

The Personal Injury AI-Readiness Study 2026

A census of 627 personal injury law firm websites,
measured for the search engines that answer instead of rank.

627

firms measured

0

A grades

69.4

average score

32.7

point AI gap

Executive summary

627

firms measured

0

A grades

98.9%

GEO below technical

45%

graded D or F

Between August 4–5, 2026, we measured the homepage of every personal injury law firm in the BrandTerritory index that operates a working website — 627 firms across 205 U.S. markets — using the same public engine, weights, and method we used to grade 18 legal SEO agencies the week before. One page per firm, the page the firm controls completely.

The market is technically excellent and rhetorically unverifiable. The average firm scores 89.8 on traditional technical SEO and 57.1 on GEO — whether an AI answer engine can read, extract, attribute, and cite the page. That 32.7-point gap holds on 620 of 627 firms. It is not a weakness of small firms: the largest advertisers in the country sit on the same side of the line as single-office practices.

The cause is not blocked crawlers, and it is not missing plugins. Only 23 firms (3.7%) block AI crawlers by name. What fails is trust: 98% of firms make claims on their own homepage that no machine — and no reader — can verify, and 74% publish legal content with no named, credentialed attorney attached to it. The sentence an AI engine refuses to repeat is, in most cases, the same sentence a state bar would ask the firm to substantiate.

Nobody has won yet. The top score in the entire market is an 84. For any firm willing to fix the trust layer of its website, the category is wide open.

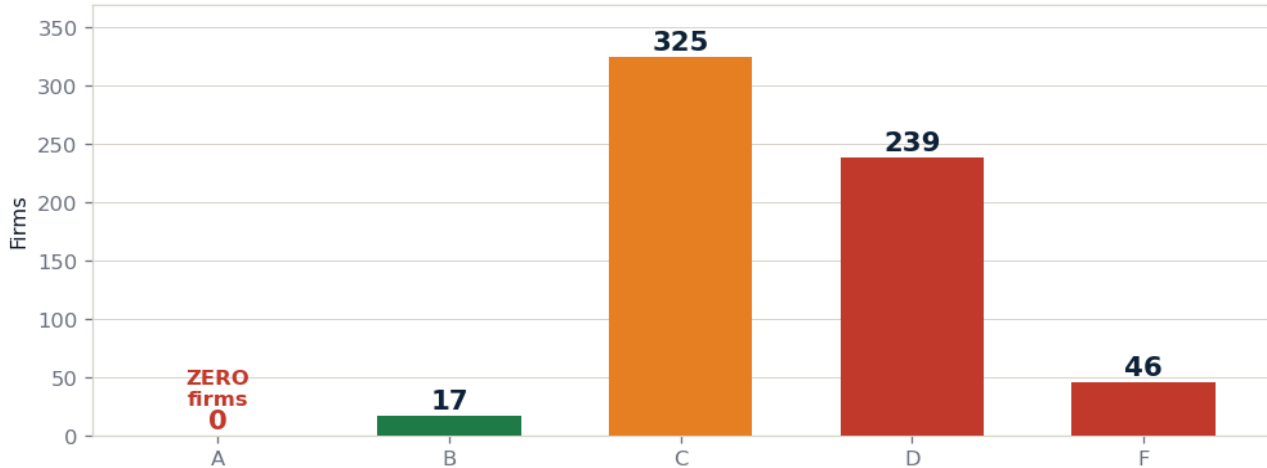
Who wrote this, and why you should discount it

This study was produced by Mass Tort Ad Agency (MTAA), a plaintiff-side advertising agency, using a scoring engine we built. We are an interested party: we sell services adjacent to the problems this study measures. Discount accordingly — and check the method instead, which is published in full, including the category weights, the sampling design, every firm we could not measure, and three errors our own verification caught before publication (p. 8). Three further disclosures: (1) our own homepage scores 93 on this engine; our site is tuned against criteria we chose, so we exclude ourselves from every table and treat that number as unearned. (2) One measured firm, Lerner & Rowe, shares a co-owner with a separate MTAA-affiliated business (PlatinumProfile); its score received no special treatment in either direction. (3) The author is not an attorney; nothing here is legal advice, and no statement in this report asserts that any firm has violated any rule of professional conduct.

Finding 1 · Zero A's, and the whole market in two grade bands

Not one A in the entire personal injury market

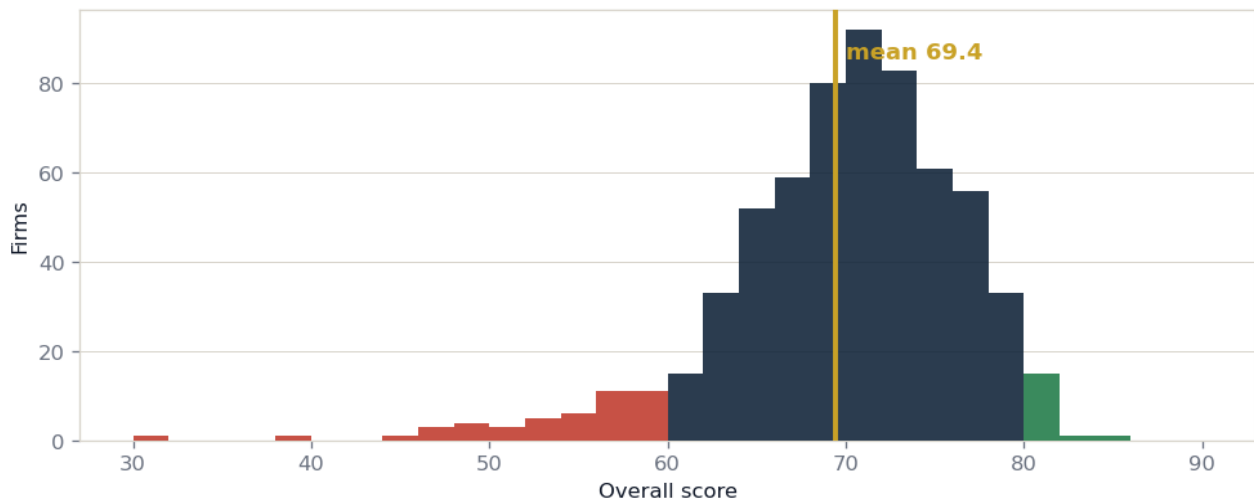
Overall AI-readiness grade, n = 627 firms · seo.mtaa.ai engine · Aug 2026



Not one of 627 firms earned an A. Seventeen firms — 2.7% of the market — reached a B. Ninety per cent of the industry sits in the C–D band, and 46 firms fail outright. The top score in the market is an 84 (Gus Anastopoulo Law Firm), the bottom a 31. The practical reading: no firm has a defensible AI moat today, and no firm is so far behind that the gap cannot be closed in a quarter.

The whole market lives in the C-D band

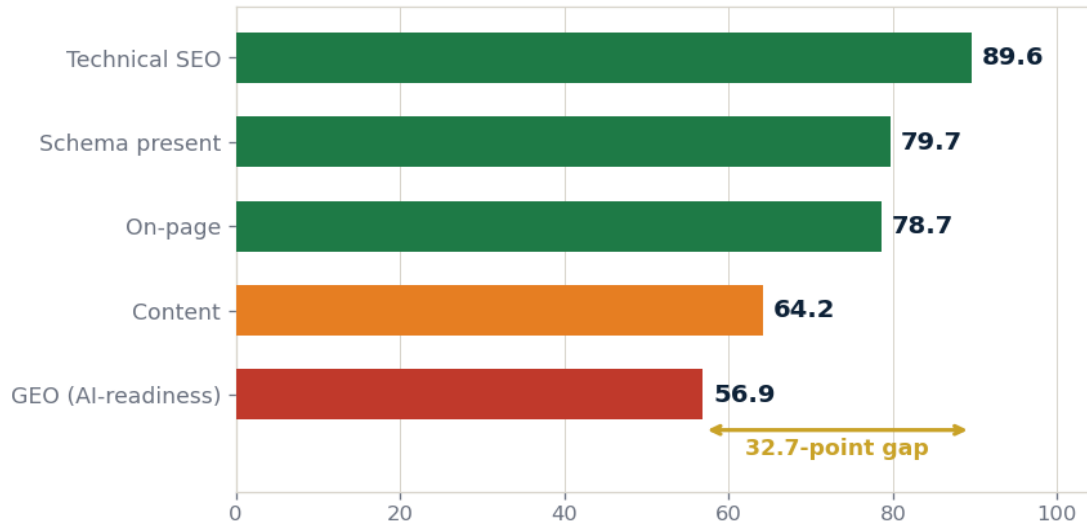
Distribution of overall scores, n = 627 · green = B or better (17 firms) · red = D/F (285 firms)



Finding 2 · A 32.7-point gap between the plumbing and the machines

Excellent plumbing, invisible to the machines

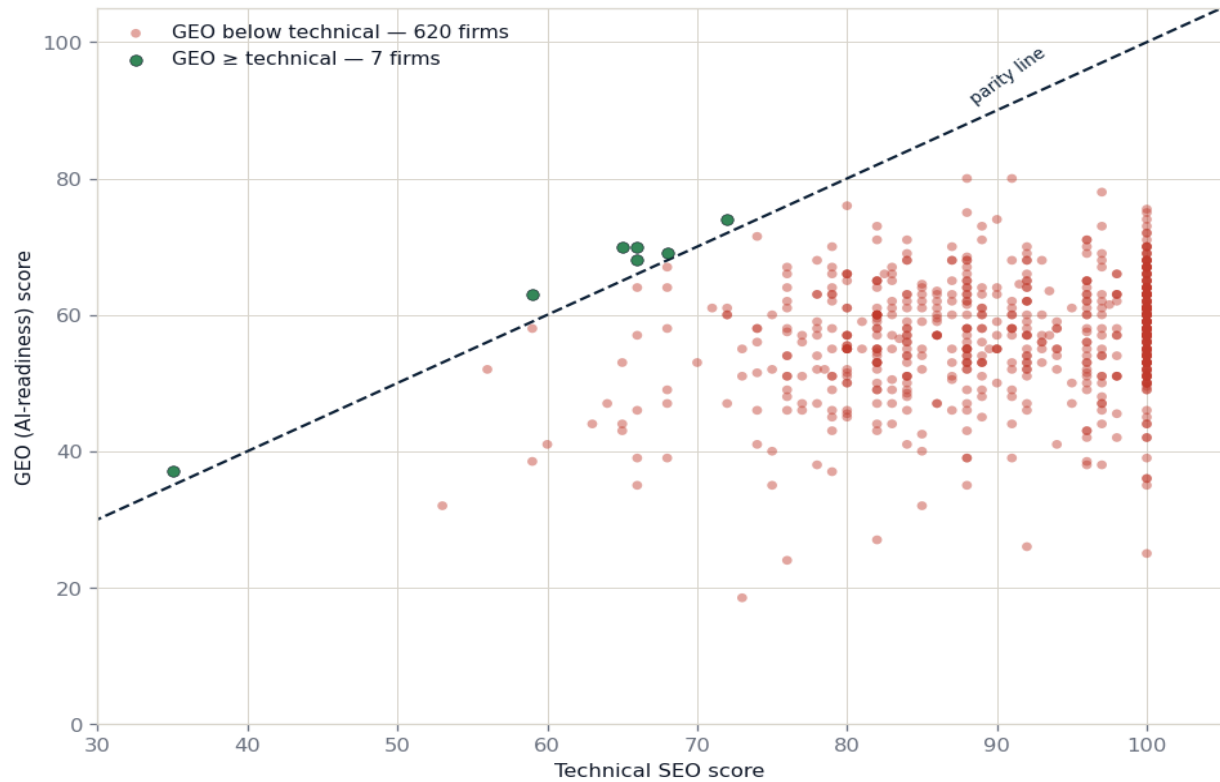
Mean category score across 627 firms · same engine, same weights as the agency audit



Twenty years of SEO investment worked. Technical scores average 89.8; 218 firms hold a perfect 100. Schema markup exists on nearly every site; on-page hygiene is broadly competent. The collapse happens on the one axis that decides whether an answer engine will quote the firm: GEO averages 57.1, and content depth 64.2. The industry built flawless websites for the search engine that ranks pages, not the one that answers questions.

620 of 627 firms sit below the line

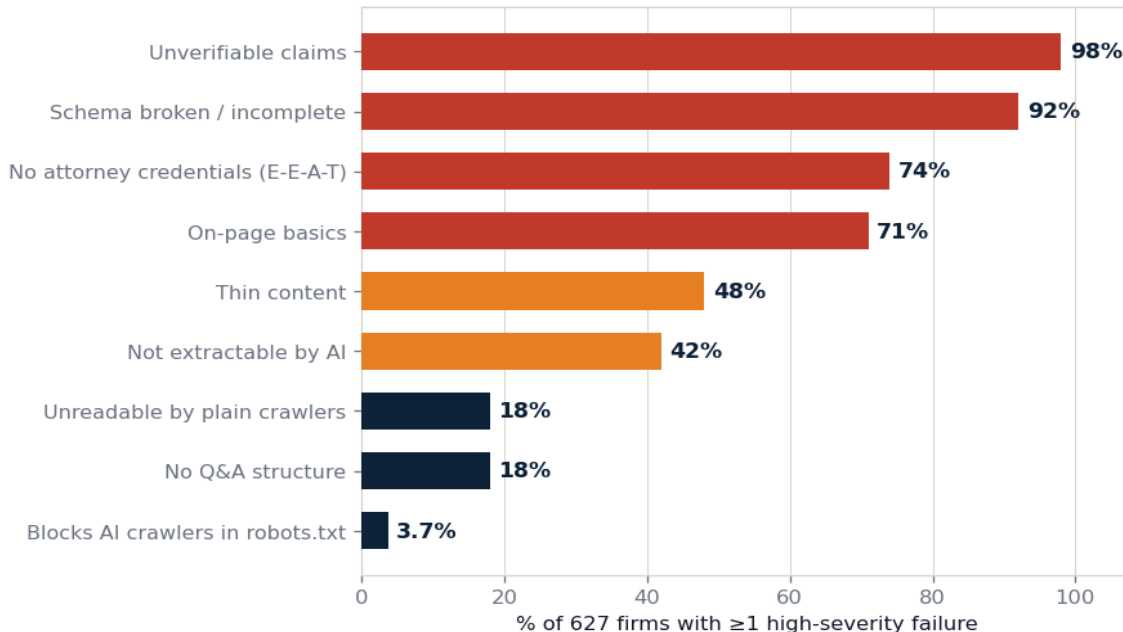
Every dot is a firm. Below the diagonal = optimized for the search engine being replaced, not the one replacing it. n = 627 · two-pass sample + full census



Finding 3 · What actually fails: trust, not technology

The failure is trust, not plumbing

High/critical failures only · model-classified taxonomy, pass/unmeasured records excluded · spot-verified against live sites



Every high-severity finding across all 627 firms was classified by a language model into a fixed ten-theme taxonomy (definitions, p. 9), with passing and unmeasurable checks excluded, and the results spot-verified against live pages. Two failures are near-universal:

Unverifiable claims — 98% of firms. "Best personal injury lawyer." "Record-breaking verdicts." "Billions recovered." Stated with no case name, no court, no date, no source. An answer engine that cannot verify a claim will not repeat it — so the firm's chosen differentiator is precisely what gets dropped from every AI answer.

No attorney credentials — 74% of firms. Legal content with no named author, no bar admissions, no jurisdictions, no verifiable experience. The page asserts a law firm wrote it, and never proves a lawyer did. Google calls this E-E-A-T; answer engines enforce it mechanically.

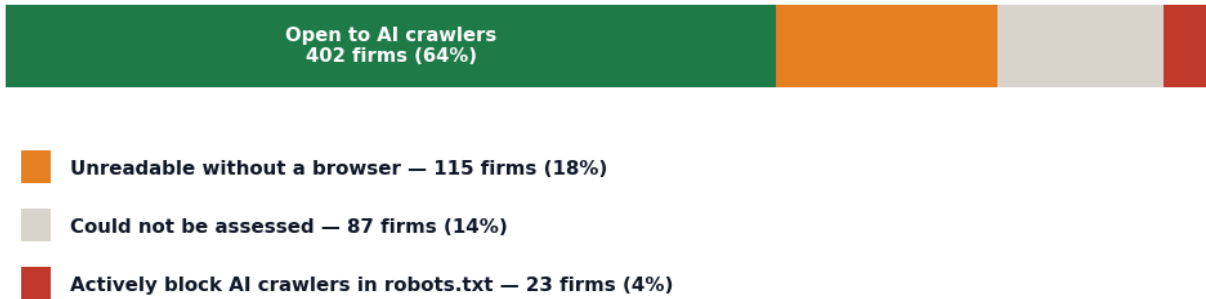
An observation on the rules — stated carefully

ABA Model Rule 7.1 and its state analogues have prohibited false or misleading lawyer advertising since 1977, and bar guidance has long treated unsubstantiated superlatives and comparisons as suspect unless they can be substantiated. Whether any specific page violates any specific state's rule depends on wording, disclaimers, jurisdiction, and substantiation we cannot see — this report makes no such claim about any firm. The observation is narrower and, we think, more useful: **AI answer engines and bar advertising rules are now testing the same sentence, and the test is the same question — can you back this up?** AI did not invent a new standard for legal marketing. It started mechanically enforcing an old one.

Finding 4 · The crawler myth, measured

Almost nobody blocks the AI — most just aren't worth reading

robots.txt + plain-fetch tests, n = 627 · all 6 sampled blockers independently confirmed



A popular explanation for AI invisibility is that law firm sites block AI crawlers. The census says otherwise. Twenty-three firms — 3.7% — block AI crawlers by name in robots.txt; each blocker we sampled was independently confirmed by fetching the live file. Most of the 23 share an identical eight-bot blocklist (GPTBot, ClaudeBot, Google-Extended, CCBot and peers), which is the signature of a single security plugin's default setting rather than 23 deliberate decisions. A further 115 firms (18%) serve pages that a plain, non-browser crawler cannot read at all — invisible to most AI retrieval the hard way. For 87 firms the robots.txt file could not be fetched by our engine; consistent with our measurement rules, those are reported as unassessed, not as failures.

For the 23: this is the cheapest high-severity fix in the entire study — one edit to one text file restores visibility to every major answer engine.

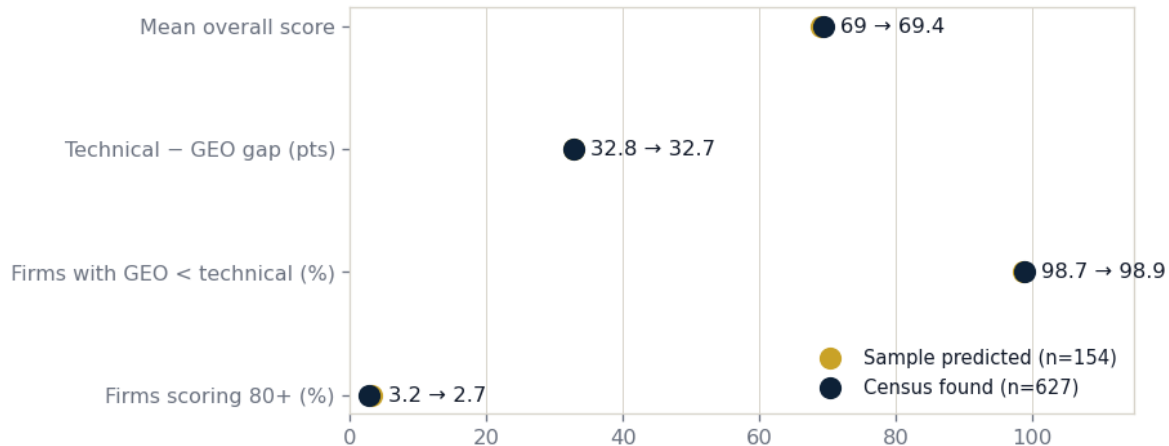
Finding 5 · Size buys plumbing, not trust

The national leaders in our two-pass sample average 96.8 on technical SEO against 88.2 for single-market firms — decades of agency spend, visible in the code. On GEO the order reverses: the single-market firms average 60.8 against the leaders' 54.1, and the best score in the study belongs to a one-market firm. Brand demand and AI-readiness are, today, essentially uncorrelated. For every firm that is not Morgan & Morgan, that is the most hopeful sentence in this report: the moat has not been dug yet.

Method — and how the study checked itself

We didn't trust the number — so we ran every firm

A 1-in-4 systematic sample predicted the market; the full census confirmed it within half a point.



Population and dedupe.

The frame is the BrandTerritory index of U.S. personal injury advertisers: 659 firm records, of which 652 active records carried a website. Domain normalization exposed 21 records that were duplicate entries of the same firm under name variants; these were merged, leaving 631 unique firms. Four could not be measured — two malformed records (excluded and corrected in the source data), one site returning 503, and one (arnolditkin.com) that rejects all non-browser traffic — leaving $n = 627$.

Instrument.

All scores come from one engine — the public auditor at seo.mtaa.ai — under one weighting: GEO 35%, content 25%, technical 20%, schema 15%, on-page 5%, identical to our August 2026 agency audit. Each audit measures the firm's homepage with browser rendering where possible (plain-fetch fallback is recorded per firm), fires the same deterministic checks per category, and uses a language model for the two judgment lenses (content quality, GEO citability). Checks the engine cannot complete — for example, a robots.txt it cannot fetch — are flagged *unmeasured* and are never counted as findings.

Sample first, then census.

We first drew a 1-in-4 systematic sample (158 firms) and measured each twice. Pass-to-pass repeatability: mean absolute difference 1.53 points (median 1; 7 of 153 firms moved ≥ 5). The sample estimated a population mean of 69.0 (95% CI 68.0–69.9). The number was bad enough that we ran the census rather than publish an estimate. The census answered 69.4 — inside the interval. One sample statistic missed: the share below a C (predicted $51\% \pm 6.9$; census 45%), a reminder that proportions are noisier than means at $n = 154$. Census firms were measured once, on the strength of the established repeatability.

Three errors we caught, kept in print.

1. Our first failure taxonomy used keyword matching; model classification of all 5,975 high-severity findings moved several headline numbers materially (schema 42%→92%; Q&A 95%→18%) and the keyword version was discarded. **2.** An early extract counted the engine's *couldn't-fetch-robots.txt* records as "blocking AI crawlers," inflating that figure to 45%; the true, independently confirmed rate is 3.7%, and the instrument now flags unmeasured checks explicitly. **3.** During hand-verification of 22 firms against live pages we initially mis-attributed the robots issue to the engine itself; the engine was correct. All three corrections were found by our own verification chain — classifier, pass/unmeasured separation, and live-site spot checks — before publication, which is the reason the chain exists.

Limits.

One page per firm (the homepage), scored by an instrument we built, with LLM lenses that carry ± 1.5 points of run-to-run noise — read scores as bands, not rankings. The taxonomy classifier was validated by spot-checks, not exhaustively. Eleven firms returned server errors during the census window and were retried once; those still failing are listed as unmeasured. Measurement dates: August 4–5, 2026. Websites change; numbers cited from this study should carry the date.

Appendix A · Taxonomy definitions

Unverifiable claims	Superlatives, comparisons, verdict/settlement figures with no case, court, date, or source.
Schema broken / incomplete	Structured data absent, unparseable, or missing the entities a legal page needs.
No attorney credentials	No named author; no bar admissions, jurisdictions, or verifiable experience attached to legal content.
On-page basics	Page title, meta description, or heading structure defective.
Thin content	Insufficient depth to answer the question the page exists for.
Not extractable	The answer exists but is buried or hedged; no self-contained statement a machine can lift.
Unreadable by plain crawlers	Page requires JavaScript/browser rendering; plain fetch returns no usable content.
No Q&A; structure	No FAQ or question-answer organisation; does not answer the questions clients ask.
Blocks AI crawlers	robots.txt disallows named AI user-agents (GPTBot, ClaudeBot, Google-Extended, etc.).
Other	Fits none of the above (6% of classified findings).

Appendix B · What a firm should do first

1. Substantiate or delete. Every superlative and every dollar figure on the homepage gets a source — case name, court, year — or gets cut. This is the 98% failure, it is free, and it is the same standard your bar already holds you to. **2. Put a lawyer on the page.** Named attorney, bar admissions, jurisdictions, in visible text and in Person schema. **3. Validate your markup** — most firms have schema; almost nobody has checked it parses. **4. If you are one of the 23:** remove the AI-crawler block from robots.txt today. **5. Then test it:** ask ChatGPT and Perplexity what your firm does. If the answer is generic or wrong, the machines are telling you what your website proves — which, today, is very little.

The Personal Injury AI-Readiness Study 2026 · Mass Tort Ad Agency · masstortadagency.com · Measured August 4–5, 2026 with seo.mtaa.ai. Reuse permitted with attribution and date. Firm-level scores available at brandterritory.ai.