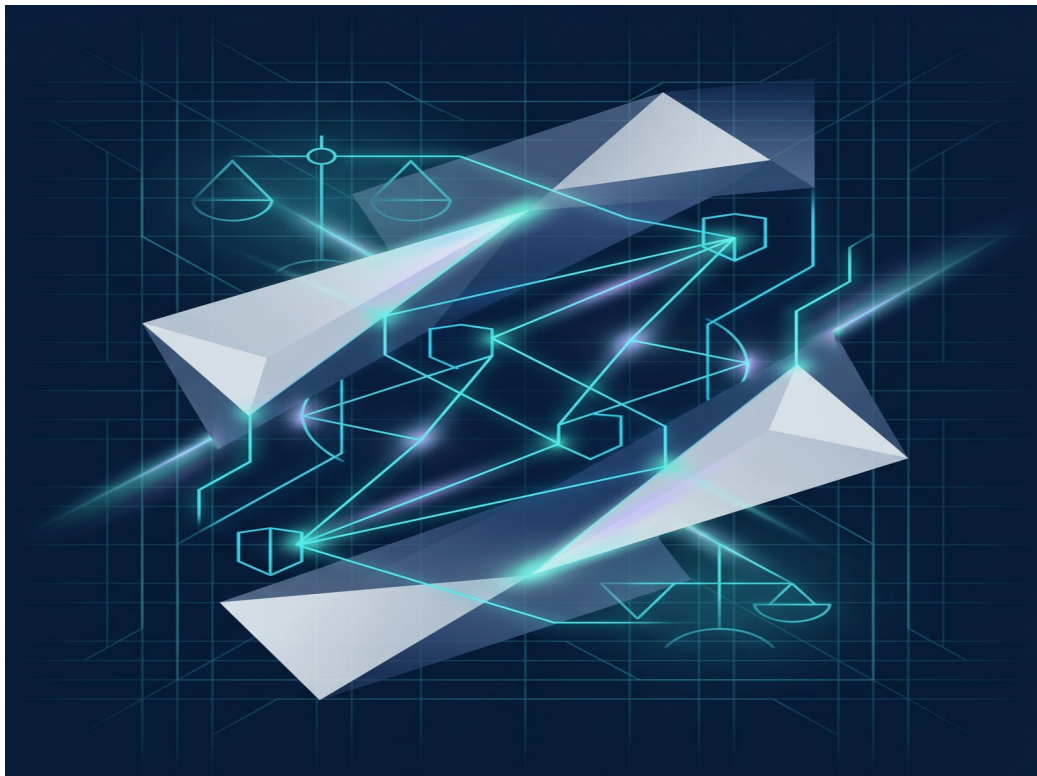


AI PLATFORM DECISION GUIDE · 2026 EDITION

The Legal LLM Buyer's Guide

A vendor-neutral 2026 comparison of Claude for Legal, ChatGPT Enterprise / Harvey, and Perplexity for Legal.



PREPARED BY THE SECUREJUSTICEAI CONSORTIUM

For SMB law firms and NGOs. Specialized service for women-led law firms and personal injury practices.

July 2026 Edition

00

00 — Contents

01 — Executive Summary	3
02 — The Landscape: Why This Decision Matters Now	4
03 — The Architectural Divide: RAG vs. LLM-Native	5
04 — Full Comparison Table	6
05 — Deep Dive: Claude for Legal	7
06 — Deep Dive: ChatGPT Enterprise / Harvey	9
07 — Deep Dive: Perplexity for Legal	11
08 — Decision Framework	13
09 — Governance Guardrails Every Firm Needs	15
10 — Sources & Methodology	16
11 — About SecureJusticeAI	17
12 — Legal Disclaimer	18

01 — Executive summary

The generative AI decision facing your firm today is not *whether* to adopt AI. It is ***which platform to trust with client matter, how to govern it, and how to protect the firm from the failure modes that have already produced 667+ documented attorney sanctions in US courts through July 2026.*** This guide gives you the vendor-neutral framework to make that decision.

Bottom line up front

- **All three leading platforms hallucinate.** Stanford RegLab benchmarks (2024) documented hallucination rates of 17–37% across leading legal LLMs. No AI-generated citation, quote, or authority should ever be filed without independent human verification.
- **The right platform depends on the work.** Claude excels at agentic reasoning and multi-step workflows. ChatGPT/Harvey dominates transactional drafting at AmLaw scale. Perplexity's RAG-first architecture makes it the only platform designed for citation-linked research.
- **The platform is only 30% of the decision.** Governance guardrails, deployment configuration, security hardening, and attorney training determine 70% of realized ROI. This is where firms without a governance partner leave value on the table — and where they accumulate compliance risk.

Why this guide, why now

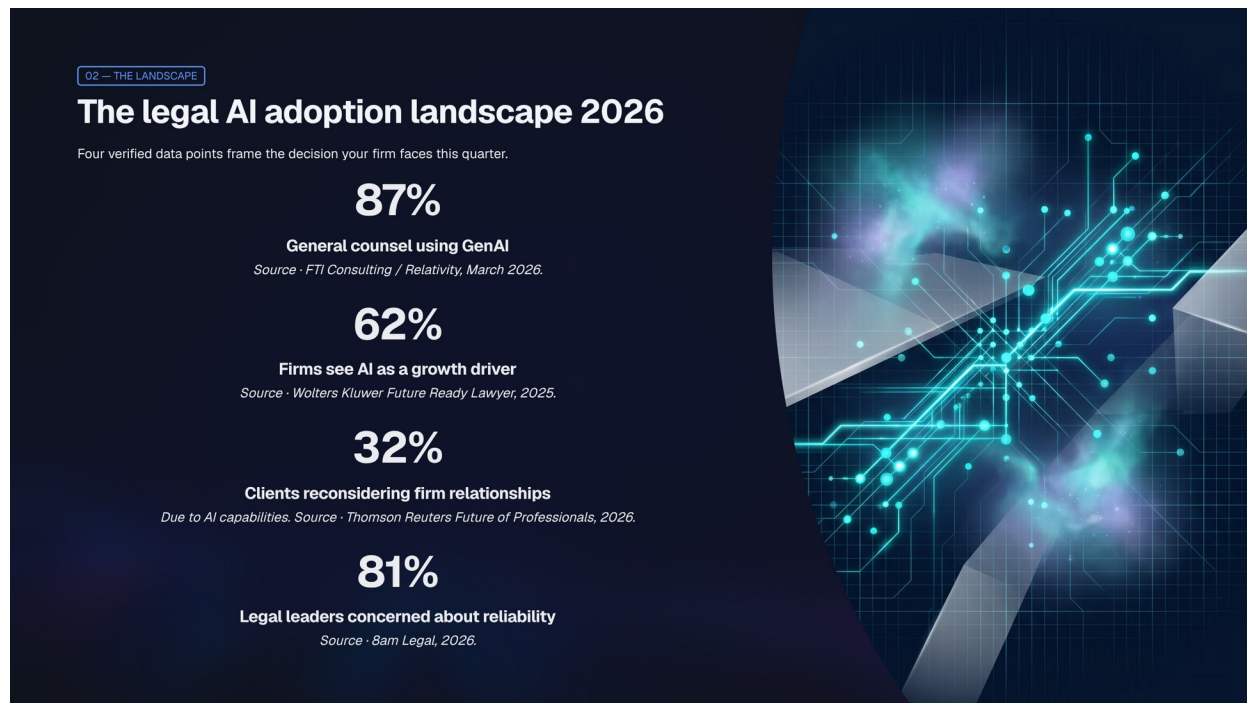
87% of US general counsel now report using generative AI in their day-to-day work (FTI Consulting / Relativity, March 2026). Simultaneously, 32% of legal clients are actively reconsidering firm relationships based on their counsel's AI capabilities (Thomson Reuters Future of Professionals, 2026). The gap between AI-forward firms and their competitors is compounding every quarter.

The stakes for choosing wrong — or delaying — are asymmetric. On the downside: sanctions, malpractice exposure, and client attrition. On the upside: dramatic operational leverage, faster client work, and pricing power. This guide is written to help managing partners at small and mid-sized law firms make this decision with clarity and without vendor spin.

Who this guide is written for

Managing Partners at SMB law firms (2–50+ attorneys). In-house General Counsel at mid-market companies (\$10M–\$150M). Executive Directors and technology leads at legal services organizations and NGOs. Specialized guidance for women-led law firms and personal injury practices.

02 — The landscape



The AI adoption gap is widening

Four verified data points from 2025–2026 tell the story:

- **87% of general counsel now use GenAI** (FTI Consulting / Relativity, March 2026). Your clients are already comfortable with AI-driven counsel.
- **62% of firms see AI as a growth driver** (Wolters Kluwer Future Ready Lawyer, 2025). Firms that have adopted AI report higher revenue growth than firms that have not.
- **32% of clients are reconsidering firm relationships** based on their firm's AI capabilities (Thomson Reuters Future of Professionals, 2026). This is a client-driven pressure, not a vendor-driven one.
- **81% of legal leaders are concerned about AI reliability** (8am Legal, 2026). Concern is universal. What distinguishes AI-forward firms is that they concern themselves with reliability while adopting, not instead of adopting.

The downside risk is documented

The historical record on AI failure in legal practice is no longer theoretical. Damien Charlotin's AI Hallucination Cases database, refreshed monthly and current as of July 2026, records:

- **667+ US attorney sanctions cases** arising from AI-generated fabrications in filed briefs.
- **1,187+ total US proceedings** where AI hallucination was cited as a factor.
- **1,725+ worldwide cases** in the same database.

- **Single-case sanctions up to \$110,000** (2026), with ~\$145,000 in aggregate sanctions in Q1 2026 alone.

The cyber risk compounds the AI risk

AI adoption without corresponding cybersecurity discipline is a liability multiplier. Three data points frame the environment your platform decision sits within:

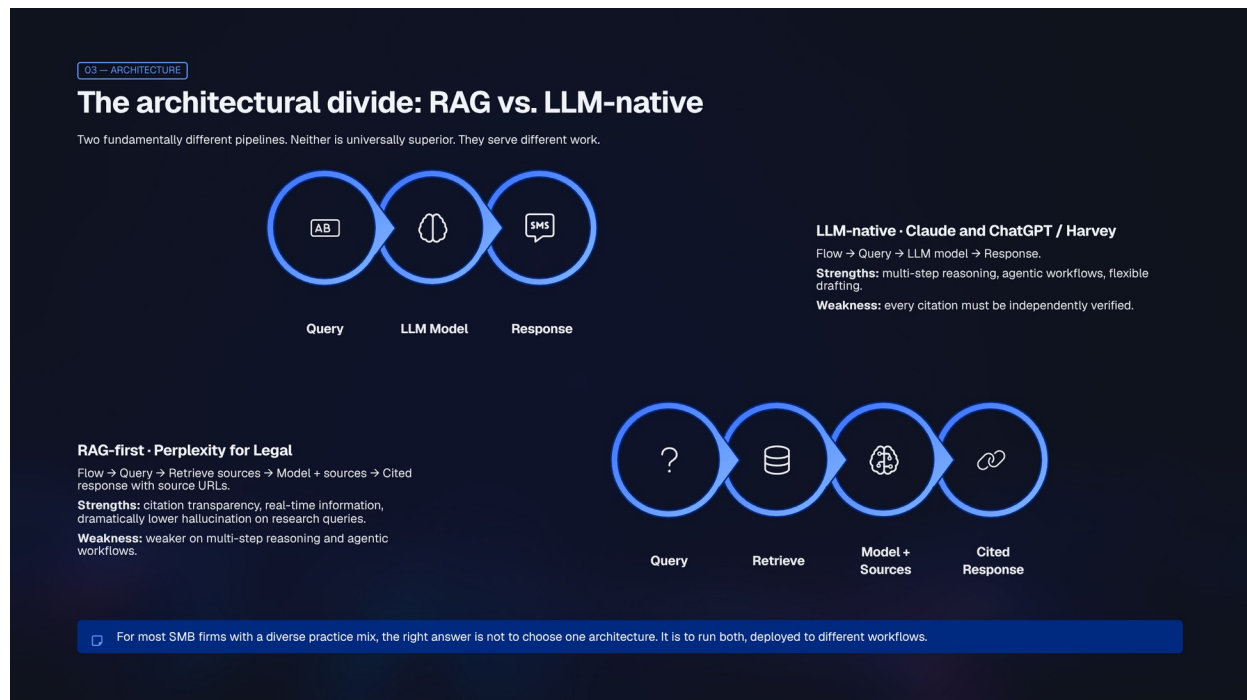
- **29% of law firms were hit by a cyber attack in 2024** (ABA TechReport, 2024). Nearly one in three.
- **\$5.08M average ransomware attack cost** per legal firm (IBM Cost of a Data Breach, 2025).
- **54% increase in ransomware attacks targeting legal organizations year-over-year** (Comparitech, January 2026). 225 attacks in 2024 grew to 346 in 2025.

Why vendor neutrality matters

Every LLM vendor has a reseller channel. Every reseller earns commissions. When a reseller recommends a platform, that recommendation is compromised by the incentive to sell what pays them, not what serves your practice.

SecureJusticeAI holds no partner incentive with Anthropic, OpenAI, or Perplexity. Our LLM Enablement service exists to help your firm succeed with whichever platform best fits your practice, matter mix, and budget. This guide is written from that same neutrality.

03 — The architectural divide



Before comparing platforms feature-by-feature, understand the single most important architectural distinction in this decision: RAG-first vs. LLM-native.

LLM-native (Claude and ChatGPT/Harvey)

Large language models generate responses from parametric knowledge — the patterns learned during training. When you ask a legal question, an LLM synthesizes an answer from that internal representation. This produces:

- **Strengths:** Flexible reasoning, strong multi-step logic, natural writing, ability to synthesize across many concepts, agentic workflows.
- **Weaknesses:** Hallucinations when the LLM lacks or misremembers specific facts. Citations must be verified. Cutoff dates matter.

RAG-first (Perplexity for Legal)

Retrieval-Augmented Generation is a fundamentally different pipeline: fetch source documents first, then generate a response grounded in the retrieved text, with every claim linked back to its source URL. This produces:

- **Strengths:** Citation transparency, real-time information, source verifiability, dramatically lower hallucination on research-style queries.
- **Weaknesses:** Only as good as retrieval quality. Weaker on complex multi-step reasoning. Limited agentic capabilities. Long-form drafting is not the primary use case.

Why this matters for your practice

These two architectures are not interchangeable, and neither is universally superior. They serve different work:

- **For legal research with cited authorities:** RAG-first (Perplexity) is the natural fit.
- **For document drafting, contract review, and multi-step reasoning:** LLM-native (Claude or ChatGPT/Harvey) is superior.
- **For agentic workflows (intake automation, matter triage, cross-file synthesis):** LLM-native is required; RAG is insufficient.

A practical implication most firms miss

SMB firms with diverse practice mixes often benefit from using more than one platform — one for drafting and reasoning, one for citation-linked research. The cost of running two enterprise LLM subscriptions is frequently less than the productivity cost of forcing every workflow through a single tool.

04 — Full comparison table

04 — PLATFORMS

The three enterprise legal LLM platforms at a glance

Vendor-neutral positioning. Same neutrality you should expect from your consortium partner.

 Claude for Legal · Anthropic Strength: reasoning + agentic. Best general-purpose foundation for SMB firms. Safety-first architecture, MCP tool ecosystem. Best fit: SMB firms, agentic workflows, ethics-heavy practices.	 ChatGPT / Harvey · OpenAI + Harvey Strength: drafting + breadth. AmLaw default for transactional work. Priced for Big Law budgets. Best fit: AmLaw 100 and 200, transactional practices, Microsoft-native firms.	 Perplexity for Legal · Perplexity Strength: research + citations. RAG-first architecture — the only platform designed for source-linked research. Best fit: research-heavy practices, solo firms, citation-critical work.
--	---	--

A 15-factor side-by-side comparison as of July 2026. Platform features change frequently; verify current specifications with vendors before contracting.

Factor	Claude for Legal	ChatGPT / Harvey	Perplexity for Legal
Model architecture	LLM-native (Claude 4)	LLM-native (GPT-4/5)	RAG-first (retrieval + LLM)
Context window	200K tokens	128K–1M tokens	~32K tokens
Primary strength	Reasoning, agentic	Drafting, breadth	Research, citations
Citation model	Generative (verify)	Generative (verify)	Source-linked URLs
No-training guarantee	Yes (Enterprise/Team)	Yes (Enterprise/Team)	Yes (Enterprise)
SOC 2 Type II	Yes	Yes	Yes
ISO 27001	Yes	Yes	Partial
HIPAA (BAA)	Yes	Yes	Limited
Enterprise SSO / SCIM	Yes	Yes	Yes
Data residency	US, EU	US, EU, custom	US primary

DMS integrations	Via MCP + partners	Broad (via Harvey)	Limited native
State bar posture	Widely accepted	Widely accepted	Widely accepted
Enterprise pricing	\$\$	\$\$\$-\$\$\$\$	\$
Best-fit firm	SMB, agentic-heavy	AmLaw, drafting-heavy	Research-heavy

PRICING TIERS ARE RELATIVE: \$ = UNDER \$30/USER/MO · \$\$ = \$30-60 · \$\$\$ = \$60-150 · \$\$\$\$ = \$150+ BESPOKE

05 — Claude for Legal (Anthropic)

Positioning

Anthropic's Claude occupies the "safety-first, reasoning-heavy" position in the market. The Claude 4 family (200K context window as of July 2026) is engineered around Constitutional AI principles, which translate into practical governance advantages for legal work: clearer refusal behavior on out-of-bounds requests, more reliable adherence to system prompts, and stronger safety guardrails against prompt injection.

Where Claude wins

- **Multi-step reasoning and agentic workflows.** Claude is the strongest of the three for tasks that require sequencing (intake screening → conflict check → client letter draft, all in one workflow).
- **Long-context matter review.** 200K tokens equals roughly 500 pages of case material analyzed in a single prompt. Useful for discovery review, deposition analysis, and matter summarization.
- **Ethics-heavy practices.** Family law, immigration, and personal injury firms handling sensitive client information benefit from Claude's more conservative default posture.
- **MCP (Model Context Protocol) ecosystem.** Claude leads on standardized tool integration, which shortens custom-agent deployment timelines.
- **Pricing efficiency for SMB firms.** Claude Team and Enterprise tiers deliver strong value in the \$25–65/user/month range, aligning with SMB budgets.

Where Claude struggles

- **No native real-time retrieval.** Unlike Perplexity, Claude does not include a citation-linked research pipeline out of the box. Add Perplexity or wire retrieval via MCP for research-heavy work.
- **Smaller partner ecosystem than OpenAI.** If your firm is deeply invested in Microsoft 365 with Copilot, or in the Harvey ecosystem, Claude requires more custom integration work.
- **Hallucinates.** Like every LLM. Stanford RegLab benchmarks put Claude in the 17–37% hallucination range on legal-specific queries. No output goes to court unverified.

Best-fit firm profiles

- **2-25 attorney litigation, PI, or family-law firms** with limited IT budget and high sensitivity to reliability.
- **Firms building custom agentic workflows** (intake automation, matter triage, discovery review, TrialLift-style citation validation).
- **Firms with strong governance culture** where a safer-default LLM reduces training overhead.
- **Legal aid and mission-driven organizations** where pricing efficiency and ethics-first design align with mission.

The verdict on Claude

For most SMB law firms with diverse practices, Claude for Legal is the strongest general-purpose foundation. Pair it with Perplexity for research and you cover roughly 90% of practice-area needs at a price point that fits SMB budgets.

06 — ChatGPT Enterprise / Harvey

Positioning

OpenAI's ChatGPT Enterprise is the market's most broadly-adopted general-purpose LLM. In legal, its most consequential expression is through Harvey, a fine-tuned OpenAI-based platform that has become the default choice for AmLaw 100 and AmLaw 200 firms. Together, they represent the "breadth and drafting" position.

Where ChatGPT/Harvey wins

- **Transactional drafting at scale.** Contract generation, redlining, clause-library management. Harvey's legal-specific tuning makes it the strongest of the three for high-volume transactional work.
- **Harvey ecosystem.** Deep integration with Big Law workflows, established relationships with major DMS providers (iManage, NetDocs), and enterprise-grade support.
- **Extremely large context in top tier.** GPT-5-family tiers exceed 1M tokens, useful for full-matter review or M&A due diligence at scale.
- **Broad third-party tool integration.** OpenAI's ecosystem is the biggest in the industry — more plugins, more connectors, more community-built solutions.

Where ChatGPT/Harvey struggles

- **Cost.** Harvey Enterprise is priced for AmLaw budgets. Firms under \$10M in revenue often find the ROI difficult to justify, especially when Claude Team or Enterprise delivers 80% of the benefit at 30% of the cost.
- **Safety guardrails.** OpenAI's default posture is less conservative than Claude's. This is neither good nor bad in absolute terms, but for ethics-heavy practices it requires more prompt-engineering discipline.
- **Hallucinates.** Same 17–37% Stanford RegLab range as Claude. Verification is non-negotiable.
- **Vendor concentration.** Betting the firm on OpenAI creates concentration risk. Recent regulatory scrutiny of the largest AI vendors amplifies that risk.

Best-fit firm profiles

- **AmLaw 100 / AmLaw 200 firms** with the budget for Harvey and the scale that justifies it.
- **Transactional practices** with high volumes of contract work, M&A, or corporate diligence.
- **Firms already deeply invested in Microsoft 365 / Azure** where the OpenAI partnership makes integration frictionless.
- **Firms with dedicated in-house AI teams** that can manage the ecosystem breadth productively.

The verdict on ChatGPT / Harvey

Harvey is the correct default for AmLaw-scale transactional practices with the budget to support it. For SMB firms, the price/value equation is harder to justify unless the practice mix is heavily transactional and the firm is already Microsoft-native. Most SMB firms will find better ROI with Claude — with or without Perplexity added for research.

07 — Perplexity for Legal

Positioning

Perplexity Enterprise Pro is architecturally distinct from Claude and ChatGPT: it is a Retrieval-Augmented Generation (RAG) system built for real-time, citation-linked research. Every response is grounded in retrieved sources, and every claim carries a link back to the source URL. For research-heavy practices, this changes the fundamental value proposition.

Where Perplexity wins

- **Citation transparency.** This is the flagship feature. Every claim in a Perplexity response is source-linked. For legal research, this dramatically reduces hallucination risk on citation-heavy queries.
- **Real-time information.** Unlike LLMs constrained by training-data cutoffs, Perplexity fetches current information. Useful for regulatory research, recent case law, and up-to-date market intelligence.
- **Ease of adoption.** Attorneys who are skeptical of "black box" AI often prefer Perplexity because the source visibility makes verification native to the workflow.
- **Pricing.** Enterprise Pro pricing is the lowest of the three platforms, making it accessible even for solo practitioners and small firms.

Where Perplexity struggles

- **Long-form drafting.** Perplexity is not designed to draft a 40-page motion. It is designed to help you research the authorities that support the motion you are drafting elsewhere.
- **Agentic workflows.** The RAG-first architecture is a poor fit for multi-step reasoning tasks. Automation, matter triage, and cross-file synthesis require an LLM-native platform.
- **EU data residency.** Perplexity's data residency options are more limited than Claude's or ChatGPT's. Firms with EU clients should verify DPA availability before contracting.
- **DMS integration depth.** Native integrations with iManage, NetDocs, and other legal DMS platforms are less mature than Harvey's.
- **Still hallucinates.** RAG dramatically reduces hallucination risk on research queries, but does not eliminate it. Retrieval failures and misinterpretation of sources still produce errors. Verify.

Best-fit firm profiles

- **Research-heavy practices** (regulatory, tax, IP, appellate litigation, complex commercial).
- **Firms that want an AI research assistant without changing their drafting workflow.**
- **Solo practitioners and 2-5 attorney firms** where pricing efficiency is decisive.

- **Firms with skeptical or verification-first culture** where source visibility overcomes AI adoption resistance.

The verdict on Perplexity

Perplexity is not a replacement for Claude or ChatGPT for firms doing drafting-heavy or agentic work. It is a specialized research complement. The strongest positioning for most SMB firms: pair it with Claude, use Perplexity for citation-linked research, use Claude for drafting and agentic workflows.

08 — Decision framework

06 — RECOMMENDATIONS

Which platform for which firm

Vendor-neutral recommendations by firm profile. Final calls depend on your practice mix.

Solo or 2-attorney PI firm Claude Team + Perplexity Best value for reasoning plus citation research under \$50 per user per month.	5–15 attorney litigation firm Claude Enterprise + Perplexity Agentic workflows plus research, both under \$30K per year.	Mid-market in-house GC Claude Enterprise Best cross-team governance and agentic support.
AmLaw 100 / 200 ChatGPT Enterprise + Harvey Ecosystem depth, drafting scale, existing DMS relationships.	Research-heavy specialty Perplexity Pro + Claude for narrative RAG citations plus reasoning synthesis.	Women-led litigation boutique Claude Enterprise + Perplexity Best reasoning plus citation stack at SMB pricing.


SECUREJUSTICEAI — TRUSTED AI · REAL JUSTICE · SECURE BY DESIGN · July 2026 Edition

By firm profile

The following matrix gives our vendor-neutral recommendation by common SMB firm profiles. Recommendations assume budget-appropriate tiers; final calls depend on your practice mix.

Firm profile	Primary recommendation	Why
Solo or 2-attorney PI firm	Claude Team + Perplexity	Best value for reasoning + citation research under \$50/user/month
5–15 attorney litigation firm	Claude Enterprise + Perplexity	Agentic workflows + research; both under \$30K/yr
15–50 attorney general practice	Claude Enterprise + Perplexity	Cross-team governance + research at SMB pricing
Mid-market In-House GC	Claude Enterprise	Best cross-team governance and agentic support
Legal aid / NGO	Claude Team or ChatGPT Team	Cost efficiency + strong safety guardrails
AmLaw 100 / 200	ChatGPT Enterprise + Harvey	Ecosystem depth, drafting scale, existing relationships

Research-heavy specialty	Perplexity Pro + Claude for narrative	RAG citations + reasoning synthesis
Women-led litigation boutique	Claude Enterprise + Perplexity	Best reasoning + citation stack at SMB pricing

By practice area

- **Personal Injury:** Claude (intake automation + client communications + demand letters). Add Perplexity for medical literature and jurisdictional research.
- **Transactional / Corporate:** ChatGPT / Harvey if AmLaw budget available; Claude Enterprise for SMB firms.
- **Litigation:** Claude for reasoning and agentic workflows + Perplexity for citation-linked research. This pairing dominates SMB litigation.
- **Immigration:** Claude — safety-critical, ethics-heavy, sensitive client information.
- **Intellectual Property:** Perplexity for prior-art research + Claude for patent narrative and drafting.
- **Regulatory:** Perplexity as primary (source-linked responses are essential); Claude as secondary for synthesis.
- **Family Law:** Claude — sensitivity and safety guardrails matter more than raw capability.
- **Estate Planning:** Claude Enterprise for drafting + Perplexity for jurisdictional variation research.

The six-question framework

If you can answer these six questions honestly, your platform decision follows naturally:

- **1. What is my primary use case?** Drafting, research, agentic workflow, or all three?
- **2. What is my budget per attorney per year?** Under \$2K/attorney/yr suggests Claude Team or Perplexity Pro. \$2K–\$8K suggests Claude Enterprise. Above \$8K opens ChatGPT/Harvey.
- **3. How ethics-sensitive is my practice?** If sensitivity is high, Claude's conservative default posture is worth the tradeoff.
- **4. Do I need EU data residency?** If yes, ChatGPT or Claude Enterprise are stronger choices than Perplexity today.
- **5. Am I already committed to a Microsoft, Google, or OpenAI ecosystem?** If deeply Microsoft-native, ChatGPT/Harvey integration is easier. Otherwise the ecosystem question is not decisive.
- **6. What is my tolerance for hallucination risk on filed work?** Every firm should answer "zero." That means verification workflows and governance, regardless of platform.

09 — Governance guardrails

05 — GOVERNANCE

The four non-negotiable governance guardrails

Regardless of platform, these controls separate AI as a productivity multiplier from AI as a malpractice risk.

Access controls

SSO with MFA on every account. Role-based access scoped to the matter team. No-training on customer data, verified contractually. Enterprise tiers only.

Data loss prevention

Client data classified before it moves. Never paste privileged content into consumer LLM interfaces. Enterprise-only for anything with PII, including practice-management data.

Human verification layer

No AI-generated citation is filed without independent verification. Prompt libraries carry governance guardrails. Attorney training on hallucination recognition, tied to actual matter work.

Audit trail

Every prompt and output logged with matter number and attorney identity. Retention aligned with client obligations. Discoverability protocols documented.

Frameworks: ISO/IEC 42001 · NIST AI RMF · ABA Formal Opinion 512 · SOC 2 Type II.

Regardless of which platform you choose, these controls are the difference between AI as a productivity multiplier and AI as a malpractice risk. They are non-negotiable in modern practice.

Access controls

- **SSO with MFA** — on every account. No exceptions.
- **Role-based access** — partners, associates, staff each get scoped permissions. Client-privileged matter is not accessible outside the matter team.
- **No-training on customer data** — verified contractually. Enterprise tiers only. Never use free consumer tiers on client matter.

Data loss prevention

- **Client data classification** — privileged, confidential, work-product, and public tiers. Attorneys must know the tier before pasting.
- **Never paste privileged content into consumer LLM interfaces** — firm policy, enforced through DLP tooling.
- **Enterprise-only for anything with PII** — no exceptions, including practice-management data.

Human verification layer

- **No AI citation is filed without independent verification** — preferably by a second attorney, always by a human.

- **Prompt libraries with governance guardrails baked in** — not generic templates.
- **Attorney training on hallucination recognition** — not a one-hour Zoom demo. Recurring, practice-specific, and tied to actual matter work.

Audit trail

- **All prompts and outputs logged** — with the matter number and attorney identity attached.
- **Retention aligned with client obligations** — including hold protocols for AI outputs on active matters.
- **Discoverability protocols** — the firm must know what to produce and what to withhold on AI-related discovery.

ABA Formal Opinion 512 compliance

The ABA's July 2024 Formal Opinion 512 clarifies attorney obligations when using generative AI. Three duties are load-bearing:

- **Attorney competence (Rule 1.1)** requires understanding of the AI tool being used, its limits, and its failure modes.
- **Client confidentiality (Rule 1.6)** requires no-training guarantees, appropriate DLP, and client disclosure where circumstances warrant.
- **Supervisory duties (Rules 5.1 and 5.3)** extend to associates' and staff's use of AI. Firms must have written policies and training programs.

Where SecureJusticeAI adds value

The four-pillar consortium delivers AI Strategy, AI Governance, Zero Trust Cybersecurity, and IT Integration & AI Solutions in a single 90-day engagement. Our Enterprise LLM Enablement service specifically deploys, configures, and governs whichever platform (Claude, ChatGPT/Harvey, or Perplexity) best serves your practice. Vendor-neutral by design.

10 — Sources & methodology

Primary sources

- Anthropic — Claude for Enterprise / Claude for Legal documentation (anthropic.com)
- OpenAI — ChatGPT Enterprise documentation (openai.com)
- Harvey AI — Public product documentation and press releases (harvey.ai)
- Perplexity — Enterprise Pro documentation (perplexity.ai)

Benchmarks and empirical studies

- Stanford RegLab (2024) — Legal LLM hallucination benchmarks, 17–37% range across leading platforms
- Damien Charlotin AI Hallucination Cases database (July 2026) — 667+ US attorney sanctions cases; 1,187+ US proceedings; 1,725+ worldwide
- IBM — Cost of a Data Breach Report (2025) — \$5.08M average legal firm ransomware cost

Legal industry surveys

- FTI Consulting / Relativity — General Counsel AI Adoption Survey (March 2026) — 87% GC adoption
- Thomson Reuters — Future of Professionals Report (2026) — 32% client reconsideration
- Clio — Legal Trends Report (2026)
- ABA — TechReport (2024, 2025) — 29% law firm cyber attack rate (2024)
- ILTA — Legal Technology Survey (2025)
- Bloomberg Law — Legal Trends (2025)
- Wolters Kluwer — Future Ready Lawyer Report (2025) — 62% AI as growth driver
- 8am Legal — Leadership Survey (2026) — 81% AI reliability concern
- Comparitech — Legal Ransomware Report (January 2026) — 54% YoY increase, 225→346 attacks

Regulatory and ethics

- American Bar Association — Formal Opinion 512 (July 2024) on GenAI use by lawyers
- State Bar Advisory Opinions — New York, California, Florida, Illinois, Texas, District of Columbia, Virginia, Washington, Oregon, New Jersey (as of July 2026)

Methodology

This guide draws on vendor documentation, third-party benchmarks, published legal industry surveys, court records via Charlotin's database, and state bar advisory opinions. Platform features, pricing, and capabilities change frequently — verify current specifications with vendors before contracting. SecureJusticeAI holds no partner incentive with Anthropic, OpenAI, Harvey, or Perplexity. Where the guide expresses a preference, it reflects independent research judgment.

11 — About SecureJusticeAI

SecureJusticeAI is a four-partner consortium built for SMB law firms and NGOs, with specialized service for women-led law firms and personal injury practices. We deliver AI strategy, governance, cybersecurity, and IT infrastructure as a single integrated 90-day deployment.

The consortium

- **Richard Schreiber** — Consortium Lead and Principal Strategist. Founder, Law Firm AI Expert and TrialLift.
- **Chris Feamster** — AI Governance & Compliance Lead. Cognify Solutions.
- **Tony Chiappetta** — Zero Trust Cybersecurity Lead. CHIPS Cyber Defense Solutions.
- **Clyde Brown** — Chief AI Strategist & Solutions Architect. TechBridge Unlimited.

Why vendor-neutral matters

Every LLM vendor has a reseller channel. Every reseller earns commissions. SecureJusticeAI holds no partner incentive with Anthropic, OpenAI, Harvey, or Perplexity. Our Enterprise LLM Enablement service exists to help your firm succeed with whichever platform best fits your practice, matter mix, and budget — not the one that pays us the most commission.

What we do

- AI Strategy & Solutions — platform selection, workflow design, custom agent architecture
- AI Governance & Compliance — policy frameworks, ABA Opinion 512 alignment, board-ready documentation
- Zero Trust Cybersecurity — prevention-first architecture, AppGuard endpoint protection, incident response
- IT Integration & AI Solutions — Azure/Microsoft 365 hardening, legacy modernization, agentic RAG pipelines
- Enterprise LLM Enablement — Claude, ChatGPT/Harvey, or Perplexity deployment, integration, and governance

Book your Clarity Audit

A structured two-week assessment of your firm's AI readiness, cybersecurity posture, and IT infrastructure. Delivers a prioritized roadmap, a partner-ready and funder-ready ROI narrative, and a 90-day deployment plan. Starting at \$2,500, fully credited toward your full engagement.

Website: securejusticeai.com

Email: richard@lawfirmaiexpert.com

12 — Legal disclaimer

This guide is provided for informational purposes only and does not constitute legal advice. Comparisons, recommendations, and any statements of platform features are based on publicly available information, vendor documentation, and independent research current as of the publication date. Platform features, pricing, and capabilities change frequently; verify current specifications with vendors before contracting or making any procurement decision.

SecureJusticeAI is not a legal service provider. Nothing in this guide should be construed as legal advice or as the establishment of an attorney-client relationship. For advice on your firm's specific obligations under state bar rules, professional responsibility standards, or client confidentiality requirements, consult qualified counsel in your jurisdiction.

No AI platform is fault-free. All firms are ethically and professionally obligated under ABA Formal Opinion 512 and analogous state bar guidance to verify AI-generated content before filing, sharing, or otherwise relying on it in professional capacity. Attorneys remain solely responsible for the accuracy and appropriateness of AI-assisted work product.

SecureJusticeAI holds no partner incentive, reseller commission, or financial relationship with Anthropic, OpenAI, Harvey AI, or Perplexity as of the publication date. All product names, logos, and brands referenced in this guide are the property of their respective owners. Use of these names does not imply endorsement.

Statistics and citations in this guide are drawn from publicly available sources listed in Section 10. Where a statistic could not be verified to primary source within acceptable precision, the closest defensible substitute has been used with the substitution flagged in the source note.

© 2026 SecureJusticeAI. All rights reserved.