fact itself is accurate, reliable, and consistent with the source document.

Takeaway: *The technique under review is not autonomous AI decision-making, but computer-assisted expert document review with a full, citation-anchored audit trail that is more verifiable than traditional manual review.*

FACTOR 2: PEER REVIEW AND PUBLICATION

Whether the technique or theory has been subject to peer review and publication.

Two distinct bodies of published, peer-reviewed literature validate Dr. [Expert]'s methodology:

- **Computer Science and Information Retrieval Literature:** The engineering foundation of Argent – specifically Retrieval-Augmented Generation (RAG) and semantic vector search – is not experimental or novel. RAG has been extensively documented, tested, and peer-reviewed in the primary computer science literature as the industry-standard methodology for secure, factual document analysis. See Lewis et al., *supra* (2020); Guu, K. et al., *REALM: Retrieval-Augmented Language Model Pre-Training*, arXiv:2002.08909 (2020); Izacard, G. & Grave, E., *Leveraging Passage Retrieval with Generative Models for Open Domain Question Answering*, arXiv:2007.01282 (2021). These peer-reviewed works collectively establish that RAG architecture is a well-validated, reproducible, and empirically tested methodology.
- **Forensic and Clinical Literature:** While the literature base specifically addressing AI-assisted tasks in forensic psychology and psychiatry is emerging, it is rapidly growing. Scholars and practitioners have begun publishing on the ethical, practical, and clinical boundaries of AI tools in forensic evaluations. Notably, Argent's founder has published peer-reviewed specifically addressing AI in forensic psychological practice – see Rilen, S. *Practical Frameworks for Integrating Artificial Intelligence in Forensic Psychology*, *Journal of Forensic Psychology Research and*

Practice (2026), where the author explicitly discusses the use of AI-assisted record review tools. This establishes a direct line of academic inquiry into the responsible, human-in-the-loop use of AI in high-stakes legal evaluations.

Under *Daubert*, the lack of a decades-long clinical trial history does not equate to a lack of reliability. As the Supreme Court noted, emerging technologies may be underrepresented in the literature and still be scientifically valid. Where literature is still developing, courts look to alternative indicators of reliability – such as the human-in-the-loop framework – which relies on a well-established methodological tradition of careful clinical verification and detail-oriented forensic focus on transparency and accuracy in record review.

Takeaway: *The underlying engineering principles of RAG and vector database search are universally peer-reviewed. Peer-reviewed forensic scholarship specifically addressing AI in forensic practice is emerging, anchored by publications including those of the platform's founder.*

FACTOR 3: KNOWN OR POTENTIAL RATE OF ERROR

The known or potential rate of error of the technique or theory.

It might be argued that the exact error rate of an Argent-assisted clinician has not been empirically quantified. However, this argument ignores both the standard of comparison and the platform's built-in safeguards.

First, the baseline error rate of *unaided, manual review* is also mathematically unknown, though empirically understood to be vulnerable to cognitive fatigue, overlooked documentation, potential inadvertent biases, or diagnostic omissions when reviewing files spanning thousands of pages. Research on human cognitive limitations on complex document review is well-established. See, Sweller, J. *Cognitive Load During Problem Solving: Effects on Learning*, 12 Cognitive Science 257 (1988), establishing theoretical foundation for cognitive overload in complex tasks. Argent is designed to *exceed* the comprehensiveness of manual review by acting as an exhaustive semantic index of the entire record set.

Second, it is granted that some element of error is introduced through machine-mechanisms. Because large language models are inherently non-deterministic, they are capable of generating errors that include misattributions or “hallucinations.” Rather than concealing this technical reality, Argent's methodology is designed entirely around mitigating it. The platform enforces a human-in-the-loop verification protocol. Because every output must carry page-level citations, the clinician must physically inspect the original document to verify each factual claim before incorporating it into their final report. Because the final certifier is the human expert, the effective error rate of the evidence presented in court is reduced to the expert's own professional standard of care.

The platform also maintains a citable audit trail so that the court or parties could theoretically reproduce the results in question to examine when and how an AI-Assisted Artifact was produced, in a manner impossible to replicate with a traditional manual review.

Lastly, Argent maintains informal tracking of output error rates by feature type and user-reported errors to drive platform improvements, creating an ongoing mechanism for error rate monitoring consistent with sound scientific practice.

Takeaway: *Argent is a tool designed to reduce human error rather than introduce machine error, as long as it is used as designed – with mandatory citation verification by the expert before any output informs their analysis.*

FACTOR 4: EXISTENCE AND MAINTENANCE OF STANDARDS AND CONTROLS

The existence and maintenance of standards controlling the technique's application or operation.

A reliable methodology must be governed by standardized procedures that minimize variability, bias, and technical drift. This factor, clarified and emphasized in FRE 702, asks not only whether standards exist, but whether they are actively maintained. Rather than claiming absolute, technically impossible mathematical determinism, Argent's standards and controls are established through a series of robust operational safeguards:

- **Immutability of the Citation Pipeline:** The requirement for page-level, document-specific citations is hardcoded into the platform's core architecture. It cannot be deactivated, bypassed, or modified by the user, ensuring an unbroken audit trail for every case file and every output produced.
- **Standardized Procedural Protocol:** The platform mandates a strict, multi-step clinical process: *Secure Ingestion* → *OCR Processing* → *Semantic Retrieval* → *Page-Level Verification* → *Disclosed Report Extraction*. This sequence reduces practitioner variability and enforces systemic rigor across all users and all cases.
- **Compliance and Data Governance Controls:** Platform operations are strictly bound by the technical and administrative safeguards of the Health Insurance Portability and Accountability Act (HIPAA). Data is strictly isolated to account-specific environments protected by AWS Key Management Services (KMS) decoupled decryption keys, and governed by contractually binding Business Associate Agreements (BAAs) with the platform and all sub-processors that handle Protected Health Information.
- **Professional Ethics Framework:** The platform's design and operation are explicitly aligned with the American Psychological Association (APA) Ethics Code

(specifically Sections 8.01 *Institutional Approval* and 9.01 *Bases for Assessments*, as well as the Specialty Guidelines for Forensic Psychology (SGFP). Practitioners using the platform are provided with professional resources including a completed APA AI Tool Evaluation Checklist (APA Office of Healthcare Innovation, October 2024), informed consent templates, and report disclosure language consistent with these standards. These practitioner-facing controls constitute an additional layer of operational governance.

- **Formal Error Reporting and Quality Control:** Argent maintains a systematic quality control feedback loop. Users can report factual inaccuracies directly within the platform interface, triggering targeted technical reviews. Every reported error receives a direct response, and error rates are tracked internally by output type. This constitutes the ongoing standards maintenance this *Daubert* factor contemplates.

Takeaway: *Standardized controls reside in rigid system constraints -such as mandatory page citations and a documented multi-step verification protocol - combined with a comprehensive legal governance framework through HIPAA compliance and BAA agreements, and professional ethics alignment through APA and SGFP standards.*

FACTOR 5: GENERAL ACCEPTANCE IN THE SCIENTIFIC COMMUNITY

Whether the technique or theory has been generally accepted in the relevant scientific community.

While Argent is one of the first AI-assisted platforms built specifically for forensic psychology and psychiatry, the underlying methodology of computer-assisted document review has achieved overwhelming general acceptance across adjacent professional and legal fields.

Within the broader legal community, the use of Technology-Assisted Review (TAR) and e-discovery algorithms to search, categorize, and summarize voluminous document sets has been universally accepted for over a decade. See *Da Silva Moore v. Publicis Group*, 287 FRD 182 (SDNY 2012) (“computer-assisted review is an available tool and should be seriously considered for use in large-data volume cases where it may save the producing party (or both parties) significant amounts of legal fees in document review”); *United States v. O’Keefe*, 537 F Supp. 2d 14 (DDC 2008) (recognizing computer-assisted search and review as standard methodology for large document productions); *Equity Analytics, LLC v. Lundy*, 248 FRD 54 (DDC 2008) (court outlining framework for e-discovery search-reliability). These decisions establish that computer-assisted document review is not merely tolerated but affirmatively

endorsed by federal courts as a more reliable alternative to purely manual review of large document sets.

Within clinical practice, the trajectory toward acceptance is clear and accelerating. The American Psychological Association (APA) has published formal guidance for clinicians evaluating AI software. Inherent in this is the acknowledgement that generative AI use is acceptable in many facets of clinical practice, accompanied by appropriate ethical constraints and properly applied decisional frameworks. One such tool the APA has released is the *Companion Checklist: Evaluation of an AI-Enabled Clinical or Administrative Tool* (APA Office of Health Care Innovation, October 2024), again presupposing that psychological AI tools are a legitimate subject of professional evaluation and judgement, rather than per se inappropriate. The Specialty Guidelines for Forensic Psychology do not prohibit technological assistance in record review. And the American Board of Professional Psychology has released an AI-policy which recognizes record review as a legitimate option as well, assuming appropriate disclosure; this reflects the field's active and ongoing engagement with AI rather than rejection of AI assistance.

The frank characterization of the general acceptance factor is that the forensic psychology community is in an active period of methodological development regarding AI tools, which is working toward guardrails, constraints, and specialized decisional frameworks rather than rejection. Practitioners who adopt responsible AI assistance with appropriate disclosure, citation architecture, and human verification are practicing at the leading edge of an emerging professional standard, not in deviation from an existing one.

Takeaway: *Legal parties generally accept AI to search and summarize legal discovery, and federal courts have affirmatively endorsed computer-assisted document review as potentially more reliable, or at least functionally equivalent, to manual review. Argent applies that same peer-reviewed methodology as a professional forensic instrument, within a regulatory and ethical framework that the relevant scientific community is actively developing.*

IV. Admissibility Under the Frye General Acceptance Standard

In jurisdictions applying the *Frye* standard, the sole inquiry is whether the expert's methodology has gained "general acceptance" within the relevant field. *Frye v. United States*, 293 F. 1013, 1014 (DC Cir. 1923).

Under *Frye*, it might be argued that AI-assisted record review has not yet reached consensus acceptance among forensic psychologists. This argument commits a foundational analytical error: it conflates the **scientific methodology** with the **instrument of review**.

The “scientific methodology” being presented in this case is **forensic record review and clinical synthesis** – a technique that has enjoyed universal acceptance for nearly a century. Argent is simply a modern, highly efficient instrument used to execute that accepted methodology. The relevant scientific community has not adopted a consensus position that forensic psychologists may only review records by hand. To the contrary, the APA and SGFP affirmatively contemplate the use of technological tools in clinical and forensic practice.

Furthermore, forensic and clinical experts routinely rely on sophisticated, software-based instruments to organize, score, and analyze raw patient data. For example, clinicians routinely utilized computerized scoring assistants for complex diagnostic instruments – such as the MMPI or WAIS – to calculate scores and generate initial interpretive hypotheses without these software tools replacing the expert’s clinical judgement. While it is acknowledged that those systems rely on deterministic and repeatable algorithms, rather than the probabilistic engines used in generative AI models, most clinicians would not be able to meaningfully describe, explain, or certainly understand that distinction because it is largely irrelevant to the way these tools are utilized. In all of these cases, the software produces a result but the expert takes responsibility for the meaning, interpretation, and clinical decision making that results.

There are other parallels. Similarly, in the medical domain, Computer-Aided Diagnosis (CAD) software is widely accepted to assist radiologists in identifying patterns in raw MRI or CT scans. Just as Argent surfaces patterns or points of potential focus, the clinicians using CAD technology use the software as a starting point, not a final determination, diagnosis, or conclusions. And so just as a clinician’s testimony is not excluded because they used standard software instruments to score psychological tests or flag radiological patterns, a forensic expert should not be excluded for utilizing a secure, page-level cited semantic processing instrument to search and organize raw collateral record history.

Just as a court does not require a full *Frye* hearing when an expert switches from handwriting notes on a legal pad to typing them into Microsoft Word, the court should not exclude testimony because an expert utilized a more traceable and comprehensive semantic search engine to organize the collateral records. The instrument changed. The methodology, grounded in rigorous, expert-verified forensic review, did not.

Because the expert independently verifies every document and remains the sole author of the forensic opinions, the underlying methodology remains grounded in tradition, generally accepted, and admissible.

V. Distinguishing Adverse Authority

State of Washington v. Puloka, Superior Court of Washington for King County (March 29, 2024), is the most frequently cited judicial decision involving AI evidence exclusion and warrants direct address. In *Puloka*, the court excluded an AI-enhanced video under the *Frye* standard, finding that the AI tools used (commercial video editing software) had not been peer-reviewed by the forensic video analysis community, were not reproducible by that community, and employed opaque methods that represented "what the AI model thought should be shown" rather than what actually occurred.

The *Puloka* court's concerns are instructive precisely because Argent Forensics is designed to address every one of them. The court in *Puloka* was troubled by three characteristics of the AI tool at issue:

- Opacity. The AI altered source material without a traceable audit trail;
- Non-reproducibility. The outputs could not be independently verified
- Substitution. The AI generated content to replace, rather than retrieve from, the original source.

Argent operates on opposite principles. It does not alter, enhance, or reconstruct source material. It retrieves information from original, unmodified records. Every output is linked to the specific source document and page number from which it was drawn — making every extraction independently verifiable by any party who has access to the underlying records. The expert does not present AI-generated content as a substitute for primary evidence. The expert presents verified extractions from primary evidence, organized by a tool that leaves a complete and transparent audit trail.

Puloka does not stand for the proposition that AI tools are categorically inadmissible in legal proceedings. It stands for the proposition that AI tools that alter source material through opaque, non-reproducible processes fail reliability and general acceptance standards. Argent's methodology is designed to satisfy precisely the standards that *Puloka* found wanting.

VI. Conclusion

Dr. [Expert]'s use of Argent Forensics does not represent a novel, untestable scientific experiment. It is not a "black box" algorithm, or an unconstrained publicly available AI chatbot. It represents a highly disciplined, technologically assisted application of traditional forensic record review – grounded in peer-reviewed AI architecture, governed by mandatory citation controls, subject to expert verification at every step, and consistent with the APA Ethics Code, the Specialty Guidelines for Forensic Psychology, and the trajectory of professional and judicial acceptance of computer-assisted document review.

Because the platform enforces a transparent, page-cited audit trail, and because the expert manually verified every fact before forming their clinical opinion, the methodology is highly reliable, thoroughly testable, and legally defensible.

For the foregoing reasons, Dr. [Expert]'s expert testimony should be ruled fully admissible.

Dated: _____

Respectfully submitted, _____

[Retaining Attorney]

Counsel for [Plaintiff/Defendant]

AUTHORITIES CITED

Cases

- Da Silva Moore v. Publicis Groupe, 287 F.R.D. 182 (S.D.N.Y. 2012)
- Daubert v. Merrell Dow Pharmaceuticals, Inc., 509 U.S. 579 (1993)
- Equity Analytics, LLC v. Lundy, 248 F.R.D. 54 (D.D.C. 2008)
- Frye v. United States, 293 F. 1013 (D.C. Cir. 1923)
- General Electric Co. v. Joiner, 522 U.S. 136 (1997)
- Kumho Tire Co. v. Carmichael, 526 U.S. 137 (1999)
- Mata v. Avianca, 678 F. Supp. 3d 443 (S.D.N.Y. 2023)
- United States v. O'Keefe, 537 F. Supp. 2d 14 (D.D.C. 2008)
- State v. Puloka, No. 21-1-04851-2 KNT, (Wash. Super. Ct. King Cnty. Mar. 29, 2024)

Rules

- Fed. R. Evid 702

Professional Standards

- American Psychological Association, APA Ethics Code (2017)
- American Psychological Association, Office of Health Care Innovation, *Companion Checklist: Evaluation of an AI-Enabled Clinical or Administrative Tool* (October 2024)
- American Psychological Association, *Specialty Guidelines for Forensic Psychology* (2013)
- American Board of Forensic Psychology, *Candidate Manual; AI Policy* (2026)

Scientific Literature

- Guu, K. et al., *REALM: Retrieval-Augmented Language Model Pre-Training*, arXiv:2002.08909 (2020)
- Izacard, G. & Grave, E., *Leveraging Passage Retrieval with Generative Models for Open Domain Question Answering*, arXiv:2007.01282 (2021)
- Lewis, P. et al., *Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks*, arXiv:2005.11401 (2020)
- Sweller, J., *Cognitive Load During Problem Solving: Effects on Learning*, 12 Cognitive Science 257 (1988)
- Rilen, S. *Practical Frameworks for Integrating Artificial Intelligence in Forensic Psychology*, Journal of Forensic Psychology Research and Practice (2026)