



RESEARCH ARTICLE

Google and AI do not show the same Lubbock.

Texas AI Search Visibility Index: Lubbock Edition

6,000 selected usable responses. 20 business verticals. Five deliberately selected question pairs per vertical, repeated ten times on six configurations.

Google Organic, Google Maps, ChatGPT Search, Gemini, Claude and Perplexity. AI Overview panels come from the same Google Organic requests, not additional independent searches.

57%

About 57% (139 of 244) of tracked business website groups that appeared in Google's sampled top ten were rarely detected on at least one AI platform. This is low detection in sampled answers, not complete absence from AI.

Rare means fewer than 10% of sampled answers on at least one of four AI platforms: fewer than 5/50 for a single-vertical group, or 10/100 when two verticals are pooled. It does not mean absent.

By Caleb Monroe | Founder & Partner, Monsoon

Collection: September 9 to 10, 2026, Lubbock time. Editorial update: September 11. The same 6,000 observations; no new collection. Not independently peer reviewed.

[Read the interactive edition and evidence ledger](#)



LOCAL STORIES BEHIND THE CITYWIDE FINDING

Which AI assistant you ask can change which Lubbock businesses you hear about

CoNetrix was named in 50/50 sampled ChatGPT answers and 0/50 Claude answers. Gorilla Law Firm showed the opposite contrast: 1/50 on ChatGPT and 40/50 on Claude. These are examples of platform differences, not a ranking of the assistants.

A practical consumer story: compare the options returned by different assistants when looking for a local service. These two examples illustrate variation, not typical performance for every business.

[Evidence and denominator](#)

The Google-to-AI gap reaches restaurants and dentists

The pattern also appears in everyday local services: 14 of 21 tracked Google-visible restaurant website groups and 16 of 25 dentist groups were detected fewer than 5/50 times on at least one AI platform, even when known uncertain matches count as possible detections.

Make the finding tangible through the businesses residents use. Interview restaurants and dental practices about discovery, without presenting the sampled groups as every local provider or judging their service quality.

[Evidence and denominator](#)

The 244 tracked business website groups each appeared at least once in a sampled Google top ten. They are not a census of Lubbock businesses or 244 consistently high-ranking businesses. Names and domains are different measurements; neither is a recommendation or business-quality score.

[Definitions, reviewed counts and original documents](#)



SUPPORTING EVIDENCE: ORIGINAL TEN FINDINGS

FINDING 01

A Google presence is not an AI presence

50/50 vs 1/50

Gorilla Law Firm's domain appeared in the Google top ten in 50/50 sampled organic results. Its name matched 1/50 usable ChatGPT answers and 40/50 Claude answers.

Interpretation: Visibility needs to be checked by surface. This does not explain the cause or establish a business-wide absence.

[Evidence: gorilla:organic, gorilla:chatgpt, gorilla:claude](#)

FINDING 02

AI surfaces can name different businesses

50/50 vs 0/50

CoNetrix matched 50/50 ChatGPT answers and 0/50 Claude answers in the managed IT prompt set.

Interpretation: A single AI platform cannot represent every platform. These configurations are not a consumer market-share sample.

[Evidence: conetrix:chatgpt, conetrix:claude](#)



SUPPORTING EVIDENCE: ORIGINAL TEN FINDINGS

FINDING 03

Some businesses appear across both channels

49/50

Southwest Cleaning's domain appeared in 49/50 Google top tens. Its name matched 50/50 ChatGPT and 50/50 Perplexity answers.

Interpretation: This is an illustrative overlap case, not a category-winner claim or a recommendation rate.

[Evidence: southwest:organic, southwest:chatgpt, southwest:perplexity](#)

FINDING 04

Yelp occupies much of the sampled Google front page

66.5%

A Yelp domain appeared in 665/1000 sampled Google top tens and 809/1000 top twenties. Mobile and main domains are deduplicated within each result.

Interpretation: Discovery can take place inside a third-party result. Presence alone does not establish influence or an available placement.

[Evidence: yelp.com:top10, yelp.com:top20](#)



SUPPORTING EVIDENCE: ORIGINAL TEN FINDINGS

FINDING 05

Facebook is a frequent organic discovery surface

57.4%

Facebook domains appeared in 574/1000 sampled Google top tens and 785/1000 top twenties.

Interpretation: Social domains also appear in traditional search. These counts do not measure clicks, conversions or individual businesses.

[Evidence: facebook.com:top10, facebook.com:top20](#)

FINDING 06

Tripadvisor is highly visible in hospitality searches

96.0%

Tripadvisor appeared in 96/100 Google top tens across the five hotel and five restaurant prompts, each repeated ten times.

Interpretation: The observed pattern is specific to these hospitality intents, not every local industry.

[Evidence: tripadvisor:top10](#)



SUPPORTING EVIDENCE: ORIGINAL TEN FINDINGS

FINDING 07

Google AI Overviews vary by vertical

599/1000

Google returned AI Overview panels in 599/1000 organic observations. Foundation repair had 49/50, versus 2/50 for personal injury law. One returned panel had no answer text.

Interpretation: An AI Overview panel is not an additional independent search or proof that a business was cited.

[Evidence: aio:all, aio:foundation-repair, aio:personal-injury-law](#)

FINDING 08

AI answers link different source ecosystems

251 vs 1

BBB was answer-linked in 251/1000 usable ChatGPT responses and 1/1000 Claude responses. Expertise appeared in 236/1000 and 130/1000, respectively.

Interpretation: These are observed answer links, not the platforms' complete retrieval histories or proof that a source caused an answer.

[Evidence: bbb.org:chatgpt, bbb.org:claude, expertise.com:chatgpt, expertise.com:claude](#)



SUPPORTING EVIDENCE: ORIGINAL TEN FINDINGS

FINDING 09

The question matters, even within one industry

10/10 vs 0/10

Southwest Cleaning matched 10/10 Claude office-cleaning answers but 0/10 floor-care answers. It matched all ten floor-care answers on both ChatGPT and Perplexity.

Interpretation: A visibility gap in one sampled intent does not prove a service, content or capability deficiency.

[Evidence: southwest:cal-commercial-cleaning-1:claude, southwest:cal-commercial-cleaning-5:claude, southwest:cal-commercial-cleaning-5:chatgpt, southwest:cal-commercial-cleaning-5:perplexity](#)

FINDING 10

Repeated Maps sets were more consistent than organic sets

98.0% vs 88.7%

Mean within-prompt set overlap was 98.0% for Maps and 88.7% for organic results. The calculation includes 4,500 Maps and 4,500 organic response pairs.

Interpretation: Set overlap is not rank-order stability. Cache behavior and statistical independence are not established.

[Evidence: overlap:maps, overlap:organic](#)



METHODOLOGY AND MEASUREMENT

What was measured

Organic: One hit per usable response when the exact domain or explicitly grouped source-domain family appears at organic rank_group 1 to 10. Raw responses returned 17 to 20 organic results. Absolute mixed-SERP positions are not organic ranks.

Name examples: A specified, disclosed alias must appear as complete normalized tokens in answer prose, excluding URLs and displayed hostnames. Two separately implemented lexical matchers must agree. This checks calculation consistency, not semantic ground truth or all possible aliases.

Answer links: A source must be linked inline, referenced by a used Perplexity citation marker, or supplied as a Claude answer citation. Retrieval-only sources are excluded. Domains are deduplicated within each response.

Repetition: Jaccard overlap is intersection divided by union of returned place-ID sets for Maps (domain or title fallback where no ID exists) and domain sets for organic. Average within each prompt, then average eligible prompts equally. Empty/empty pairs and flagged responses are excluded. Order within the set is not measured. Maps set keys are SHA-256 hashes in the download, preserving set equality without publishing branch identities.

Observed visibility tables: Rows preserve recorded names and domain groups rather than independently verified branches. Organic counts exact-domain or normalized-title matches in the top ten; Maps uses the returned top twenty. All columns count the disclosed aliases as normalized tokens in answer prose, excluding bare hostnames. Every count matches the retained original table calculation. This validates computation, not semantic accuracy or all possible aliases. Counts by question use ten eligible responses per configuration.

Commercial interpretation: Search inclusion is not traffic, demand, customer preference or business quality. Prompt selection is purposive and unweighted by search volume. Cross-platform contrasts can reflect different access modes, models and retrieval settings. No causal ranking explanation, independent consumer sample or longitudinal improvement is claimed.

Returned model labels and configuration limits

ChatGPT supplied no model label for 981/1,000 observations; 19 reported gpt-5-6. Gemini reported 3.5 Flash-Lite. Claude used claude-haiku-4-5-20251001; Perplexity used sonar. Labels are provider-reported, not independently verified. Scraped experiences and API configurations are not assumed equivalent to consumer sessions.

[Exact prompts, request controls and definitions](#)



LIMITATIONS, ATTRIBUTION AND CONTACT

The low-detection threshold was selected after collection. This is exploratory analysis, not a preregistered endpoint.

Reviewed name variants are included. Known uncertain matches count as possible detections; two identity-uncertain groups remain in the denominator but cannot qualify. Unknown aliases can still affect counts.

Sector counts use each vertical separately and should not be summed into a citywide total. Questions differ by vertical; ten repetitions are not ten independent consumer samples.

Five deliberately selected prompt pairs per vertical and ten repetitions. This pilot is not a census, a representative sample of consumers or a statewide index.

Collection timestamps: 2026-09-09T16:11:24.911Z through 2026-09-11T01:57:55.453Z, UTC (September 9 to 10 in Lubbock). Repetitions may share caches or other dependencies. No independence-based confidence intervals or significance claims are made.

Scraped ChatGPT/Gemini experiences and Claude/Perplexity APIs are different configurations. ChatGPT returned no model label for 981 of 1,000 observations. Returned labels are provider-reported, not independently verified.

Gemini supplied detectable answer links in 211 of 1,000 responses. Missing source metadata is unknown attribution, not evidence of no citations.

Six ChatGPT location-clarification responses were replaced with valid retries of the same questions. One paired wrongful-death question was replaced after observing missing Maps business results, using personal injury lawyers in Lubbock and its matched AI question on all six configurations, ten repetitions each. The original 60 paired observations and earlier failed attempts remain archived. This post-collection query choice changes the estimand; the new wording is not a recovery of the old Maps result. No results are silently reclassified.

Featured business examples are deliberately selected contrasts. The broader tables preserve recorded names and domain groups, including overlapping aliases. A name match is not a recommendation, proof of branch identity or local service eligibility. Do not sum overlapping rows as market share.

Validation is automated structural and computational checking, not human annotation, peer review or measured semantic precision/recall. Complete recommendation rates, business winners and branch-level league tables are withheld.

This edition uses five question pairs per vertical. It is a bounded descriptive study, not comprehensive vertical coverage or certified category winners. The separate full Index release policy remains unchanged.

Produced by Monsoon Research Lab. Businesses were identified through the sampled searches. Client status did not influence inclusion or the reported results, as confirmed by Monsoon. Examples illustrate observed contrasts and are not randomly selected, endorsements or business-quality rankings. Monsoon works in SEO and AI-search visibility. This study is intended as a public research resource, not a sales report.

Source: Monsoon Research Lab, Texas AI Search Visibility Index: Lubbock Edition, September 2026, descriptive study v2.

[Caleb Monroe | Founder & Partner | caleb.m@growwithmonsoon.com](mailto:caleb.m@growwithmonsoon.com)

[Permanent edition, downloads and corrections](#)