

Vendor Evaluation Matrix

E-DISCOVERY & AGENTIC REVIEW • 2026 STRATEGIC FRAMEWORK

The Shift to 'Agentic' Discovery

By Q1 2026, the E-Discovery landscape has moved beyond Technology Assisted Review (TAR 1.0/2.0) into **Agentic AI Review**. The primary challenge is no longer just volume—it is the complexity of conversational data (Slack Huddles, Microsoft 365 Co-pilot prompts) and the defensibility of generative outputs.

This evaluation matrix prioritises vendors capable of handling multi-modal data streams and providing **ISO 42001** compliant audit trails for AI decisions.

Key 2026 Drivers

- Microsoft 365 Co-pilot Log ingestion
- Hallucination Risk Scoring
- Outcome-based Pricing (vs GB)

Asset 01: The Evaluation Matrix

SCORE 1 (POOR) TO 5 (EXCELLENT)

CRITERIA CATEGORY	THE "LEGACY" STANDARD	THE "ALPODO" 2026 STANDARD
Modern Data Intake Handling conversational & AI-generated logs	Supports Email (SMTP), standard Chat strings, and basic attachments. Struggles with threaded "huddles".	Native parsing of MS Teams Loop components, Zoom Summaries, and Copilot prompt history. Preserves context/threading.
AI Autonomy Level Degree of human intervention required	TAR 2.0 (CAL): Ranks documents. Humans must review the top tier to train the model continuously.	Agentic (Level 4): Autonomous First Pass Review. Agents code for responsiveness and privilege; Humans audit the logic, not the docs.
Defensibility & Audit Judicial acceptance of the process	Statistical validation (Recall/Precision curves). "Black box" regarding why a specific doc was coded Hot.	Reasoning Logs: AI generates a "Memo" for every privilege redaction explaining <i>why</i> it applies (e.g., "Legal Advice involved").

Security & Privacy Model training and data residency	Data encryption at rest. SOC 2 Type II compliance.	Zero-Retention Inference: No data is used to train base models. ISO 42001 (AI Management) certified.
--	--	---

Asset 02: Implementation Roadmap

DURATION: 10 WEEKS

WEEKS 1-2

Data Cartography & Scope

Map 2026-era data sources (Slack Huddles, Zoom AI Summaries, MS Teams Co-pilot logs) to define technical requirements. Establish 'outcome-based' success metrics (e.g., "Review Accuracy per Dollar") rather than volume-based metrics.

WEEKS 3-4

The 'Agentic' Market Scan

Shortlist vendors offering 'Agentic AI' (autonomous review agents) versus legacy TAR. Issue RFP focusing specifically on GenAI hallucination rates, citation accuracy, and ISO 42001 alignment.

WEEKS 5-6

Bias & Security Stress Test

Conduct a 'Red Team' exercise on vendor LLMs to test for privilege leakage and algorithmic bias. Validate compliance with EU AI Act (2025/2026 enforcement phase).

WEEKS 7-8

The 'Shadow' Pilot

Run the top two contenders in parallel with your current workflow on a closed matter (The "Shadow" Pilot). Compare 'Time-to-Fact' and 'Review Accuracy' directly against human-only benchmarks.

WEEK 9

Commercial Negotiation

Pivot away from pure GB/hosting fees. Negotiate 'Capped Review' or 'Outcome-Based' pricing to prevent GenAI token cost spiralling during heavy litigation.

WEEK 10

Executive Ratification

Present final ROI deck. Highlight the 'Efficiency Dividend' (speed) and 'Risk Shield' (consistency) to General Counsel and IT Security for final sign-off.

Anticipating Objections: The "Black Box" Defence

OBJECTION

"How do we prove to a judge that the AI didn't hallucinate a document into relevance?"

THE ALPODO REBUTTAL

We do not rely on the AI's memory; we rely on its reasoning. The selected vendor provides **"Citation-Backed Coding"**, where every decision is hyperlinked to a specific text snippet in the source file. If the AI cannot cite the text, it cannot code the document.

EVALUATED BY

DATE

APPROVED BY (GC)