

2026 EDITION

AI and the Future of Reputation Management

How AI Search, GEO, and Agentic Systems Are Reshaping Reputation Practice

By Will Kidder, PhD · Head of Content Operations, Status Labs

2.5B

CHATGPT DAILY QUERIES

2B

MONTHLY GOOGLE AI OVERVIEWS
USERS

60%

SEARCHES THAT END WITHOUT A
CLICK

Table of Contents

- 1 Introduction
- 2 The AI Search Revolution
- 3 What Is Generative Engine Optimization?
- 4 GEO in Practice
- 5 The Age of AI Agents
- 6 Reasoning Models and the Reputation Audit
- 7 The Zero-Click Paradox: Visibility Without Traffic
- 8 AI Content Generation and the Authenticity Question
- 9 Deepfakes, Misinformation, and LLM Grooming
- 10 The Regulatory Landscape
- 11 AI-Optimized Reputation Management
- 12 Conclusion
- 13 Frequently Asked Questions
- 14 Endnotes

Key Takeaways

AI search has become a primary information layer. ChatGPT now processes 2.5 billion queries daily and reached 900 million weekly active users as of February 2026, more than double the 400 million it reported a year earlier. Google's AI Overviews reach 2 billion monthly users, and Google AI Mode surpassed 1 billion monthly users following its May 2026 rollout to all U.S. users, with

queries more than doubling each quarter.¹ Anthropic's Claude has become the dominant enterprise AI platform, capturing the largest share of enterprise LLM API usage (40%, ahead of OpenAI) and making it one of the most likely tools a professional will use to research a company or executive. Managing reputation now means managing how all these AI platforms represent you.

Generative Engine Optimization (GEO) is the new SEO. AI systems demonstrate a systematic bias toward earned media over brand-owned content, and typically cite only 2–7 domains per response compared to Google's 10 results. Where traditional SEO optimizes for ranking, GEO optimizes for being cited in answers, prioritizing factors such as authoritative third-party coverage, structured extractable content, and consistent entity data. GEO strategies can boost AI visibility by up to 40%, but a November 2025 survey found that 75% of marketers lack confidence in how their brand currently appears in AI-generated summaries.

AI agents are the next reputation frontier. A 2025 survey found that 60% of consumers expect to use AI agents for purchases within the next 12 months. And when an AI agent researches a company before a meeting or shortlists vendors for a contract decision, it can form a judgment that bypasses traditional brand channels. GPT, Claude, and Gemini have all integrated AI agents, which can autonomously browse, shortlist, and transact on users' behalf. Google foregrounded Search agents at I/O 2026, bringing agentic query resolution into the world's most widely used search platform.²

The zero-click era demands new success metrics. Presence in AI-generated answers carries measurable business value even when it drives zero clicks. Approximately 60% of searches now end without a website visit, and Google's AI Mode ends sessions without a click 93% of the time. Yet AI-referred visitors convert at markedly higher rates: Semrush found the average AI search visitor to be 4.4 times as valuable as an organic search visitor.

Reasoning is now the default user experience, and AI research capabilities have democratized the deep background check. Today's frontier models (GPT-5.5, Claude Opus 4.8 and Sonnet 4.6, and Gemini 3.5) do not simply retrieve information. They reason, cross-reference, and weigh conflicting claims as a standard feature. And their integrated deep research capabilities mean any customer, investor, journalist, board member, or counterparty can commission a thorough multi-page research report on a person or organization in minutes, drawing from an extensive source record. For brands with inconsistent or poorly sourced digital footprints, this is a structural liability. For those with strong, well-sourced entity data, it is a structural advantage.

Regulatory divergence is accelerating. Donald Trump's January 2025 executive order revoked the Biden administration's AI safety framework in favor of an innovation-first approach, while the EU has been more cautious, with the AI Act's General Purpose AI obligations taking effect in August

2025. Companies operating across jurisdictions face incompatible regulatory environments, and the legal questions around AI defamation remain unsettled in courts on both sides of the Atlantic.

1 Introduction

When we published our 2025 white paper on AI and reputation management, we described the challenge facing practitioners as a parallel optimization demand. The argument was that organizations needed to manage their Google search presence and, separately, begin monitoring what AI platforms were saying about them. The two channels had different mechanics, different audiences, and different timelines.

Of course a Google search presence still matters, both on its own and for AI search optimization, but last year's framing needs to be updated. The parallel demand to focus on AI search has, in many ways, become primary.

ChatGPT now processes 2.5 billion queries per day across 900 million weekly active users.³ Google's AI Overviews appear in roughly one-quarter of all Google searches and reach 2 billion monthly users.⁴ Even Perplexity, which barely registered as a consumer product when our 2025 paper was published, grew at more than 20% month over month and now processes more than 780 million monthly queries.⁵ Google Gemini holds 21.5% of the AI chatbot market, up from 5.4% a year earlier.⁶

It's abundantly clear that the question for reputation practitioners is no longer whether AI matters. It's whether they have adapted to a world where AI is the first source many users encounter about a person or organization.

Several developments since our 2025 paper have sharpened the picture of how AI affects reputation. The practice of Generative Engine Optimization (GEO), which we introduced a year ago as a forward-looking recommendation, is now an established discipline with dedicated tools, agencies, and a growing body of peer-reviewed research. OpenAI launched its agent product, Operator, in January 2025. By July it had been integrated into ChatGPT's main interface.⁷ AI agents (autonomous systems that browse, research, and transact on behalf of users) have moved from laboratory demonstrations to commercial products available to all paid users of Claude and ChatGPT, while Gemini announced agents as a key feature of the future of Google search at I/O 2026.

Meanwhile, the international AI model ecosystem changed dramatically in a single week when DeepSeek released its R1 model in early 2025, achieving performance comparable to OpenAI's o1 at a fraction of the training cost, an event that wiped approximately \$600 billion from NVIDIA's

market capitalization in a single trading session. That market cap has since been regained, but the DeepSeek release suggested that AI capabilities could commoditize at an international scale faster than most forecasters had anticipated.⁸

The question for reputation practitioners is no longer whether AI matters. It's whether they have adapted to a world where AI is the first source many users encounter about a person or organization.

Writing in 2026, we are now describing a field that has moved from specialist conversation into the mainstream of cultural attention. Coverage of AI permeates media at a sometimes overwhelming scale, from podcasts and YouTube videos to mainstream legacy media outlets. New York Times feature writers are publishing accounts of building functional software through natural language alone; TIME's year-end assessment declared that AI had changed work forever in 2025.^{9,10} Underneath this flurry of coverage, the technology continues to improve. The model lineage—from GPT-5.5 to Claude Opus 4.8, to Gemini 3.5—has turned over entirely since our last paper, each generation raising the baseline of what counts as capable. This progression has yet to show signs of slowing, and it seems likely we'll be talking about entire new generations of these models, with improved and perhaps unforeseen capabilities, a year from now.

The unprecedented pace of the AI revolution in work and search has been, and will continue to be, the driving catalyst of how we approach online reputation management. AI systems that struggled to reliably contextualize individual and brand histories in earlier iterations are now much more capable and widely used. They are deployed by everyday users for first impressions, and by professionals to brief investors, screen vendors, and assess both individuals and organizations before consequential decisions are made.

The May 2026 Google I/O developer conference crystallized what had been an accelerating trend into a definitive shift. Google announced that AI Mode—its fully AI-generated search experience—had already surpassed 1 billion monthly users and was being rolled out to all U.S. users, that the traditional search box itself was being redesigned as an AI-first entry point, and that Search agents would be capable of executing complex, multi-step tasks without users ever visiting a source website.¹²

Taken together, these announcements confirm what reputation professionals have been tracking at the margins for some time now: the information layer between people and the web is now substantially AI-mediated, and that mediation is accelerating.

This substantively updated edition of *AI and the Future of Reputation Management* expands our prior analysis with these new developments at its center. The sections on AI search and GEO are

substantially new, as the field has moved considerably since 2025. The section on AI agents is entirely new. We have updated the regulatory, deepfakes, and AI content sections to reflect 2025–2026 developments and discuss what the future may look like in this new regulatory space.

The foundational argument from our 2025 paper, that reputation managers must treat AI systems as primary information gatekeepers requiring active management, has not changed. The urgency, and the role of AI in search and in our everyday lives, has.

2 The AI Search Revolution

Our 2025 paper opened with the observation that ChatGPT had reached 1 billion users faster than any consumer application in history. The platform has since grown well beyond that milestone. It now processes more than double the volume recorded a year earlier, and as of February 2026 it had reached more than double the 400 million it reported in February 2025. ChatGPT now ranks among the most visited websites in the world.³

It's a truly unprecedented growth rate, but ChatGPT, though sometimes spoken of as synonymous with AI chatbots by everyday users, is only one part of a broader shift in how a large portion of the world's population accesses information. Platforms like Claude and Gemini continue to gain users, opening up an AI search environment that is more dispersed than the standard Google search environment we've become accustomed to for the last 20 years.

Google still remains the dominant search engine by total volume, and its own AI features are expanding. AI Overviews, the AI-generated summaries that appear above traditional search results for a growing share of queries, reached 2 billion monthly users⁴ and peaked at approximately 25% of all Google searches in mid-2025, though the share had declined to under 16% by late 2025.¹¹ Regardless, that figure was still higher than the 6.5% share in January 2025 and the 13% share in March 2025.¹²

AI Overviews initially concentrated on informational queries, but by late 2025, the share of AI Overviews appearing for commercial-intent queries had risen from 8% to 18%.¹¹ Reputation managers and marketing teams alike should take note of this expansion into commercial intent. AI summaries are now a standard feature of searches that affect purchasing and evaluation decisions, not just fact retrieval.

KEY TERM --- AI OVERVIEWS.

AI Overviews are AI-generated answer summaries produced by Google and displayed at the top of search results pages. They synthesize information from multiple web sources into a direct response,

often resolving queries without the user clicking through to any website. AI Overviews draw from content ranking in the top-10 organic results the large majority of the time, making traditional SEO directly relevant to appearing in them. They reached 2 billion monthly users in 2025 and have expanded from informational queries into commercial intent searches.

Perplexity's trajectory tells a parallel story about market fragmentation. Built from the start as an AI-native search engine rather than a chatbot with search features added, Perplexity grew at more than 20% month over month to surpass 780 million monthly queries by mid-2025, approached \$200 million in annualized revenue, and reached a valuation of approximately \$20 billion by late 2025.⁵ Its citation patterns differ sharply from ChatGPT's: the platform weights real-time content heavily, drawing a large share of its citations from recently published material. Organizations that have not produced substantive recent content face a specific disadvantage with Perplexity, as recency is a direct quality signal.

Perplexity's citation patterns are worth examining not because of the platform's scale, which is dwarfed by platforms like Gemini, Claude, and ChatGPT, but as a leading indicator. As these more widely used platforms continue expanding their real-time search capabilities, some version of the recency-weighted behavior Perplexity pioneered is likely to become standard across all major platforms.

AI's competitive landscape has fragmented in a way that has direct consequences for reputation management strategy. Research examining citation patterns across AI platforms found that only 11% of domains appear in both ChatGPT responses and Perplexity responses to similar queries.¹³ The platforms draw from different source pools and apply different authority signals.

Google AI Overviews pull from content ranking in the top-10 organic results the large majority of the time. This makes them far more dependent on conventional SEO than ChatGPT. Claude's Research mode applies its own source prioritization, surfacing primary sources and flagging evidentiary conflicts in ways that can produce meaningfully different characterizations of the same organization. An organization could rank well in Google AI Overviews and be effectively absent from ChatGPT, or be well-represented in Perplexity while carrying a damaging characterization in Claude's research outputs. Managing AI reputation now requires understanding how each platform selects and weighs its sources.

The 2025 release of DeepSeek-R1 added another wrinkle. Published by a Chinese AI research company at approximately 5% of the training cost of comparable U.S. models, DeepSeek-R1 matched OpenAI's o1 on a range of benchmark tasks, triggering what observers widely described as a "Sputnik moment" in American AI development.⁸ Its open-source release, freely available for modification and deployment, accelerated the proliferation of high-capability AI models across geographies and use cases. Though more mainstream models have since continued to differentiate

themselves from DeepSeek by demonstrating more and improved capabilities, there is still an important takeaway here for reputation practitioners: as AI capabilities become commoditized and more organizations deploy AI-powered products, the number of surfaces on which a brand's narrative can be represented, distorted, or hallucinated will continue to expand at an international scale.

The challenge is no longer preparing for one dominant AI platform. Rather, it is adapting to an ecosystem of several different AI systems operating simultaneously.

3 What Is Generative Engine Optimization?

In our 2025 paper, we introduced Generative Engine Optimization as one of four strategic recommendations. Since then it has become a standalone discipline with its own agencies, measurement platforms, academic literature, and a growing volume of practitioner research. While it remains a developing field in its infancy, those who were early to adapt have developed more capable GEO infrastructure.

There is certainly very important overlap between strong SEO and GEO, but the logic of GEO is different enough from traditional SEO that it requires its own framework, particularly in terms of targeting relevant queries rather than simple keywords and phrases.¹⁴

KEY TERM --- GENERATIVE ENGINE OPTIMIZATION (GEO).

Generative Engine Optimization (GEO) is the practice of structuring content, citations, and digital presence so that AI systems — including ChatGPT, Gemini, Google AI Overviews, and Claude — accurately represent and favorably cite an organization or individual when responding to relevant queries. Unlike traditional SEO, which targets ranking algorithms that return lists of links, GEO targets AI retrieval and synthesis systems that answer queries directly, often without directing users to any website at all. The term was introduced in a 2023 Princeton University research paper and has since become the dominant framework for AI search optimization. GEO is often used interchangeably with “Answer Engine Optimization” (AEO), an older term originally coined for featured-snippet and voice-search optimization that has expanded to encompass AI search; the labels remain in flux.

The most common misreading of the AI search shift is that SEO has been demoted. In fact, the opposite is true: because every major AI search platform sits on a retrieval layer that is, in effect, a search index, weak SEO foundations now cost a brand visibility in two places simultaneously instead of one: in conventional search results and in the AI-generated answer that increasingly replaces them.

KEY TERM --- RETRIEVAL-AUGMENTED GENERATION (RAG).

Retrieval-augmented generation is the architecture underlying every major AI search platform: ChatGPT search, Claude with web search enabled, Perplexity, Google AI Overviews, and Gemini. When a user submits a query, the system does not necessarily generate an answer from the model's training data alone. It often first retrieves a small set of documents from an indexed corpus (typically the open web, plus platform-specific sources like Wikipedia or Reddit), then passes those documents to the language model as context for the generated answer. The retrieval step uses ranking signals very similar to those of traditional search: relevance to the query, source authority, recency, and crawlability. The generation step composes the answer from what was retrieved. If a brand's content is not retrieved, it cannot be cited.

This architecture has a counterintuitive consequence for SEO. The signals that determine whether a page can be retrieved (clean HTML, crawlability, indexability, internal linking, page authority, freshness, structured data) are the same signals SEO has optimized for since the early 2000s. A page that Google can't crawl will not appear in AI Overviews. A page that lacks authority signals will not be selected from a retrieval pool of thousands of candidates. A page with broken structured data will not be parsed cleanly when the model composes its answer.

GEO-specific tactics like extractable structure, earned media density, entity consolidation, and query targeting are crucial for AI search. They shape what happens once content has been retrieved. However, skip the SEO layers and the GEO layer has nothing to act on.

The practical implication: traditional SEO fundamentals remain an important substrate on which GEO sits. Google's own developer guidance makes the underlying architecture explicit. In its AI Optimization Guide, Google states that "our generative AI features on Google Search are rooted in our core Search ranking and quality systems."¹⁵ This is direct confirmation that the EEAT signals that have long governed organic search—Expertise, Experience, Authoritativeness, and Trustworthiness—remain fully operative in AI search.

"Our generative AI features on Google Search are rooted in our core Search ranking and quality systems."

— Google Search Central, AI Optimization Guide

GOOGLE AI MODE: KEY FACTS.

Google's AI Mode surpassed 1 billion monthly users after rolling out to all U.S. users in 2026.¹ The redesigned Google Search interface replaces the traditional search box with an AI-first entry point that routes queries to specialized Search agents capable of executing multi-step tasks—research, comparison, and booking—without users leaving Google.² Google Search Central confirms that these AI features "are rooted in our core Search ranking and quality systems," meaning EEAT signals and SEO authority remain a primary mechanism for earning representation in AI-generated results.¹⁵

AI crawler access is the boundary condition of the entire stack. If a site blocks GPTBot (OpenAI), Claude-SearchBot (Anthropic), OAI-SearchBot (ChatGPT's search index), PerplexityBot, or Google-Extended at the robots.txt level, or has crawl errors that prevent these bots from rendering content, no GEO investment downstream of that block will produce AI visibility. Many brands quietly block these crawlers by default through CDN security rules or inherited robots.txt entries; the cost compounds month over month as AI-mediated discovery grows. The first audit any organization should run is a crawler-access audit: which AI user agents are currently allowed, which are blocked, and is that decision deliberate?

A NOTE ON LLMS.TXT.

Proposed in 2024 as a Markdown-based standard for declaring AI-readable site content, llms.txt has been adopted by Anthropic, OpenAI, Cloudflare, and Stripe on their own developer documentation sites, but as of early 2026, no major AI platform has committed to operationally fetching it. Google's Search Advocate stated in June 2025 that "no AI system currently uses llms.txt," and an OtterlyAI traffic analysis found request volume at just 0.1% of total AI crawler activity. For now, llms.txt is a low-cost insurance policy worth implementing on documentation-heavy sites but not a substitute for the technical SEO foundations AI retrieval actually depends on.¹⁶

Strong SEO foundations also protect against hallucination. When a reasoning model is asked about a well-indexed entity (one with extensive, authoritative, well-structured source material across the open web), its retrieval step returns a rich, consistent context, and its generation step produces grounded output. When the same model is asked about a thinly-indexed entity (sparse sources, inconsistent characterizations, weak entity data), retrieval surfaces fragments, contradictions, or nothing at all, and the model fills the gap from its training data, adjacent entities, or pure pattern-completion. That is the architecture of hallucination. Brands with weak SEO footprints face two compounding risks: invisibility in AI-generated answers, and active misrepresentation when the model fills retrieval gaps from training data rather than live sources.

The Difference Between GEO and SEO

While GEO and traditional SEO are related, there are still several key differences. The major AI search platforms are built on retrieval-augmented generation (RAG) systems that pull content from the same web indexes SEO has always aimed to populate. But GEO shifts the emphasis: toward earned media over brand-owned content, toward entity recognition over keyword density, toward recency and extractable structure over link-building alone, toward answering longer and more targeted queries.

The term GEO was introduced in a 2023 paper by researchers from Princeton University, Georgia Tech, the Allen Institute for AI, and IIT Delhi, published at the KDD 2024 conference.¹⁷ Their central finding: GEO strategies can boost source visibility in AI-generated responses by up to 40%. This was based on testing nine optimization methods across 10,000 search queries. The specific

methods that drove the largest gains are worth stating precisely: for lower-ranked pages, adding citations from credible sources boosted visibility by as much as 115.1%, while incorporating quotations and statistics—the best-performing methods overall—improved visibility by 22% to 37% across the study’s metrics.¹⁷

GEO tactic	Measured increase in source visibility
Authoritative citations (lower-ranked sites)	+115.1%
Quotations from recognized sources	+37%
Integrated statistics	+22%

Source: Princeton GEO study (endnote 17).

A 2025 follow-up study from researchers at the University of Toronto systematically compared how AI search engines select sources against how Google does, across multiple verticals and languages.¹⁸ The central finding applies directly to reputation management, as the study found that AI search systems show “a systematic and overwhelming bias towards Earned media — third-party, authoritative sources — over Brand-owned and Social content.” This is a fundamental structural departure from Google, where first-party content (a brand’s own website and press materials) can rank highly with proper technical optimization. In AI search, the primary signal is often what credible third parties have written about an organization, not what that organization has written about itself.

In AI search, the primary signal is often what credible third parties have written about an organization, not what that organization has written about itself.

The scarcity of AI citations reinforces this point. Where a Google search returns 10 links for users to evaluate, AI models typically reference 2–7 domains in a given response.¹⁹ Competition for AI citations is more concentrated than competition for Google rankings. Being excluded from the set of sources an AI draws on means effectively not existing for that query. There is no second-page equivalent to discover later.

A coherent GEO strategy needs to account for several platforms, because a brand’s audience may reach it through any of them and the audiences for each platform can be importantly distinct. ChatGPT’s user base skews toward personal research and information tasks; Claude’s toward professional, enterprise, and analytical work, including the due diligence and competitive intelligence use cases where reputational accuracy matters most; Google AI Overviews can easily

reach the mainstream Google user population; Perplexity users may favor research-intensive queries where recency and source diversity carry high weight.

Content recency is a harder constraint in AI search than in traditional search: AI systems and crawlers systematically favor recently published and frequently updated content. Organizations that maintain a regular cadence of credible, well-sourced content publication are systematically favored over those that publish infrequently or whose most substantive recent coverage dates from several years ago.

Together, these platform dynamics and content signals define the competitive terrain of modern AI search. Brands and practitioners who understand each platform’s citation architecture, recency weighting, and source preferences are equipped to act on that knowledge.

There are also real platform-specific patterns that require distinct strategies:

How the major platforms select and cite sources

Platform	How it selects and cites sources
ChatGPT	Draws heavily on Wikipedia and its broad training corpus, supplemented by real-time web search for current topics. One analysis of 680 million AI citations found that, among citations going to ChatGPT’s 10 most-cited domains, 47.9% went to Wikipedia. GPT-5.4, released in March 2026, arrived in standard, Thinking, and Pro versions, with GPT Thinking offering deliberative reasoning. This has continued with GPT-5.5. Citation behavior across modes favors established, high-authority domains. ²⁰
Claude / Sonnet 4.6 and Opus 4.8	The leading enterprise AI platform, with 40% of enterprise LLM API usage—ahead of OpenAI’s 27%—by late 2025. ²¹ Through Cowork and its Research mode, it can conduct agentic multi-step web research producing detailed cited reports. It integrates with Google Workspace, Jira, Asana, and dozens of other connected services. ²² Citation behavior reflects Claude’s training: a strong preference for primary and authoritative sources, and a tendency to surface conflicting accounts rather than flatten them. A professional conducting due diligence on an executive or organization is more likely to be running that query in Claude.
Google AI Overviews / Gemini	Operates across two surfaces with distinct mechanisms. AI Overviews, powered by Gemini, cite from content already ranking in Google’s top-10 organic results the large majority of the time. Conventional SEO is therefore directly load-bearing for this platform. ¹⁵ Gemini Advanced’s Deep Research agent, rebuilt on the Gemini 3 foundation in December 2025, browses hundreds of websites and Google Workspace documents to produce multi-page reports, and is now available to developers via API. ²³ As the AI layer of the world’s most visited search interface, Gemini has a broad reach to the general population. Google’s May 2026 I/O conference accelerated this trajectory: AI Mode—Gemini’s fully AI-generated search experience—surpassed 1 billion monthly users and was extended to all U.S. users, while the traditional search interface was redesigned to route queries

Platform	How it selects and cites sources
	through AI agents capable of completing multi-step tasks without users visiting source websites. ¹²
Perplexity	Pulls from Reddit (46.7% of its 10 most-cited sources), ¹³ community forums, and recently published editorial content, with a real-time search architecture that makes recency a dominant quality signal. Its user base skews toward research-intensive queries, and its citation behavior is more volatile than other platforms. It responds quickly to fresh, well-sourced content.

4 GEO in Practice

The measurement infrastructure for GEO is now substantial. A dedicated category of AI visibility platforms has emerged to track how individuals and organizations appear across platforms like ChatGPT, Gemini, and Claude, measuring citation frequency, sentiment, and competitive share of voice. The analogy to Google Search Console is instructive: just as the SEO profession developed rigorous measurement frameworks alongside search engine optimization, GEO is now benefiting from comparable tooling.

The operational question is no longer whether to monitor AI representations, but how systematically and across which platforms to do so.

Our 2025 paper's section on AI strategy identified four broad priorities: monitoring AI profiles, curating authoritative content, integrating AI tools, and optimizing for GEO. Those priorities remain sound as a foundation.

What has changed is the tactical specificity now available, along with the urgency to implement GEO strategies in an increasingly AI-mediated search environment. A November 2025 Optimizely report found that 75% of marketers lack confidence in how their brand currently appears in AI-generated summaries, despite the fact that 67% of consumers now use AI tools to research products and services.²⁴ That gap will not close without deliberate action.

1. Entity Consolidation

Entity consolidation is the foundation of any GEO effort. AI systems build representations of people and organizations by aggregating entity data: structured information about who someone is, what they do, and how they are characterized across credible sources. Inconsistencies in that data create ambiguity that models resolve unpredictably, often defaulting to whatever source they find most authoritative regardless of accuracy.

Wikipedia remains one of the most important sources across AI search, and it is the single most-cited source for ChatGPT, accounting for nearly half (47.9%) of its 10 most-cited sources.¹³ The Wikimedia Foundation's statement that "every LLM is trained on Wikipedia content, and it is almost always the largest source of training data in their data sets" applies with particular force here. Maintaining an accurate, balanced, well-sourced Wikipedia article is not optional for any organization with substantial AI reputation exposure. For executives and individuals who meet Wikipedia's notability standards, maintaining an accurate biographical entry carries disproportionate weight in AI-mediated representation.

2. Earned Media Dominance

Earned media dominance is the strategic insight that most clearly differentiates GEO's optimization priorities from traditional SEO's. Because AI systems systematically bias toward third-party coverage over brand-produced content, the most effective GEO investment may not be better technical SEO on company websites but more credible external coverage. Chen and colleagues at the University of Toronto characterized this bias as "systematic and overwhelming."¹⁸

For reputation managers, this shifts emphasis toward media relations, authoritative industry placements, and substantive presence in community platforms like Reddit. Editorial placements in credible publications, research citations in industry reports, and positive coverage in relevant trade media all carry direct AI visibility value.

None of this eliminates the value of brand-owned content. Targeted on-site publishing and optimization still carry real GEO weight, particularly for branded queries—the "What does [company] do?" questions—where a clear, well-structured company website and owned profiles remain among the primary sources AI systems draw on.

CAUTION --- TACTICS THAT BACKFIRE.

Several tactics that gained traction through 2025 under the GEO banner can each damage the SEO foundation on which AI search visibility depends, including artificially refreshing publication dates to exploit recency signals, publishing overly self-promotional comparison listicles engineered to capture AI citation patterns, and embedding hidden text intended to steer AI-generated summaries. Date refreshing without substantive content improvement is detectable, while self-promotional formats face increasing editorial scrutiny from both human reviewers and AI systems. When these tactics reduce an organization's organic search rankings, they simultaneously reduce its AI retrieval visibility. The most extreme end of this spectrum has been formally classified by Microsoft Research as "AI Recommendation Poisoning": the embedding of hidden prompts in website content to steer AI-generated outputs toward favorable recommendations.²⁵ GEO that undermines the SEO foundation is not optimization.

3. Content Structure for Extractability

Content structure for extractability is the technical complement to earned media strategy. AI systems are most likely to cite content they can parse and extract cleanly. Research shows that hierarchical headings, standalone answer paragraphs (40–60 words at section starts), integrated statistics with attribution, FAQ formats, and comparison tables all correlate with higher citation rates.¹⁷ Organizations producing thought leadership content, executive profiles, or technical documentation should apply these structural principles: not as optimization tricks but because they reflect genuine clarity of communication. Content that is easy for a human to navigate is, almost by definition, easy for an AI to cite accurately.

4. Monitoring and Correction

Monitoring and correction close the loop. Regular systematic queries to major AI platforms (“Who is [executive name]?” “What is [company name] known for?” “What are the main criticisms of [company name]?” “Evaluate the evidence for [company]’s claims about [topic]”) provide an ongoing audit of AI representation that should complement traditional media monitoring. When errors appear, the correction pathway runs through the underlying source material, not through the AI platform itself. Though most major AI developers now have formal processes for submitting factual corrections, their success varies and you can’t directly edit an AI system’s output.²⁶ However, you can work to correct the erroneous Wikipedia entry, publish a well-sourced rebuttal of a false claim, or update the outdated press coverage that is driving the inaccurate representation.

5 The Age of AI Agents

On Jan. 23, 2025, OpenAI launched Operator, an AI agent that can take control of a web browser and autonomously perform tasks, from filling out forms to completing bookings.⁷ OpenAI CEO Sam Altman had signaled months earlier that 2025 would be the year AI agents moved from research to deployment. By July 2025, Operator had been integrated into ChatGPT as “agent mode,” initially available to paying Pro, Plus, and Team subscribers via the Tools menu.²⁷

Through its Cowork and Claude API features, Claude offers computer-use capabilities allowing autonomous navigation. Gemini has its own agentic features, and Google plans to feature agent capabilities prominently in its push toward AI Mode as the default. The question for brands and public figures is whether they are prepared for a market in which agents are already operating and will likely continue to increase in importance to search.

Google joined the agentic frontier in earnest at its May 2026 I/O conference, announcing Search agents capable of executing multi-step tasks—researching, comparing, and booking—directly within the search interface.^{1,2} Where OpenAI’s Operator navigates the open web on a user’s behalf, Google’s Search agents operate within its own search ecosystem, drawing on the same index and ranking signals that govern AI Overviews and AI Mode. Again, this integration means that agent-mediated interactions in Google inherit the same EEAT architecture as AI search results: brands and individuals whose content earns organic authority are more likely to be selected, referenced, or acted upon by agents operating within Google’s environment.

By fall 2025, agents had moved from a product category requiring technical translation into a phenomenon that mainstream publications were describing for general readers. Kevin Roose’s feature in the New York Times, “Not a Coder? With A.I., Just Having an Idea Can Be Enough,” described what he called “software for one” (functional applications assembled by non-programmers through natural language conversation). The concept of “vibe coding” began reaching audiences well outside the technology press.⁹

A year later, a Times opinion essay by Paul Ford, “The A.I. Disruption We’ve Been Waiting for Has Arrived,” described building a functional, searchable database application on his phone during a subway ride by typing prompts.²⁸ MIT Technology Review’s December 2025 assessment described the year as a “great AI hype correction” that had nonetheless arrived alongside genuine capability gains, including many driven by improvements in agents.²⁹

The agentic coding category crystallized most visibly around Claude Code and OpenAI’s Codex, launched in May 2025 as a tool designed, in TechCrunch’s characterization, to work “without users ever seeing the underlying code.”³⁰ By February 2026, Codex had more than 1.6 million weekly active users. Fortune reported that token usage had grown fivefold in weeks and that enterprise deployments had reached companies including Cisco, Nvidia, and Ramp.³¹

Anthropic extended the same agentic logic to knowledge work broadly, with tools like Claude Code for developers and Claude Cowork for general business users automating tasks like document management, research, and communications through natural language.

By late 2025, Anthropic had overtaken OpenAI to lead the enterprise large-language-model market, with 40% of API usage to OpenAI’s 27%.²¹ There are reputational consequences for this adoption curve. The population of AI-mediated interactions capable of surfacing, describing, or misrepresenting a person or organization has grown far beyond the search query to encompass the routine research and evaluation workflows of knowledge workers who now regularly delegate substantive work to AI agents.

The scope of what agents do makes this a categorically different reputation challenge from anything the industry has previously faced. A search engine presents information for a user who

decides what to do with it. An agent acts. When a user activates ChatGPT's agent mode to research options for a purchase, to book a restaurant for a client dinner, or to prepare a briefing on an executive before a meeting, the agent forms a judgment that shapes real-world outcomes without direct human review of the underlying data. That judgment is drawn from whatever sources the agent treats as authoritative.

Harvard Business Review's April 2026 analysis identified three emerging interaction modes for AI agents: brand-deployed agents (companies using AI to interact with customers), consumer agents (AI acting on behalf of individual users across multiple brands), and full AI intermediation (autonomous AI-to-AI transactions with no direct human involvement).³²

While it may sound odd initially, that third mode is already functioning at commercial scale. Imagine ChatGPT's agent searches OpenTable, selects restaurants based on stated user preferences, and completes a reservation. Restaurant management AI handles the confirmation on the other end. Neither the diner nor the restaurant made a direct decision about each other. Two AI systems intermediated the transaction.

Recent research suggests consumer adoption of agentic AI is moving faster than brand preparation. A 2025 survey by Kearney found that 60% of shoppers expect to use agentic AI to make purchasing decisions within the next 12 months.³² Two-thirds of Gen Z and more than half of Millennials have already begun using LLMs to research products. It seems inevitable that those numbers will increase.

McKinsey projects that agentic commerce could orchestrate up to \$1 trillion in U.S. B2C retail revenue by 2030.³³ The reputational stakes are proportional to that scale.

The Reputational Risks of Agentic Agents

The risks, when agents get things wrong, fall into two categories. The first is representation: agents may describe products, services, or individuals inaccurately, and those inaccurate descriptions drive action rather than simply informing it.

Ali Furman and colleagues, writing in the February 2026 issue of Harvard Business Review, identified five specific agent risks: misinterpreting product specifications, acting beyond authorized spending limits, handling customer data unsafely, misrepresenting brands through inaccurate information, and failing to provide recovery pathways when errors occur.³⁴ The same study established a clear accountability pattern: when AI agents make errors, customers may attribute those failures to the brand, not to the AI system.

When AI agents make errors, customers may attribute those failures to the brand, not to the AI system.

The second category is discoverability. AI agents do not browse the open web the way humans do. They query sources they have been trained to trust, draw on real-time search where enabled, and apply ranking logic that may differ from Google's. Being invisible to an agent means exclusion from the shortlist it generates, even if the brand ranks well in traditional search. With AI models typically citing only 2–7 domains per response, an agent tasked with finding the best vendor in a category is not generating a list of options for a human to browse. It is making a recommendation, sometimes with no human step between recommendation and transaction.

AGENT MEMORY RISK.

The memory dimension of agentic AI introduces a further consideration. As agents accumulate user preferences and interaction history, they develop what amounts to a persistent model of which brands a given user trusts, has used before, or has expressed interest in. A positive brand experience registered by an agent creates a standing favorable signal for that user going forward. A negative experience, or an agent's initial miscategorization of a brand, can effectively suppress it from future consideration sets without any deliberate user decision. Managing this dynamic requires monitoring and feedback mechanisms that do not yet exist at scale, but organizations that begin building agent-specific brand data infrastructure now will be better positioned as agentic commerce expands.

Preparation requires specific operational steps. Product and service data should be structured for machine readability: consistent schema markup, clear specification sheets, explicit pricing and availability, and stated policies on returns and privacy.

Beyond data structure, organizations need to test agent workflows directly, querying different platforms' agent modes with the prompts their target customers are likely to use, evaluating what recommendations the agents generate, and identifying where their brand is absent, misrepresented, or inadequately described.

6 Reasoning Models and the Reputation Audit

The AI models most users encountered when our 2025 paper was published (GPT-4o, Claude 3.5 Sonnet, Gemini 1.5) were sophisticated retrieval and generation systems. They produced fluent, often accurate responses by predicting what text should follow a given prompt. The underlying mechanism was pattern-matching on training data. They did not necessarily carefully evaluate the quality of their own outputs before producing them.

Reasoning models are architecturally different. They deliberate: generating intermediate steps, checking conclusions against prior reasoning, revising approaches when they prove inconsistent. The outputs of that deliberative process are sometimes visible to users in “chain of thought” displays, but the underlying computation runs regardless of whether it is shown.

DeepSeek-R1, released as a preprint in January 2025 and subsequently published in Nature in September 2025, established that this reasoning capability could emerge from pure reinforcement learning, without the human-annotated reasoning examples that most prior research had assumed were necessary.³⁵ The model develops emergent behaviors including self-reflection, verification of intermediate claims, and what the research describes as “dynamic strategy adaptation.” The authors note that these capabilities arose without explicit programming: the model learned to reason because reasoning produced better outcomes during training.

OpenAI’s o1 and o3 models and Anthropic’s Claude 3.7 Sonnet with extended thinking were early exemplars of this architecture. By 2026, reasoning has been integrated into the default interface of every major platform: GPT-5.4 consolidated o-series reasoning into the standard ChatGPT experience without a mode switch, and this has continued with GPT 5.5; Claude Sonnet 4.6 and Claude Opus 4.8 apply extended thinking across Anthropic’s full capability range; and Google’s Gemini 3 Deep Think brings deliberative reasoning to the world’s most widely deployed search-adjacent AI platform.

Reasoning-capable models are no longer specialized tools users opt into. In many cases, they are the default.

Research on the behavioral consequences of reasoning architecture is beginning to clarify what this means for how AI evaluates information. One recent study examined whether reasoning models are more faithful than non-reasoning counterparts, specifically whether they acknowledge when external inputs influence their responses.³⁶ It found that reasoning models are considerably more likely to surface, rather than suppress, the factors that shape their outputs. This translates directly into greater transparency about source quality and evidentiary conflicts.

Despite the improvements in reasoning models, there remains reason to be cautious. Apple published a direct challenge to the assumption that reasoning model outputs reflect genuine deliberation. Their June 2025 paper, “The Illusion of Thinking,” found that frontier reasoning models demonstrate “complete accuracy collapse” at high-complexity tasks, perform inconsistently across similar problems, and can’t reliably employ explicit algorithms even when the task demands them.³⁷ The Anthropic team contested aspects of the methodology, and subsequent developments in Claude and other models have demonstrated significant improvements in these areas, but the core critique stands as a useful caution: the visible chain of thought does not guarantee that the model accurately applied the reasoning it presents.

Both findings matter for reputation management, and they point in the same direction. Reasoning models apply more scrutiny to their inputs than models that simply generate the most probable continuation of a prompt, though that scrutiny is not infallible and does not behave like the reliable

application of explicit rules. For brands with strong, consistent, well-sourced entity data, the additional scrutiny is an advantage: a reasoning model evaluating conflicting claims about an organization is more likely to favor the version that is best supported by credible, consistent sources. For those with thin digital footprints, outdated Wikipedia articles, or conflicting characterizations across different platforms, the scrutiny is a liability that compounds over time.

The practical response is a reasoning-model audit: a systematic evaluation of what a deliberative AI model surfaces when it reasons about a person or organization, distinct from what a passive query returns. This involves asking not just “Who is X?” but “Evaluate the available evidence about [company]’s claims regarding [topic].” “What sources contradict the mainstream characterization of [executive]?” “What would a skeptical reviewer of [company]’s public record focus on?” The outputs of such queries reveal the source landscape a reasoning model draws on and where the weak points in an organization’s AI reputation profile lie.

This kind of adversarial audit is more demanding than the standard AI profile monitoring. But it can also be more revealing, and more likely to surface vulnerabilities before they show up in more general queries and in consequential contexts such as due diligence processes, journalist research, or agent-generated reports.

What Is AI Deep Research?

ChatGPT’s Deep Research, launched in February 2025 and now integrated into GPT-5.5, autonomously browses the web for up to 30 minutes, synthesizes multiple sources, and produces detailed reports with inline citations. It was originally built for knowledge work and users doing intensive research in areas such as finance, science, policy, and engineering.³⁸

Anthropic’s Claude Research mode, launched in April 2025 and expanded in May 2025 to integrate with Google Workspace and connected enterprise services, runs comparable agentic research cycles lasting 5–45 minutes.²²

Google’s Gemini Advanced Deep Research agent, rebuilt on the Gemini 3 foundation in December 2025 and now available to developers via API, browses hundreds of websites alongside Google-native data sources to generate detailed synthesis reports.²³

The practical consequence for reputation management is immediate: any interested party or stakeholder can now commission a thorough research report on a person or organization in seconds. The report they receive synthesizes whatever is in the source record. A regulatory action that was later resolved, a critical profile from three years ago, or a Wikipedia article that records a historical controversy will appear in the output, weighted by the AI’s assessment of source credibility. These tools do not curate with an eye toward favorable representation but rather synthesize based on perceived credibility and accuracy and thoroughness of available sources.

This creates a vulnerability and a new challenge for reputation management. An organization might manage its first-page Google results effectively while remaining exposed to Deep Research outputs that find and synthesize lower-ranked but AI-credible sources. A cluster of third-party articles that mischaracterize an executive could appear as a coherent narrative in Claude's Research output regardless of whether those articles have high visibility in traditional search.

One practical response is to commission Deep Research reports on an organization or individual across all major platforms as a regular audit practice, treating the outputs as an indicator of current AI reputational exposure, and as a preview of what the next journalist, investor, or counterparty might encounter if they choose to look.

7 The Zero-Click Paradox: Visibility Without Traffic

Approximately 60% of searches now end without a click to an external website.³⁹ In Google's AI Mode, the platform's most AI-intensive search experience, roughly 93% of sessions end without a click.⁴⁰

KEY TERM --- ZERO-CLICK SEARCH.

Zero-click search refers to queries resolved directly on the results page — through AI-generated summaries, featured snippets, or knowledge panels — without the user clicking through to an external website. Approximately 60% of all searches now end without a click. In Google's AI Mode specifically, about 93% of queries resolve this way. Zero-click does not mean users received no answer. Rather, it means the answer was delivered by the platform itself, bypassing publishers entirely and decoupling visibility from website traffic.



A Bain & Company analysis published in 2025 found organic web traffic declining by an estimated 15–25% across major content categories. The analysis attributed the decline primarily to AI-generated summaries that absorbed queries previously requiring a click.³⁹ A Pew Research Center analysis of real browsing data from U.S. adults found that those who encountered an AI summary went on to click through to a website just 8% of the time, compared with 15% for those who saw only standard results—nearly halving click-through.⁴¹

The paradox is that search volume has not declined. It has grown. A larger and larger percentage of searches, however, resolve without producing traffic to any site.

Traffic and visibility have decoupled. When a user asks Claude or ChatGPT about the best product in a given area, and the chatbot delivers a paragraph recommending two or three options with brief descriptions, no click occurs. But a brand recommendation has been made. When Google’s AI Overviews include a brand’s positioning statement in a summary about an industry trend, the brand has shaped a user’s understanding without driving a visit. That recommendation carries commercial value that doesn’t appear in website analytics.

Semrush found that the average AI search visitor was 4.4 times as valuable as the average traditional organic search visitor, based on conversion rate. The study covered 500+ high-value digital marketing and SEO-related topics and subtopics and analyzed AI search experiences including Google AI Overviews, Google AI Mode, ChatGPT, Claude, and Perplexity.⁴²

This conversion premium reflects a selection effect. Users who reach a website through an AI recommendation have typically passed through a filtering and evaluation process that traditional search does not provide. They arrived via recommendation rather than a browsable list of sites.

Even small absolute volumes of AI-referred traffic can therefore carry real commercial value. Brand visibility within AI-generated answers is worth measuring even when it produces no direct traffic.

Users who reach a website through an AI recommendation have typically passed through a filtering and evaluation process that traditional search does not provide. They arrived via recommendation rather than a browsable list of sites.

There are several potential replacement metrics for the click-driven era, such as answer inclusion rate (frequency of brand content in AI summaries), entity presence index (consistent brand recognition across AI platforms), source authority score (assessment of content trustworthiness as judged by AI systems), and AI citation frequency (how often content appears quoted in AI responses).¹² None of these metrics exist natively in Google Analytics or most standard marketing reporting tools. Building the measurement infrastructure to track them requires either platform-specific tools or systematic manual auditing.

Among B2B organizations, 89% of buyers now use generative AI for self-directed research, but the companies being researched have largely not adapted their content strategies for AI-mediated discovery.¹⁹ The gap between awareness and preparation is fundamentally an execution problem, not a knowledge deficit: measurement frameworks, content strategies, and optimization workflows built for a click-driven world do not automatically transfer to an answer-driven one.

And there is a subtler development worth considering. Researchers studying proactive LLM agents describe a shift from systems that merely respond to explicit prompts toward agents that can anticipate and initiate tasks without direct human instructions.⁴³ Google Discover, personalized AI feeds, and agent-initiated research all represent this category. A user's AI assistant with access to a calendar may know the user has a meeting with an executive from a particular company tomorrow and surface a briefing on that company before the user thinks to ask. That briefing, accurate or not, shapes the interaction that follows. The population of moments at which an AI system forms or communicates a view about a person or organization has expanded well beyond the query box.

Google's May 2026 I/O announcements added another dimension to zero-search discovery: Personal Intelligence, a capability that integrates Google's AI systems with a user's Gmail, Calendar, and personal data to deliver proactive, personalized information.¹ Under Personal Intelligence, Google's AI may surface information about a person, business, or topic based on the user's prior context and relationship history—not because the user searched for it. The reputation

implications are substantial: the digital footprint management of the future may have to account for what an AI system proactively concludes and surfaces about you based on the totality of a user's data environment, not just what they type as a query.

8 AI Content Generation and the Authenticity Question

When we published our 2025 paper, we described how AI content generation was transforming production workflows and raising quality questions. By 2026, that transformation has accelerated, but quality questions remain. The question is not whether organizations use AI for content creation (most do), but whether the content they produce with AI assistance is responsible, accurate, and genuinely mediated by human creativity and editorial standards.

The hallucination problem has not been solved. A 2026 Nature article argues that large language models still produce confident, plausible falsehoods even after mitigation techniques such as retrieval, tool use, self-verification, and reinforcement learning from human feedback.⁴⁴

For an organization producing content at scale with AI assistance, that error rate is manageable with human review in place. Without review, errors that survive publication become source-of-record inaccuracies that other AI systems may subsequently cite. The cycle of AI-generated content feeding AI-generated answers creates a compounding risk when human agency and quality controls are absent.

A 2025 case illustrated the institutional dimension. Deloitte Australia produced a 237-page government report using Azure OpenAI GPT-4o in drafting, under a contract worth AU\$440,000. A Sydney Law School researcher, reading the published report, identified more than a dozen errors including fabricated academic references and an invented quotation attributed to a federal court judge.⁴⁵ Deloitte acknowledged using AI in the report's production, issued a corrected version in September 2025, and was required to refund the final installment of the contract. The firm maintained that the corrections did not affect the substance of the finding. Sen. Barbara Pocock of the Australian Greens called the report's errors the kind "a first-year university student would be in deep trouble for."

For reputation management, the Deloitte case is informative both as a demonstration of the reputational hazards of relying on AI without proper review protocols, and as an example of how AI-hallucinated content can pass through institutional review processes, reach publication, and become source material that other AI systems may subsequently index and cite.

The correction protocols described in the GEO section apply here with particular force: organizations should carefully audit AI-generated content they have published for factual accuracy. The goal is to treat AI-generated content not merely as a communication product but as a potential source that shapes their AI representation.

The authenticity question has a second dimension. Research published in 2023, in the early days of AI-generated content, found that participants attributed similar levels of credibility to AI-generated and human-written content, with AI-generated text rated as clearer and more engaging in some conditions.⁴⁶ AI systems have only improved in their ability to generate accurate and engaging content in the three years since.

This finding and others like it cut in two directions. For brands, it speaks to how AI-assisted content production can maintain quality at scale when human oversight is maintained. For the broader information environment, it means that high-confidence, authoritative-appearing content is proliferating faster than the mechanisms for verifying whether any given piece of it is accurate. This is the condition in which AI hallucinations about brands are most dangerous: a hallucinated claim presented with editorial polish is more likely to be treated as fact by both human readers and downstream AI systems.

The cycle of AI-generated content feeding AI-generated answers creates a compounding risk when human agency and quality controls are absent.

The operational implication remains the same as in our 2025 paper: human review isn't optional for AI-generated content that will be indexed and crawled. Press releases, executive profiles, policy statements, blog posts distributed at scale: all of these feed the source ecosystem that AI systems draw on to represent an organization. Quality control at the publication stage is reputation management at the source.

9 Deepfakes, Misinformation, and LLM Grooming

Our 2025 paper covered a range of deepfakes in depth: the Taylor Swift deepfake pornography incident, a \$25 million fraud loss from an AI-generated executive impersonation, various misinformation campaigns. The deepfake detection-versus-generation arms race has continued on the same trajectory: detection improving incrementally, generation capabilities outpacing it, particularly in video and image generation. But two developments since our 2025 paper warrant specific attention.

The first is LLM grooming: coordinated efforts to shape AI training and retrieval data with the explicit purpose of influencing how AI systems respond. A March 2025 investigation by the media credibility firm NewsGuard documented the scale of one such operation: a Moscow-based disinformation network called Pravda published 3.6 million articles in 2024 across a network of 150 websites. The operation targeted not human readers but AI crawlers.⁴⁷ The goal was explicit: to flood the training and retrieval data of major AI systems with pro-Kremlin narratives. The network's disinformation narratives appeared in approximately 33.5% of responses from ten major AI platforms tested, including ChatGPT, Claude, Gemini, and Microsoft Copilot.

KEY TERM --- LLM GROOMING.

LLM grooming is the coordinated seeding of AI training data and retrieval indexes with fabricated or distorted content, with the explicit goal of shaping what AI systems believe and say about a person, organization, or topic. Unlike conventional disinformation — which targets human readers through persuasion — LLM grooming targets the AI systems themselves, embedding false narratives into the sources those systems draw on when generating responses. The 2024 Pravda network operation (3.6 million articles across 150 websites designed to influence AI outputs) is the largest documented example to date.

LLM grooming is categorically different from traditional influence operations. Conventional disinformation targets human psychology; it works by persuading individual readers. LLM grooming targets algorithmic systems; it works by being indexed and processed by AI models, which then propagate the narrative at scale to users who have no awareness of the operation. The Pravda findings establish that sophisticated actors are designing campaigns with AI propagation as the explicit mechanism, not a side effect.

WHO IS VULNERABLE.

Organizations in sectors where coordinated disinformation campaigns are common (e.g., defense contractors, pharmaceutical companies, energy producers, executives in contested policy debates) face a new threat model in which adversarial content is engineered to influence AI representations rather than human opinions. The correction protocols for LLM-groomed narratives are considerably more difficult than those for traditional misinformation, because the narrative is distributed across thousands of indexed pages rather than originating from identifiable sources that can be countered.

The second development concerns the continuing evolution of synthetic media as a professional fraud vector. The Arup deepfake fraud, in which deepfakes of executives led to the company losing millions of dollars, has proven to be a leading indicator rather than an anomaly. Finance and legal teams at major companies now incorporate deepfake verification into protocols for wire transfer authorization and executive communication.

The bottom line: as the creation of convincing deepfakes, including images, video, and voices, has become more accessible to everyday AI users, and those deepfakes have become more convincing,

the risk of criminal and defamatory applications has amplified. Despite ramped-up detection efforts, this trajectory shows no signs of slowing.

10 The Regulatory Landscape

The regulatory environment for AI has diverged sharply since our 2025 paper. Two major jurisdictions are moving in explicitly different directions, with real consequences for organizations operating across borders.

In the United States, the Trump administration revoked President Biden's October 2023 Executive Order on AI on Jan. 20, 2025, the first day of the new administration.⁴⁸ Three days later, a replacement executive order titled "Removing Barriers to American Leadership in Artificial Intelligence" directed federal agencies to review and dismantle AI policies deemed impediments to innovation, required a 180-day action plan for advancing U.S. AI dominance, and explicitly framed the goal as developing AI systems "free from ideological bias or engineered social agendas."⁴⁸

The philosophical departure from the Biden framework (which had emphasized safety testing, equity assessments, and third-party auditing) was immediate and unambiguous. A December 2025 executive order sought to advance a "minimally burdensome national policy framework" for AI. It also sought to preempt state-level regulation that the administration characterized as unnecessary fragmentation.⁴⁹

A July 2025 executive order, "Preventing Woke AI in the Federal Government" (EO 14319), went further, mandating that federal agencies only procure LLMs developed under "Unbiased AI Principles" of truth-seeking and ideological neutrality.⁵⁰ OMB implementation guidance (M-26-04, December 2025) imposes transparency and decommissioning obligations on AI contractors, a procurement-side lever shaping how the largest enterprise customer in the country evaluates model behavior.

The EU AI Act is moving in the opposite direction on a staged implementation timeline. It entered into force in August 2024. Eight categories of AI practices, including harmful behavioral manipulation, social scoring, and certain uses of real-time biometric identification, were prohibited in February 2025. The General Purpose AI Model obligations, governing systems like the LLMs underlying Claude and ChatGPT, became applicable in August 2025.⁵¹

Full applicability, including transparency requirements for generative AI and mandatory disclosure when users interact with chatbots rather than humans, is set to arrive in August 2026. The deepfake-specific provisions are among the most practically relevant for reputation management:

the EU Act will require that synthetic media be labeled as AI-generated, and that users be informed when they are interacting with a chatbot rather than a human representative.

At the state level in the United States, the picture is more active than the federal inaction suggests. Texas Gov. Greg Abbott signed the Texas Responsible AI Governance Act (TRAIGA) into law on June 22, 2025, effective Jan. 1, 2026. TRAIGA prohibits AI systems designed to manipulate human behavior, engage in unlawful discrimination, or generate child sexual abuse material and deepfake pornography. Civil penalties for violations range from \$80,000 to \$200,000, with continuing violations subject to additional daily fines. Texas became the third state with this type of AI governance legislation, and the count of states with some form of AI-related law continues to rise.

Meanwhile, the defamation frontier has produced its first notable ruling. In *Walters v. OpenAI*, the Georgia Superior Court of Gwinnett County granted summary judgment for OpenAI in May 2025.⁵³ Mark Walters, a gun rights radio host, had sued after ChatGPT falsely claimed he had embezzled funds from the Second Amendment Foundation. The court dismissed on three grounds: no reasonable reader would have interpreted the AI's disclaimer-accompanied output as factual assertion; OpenAI had not acted with actual malice required for a public figure's defamation claim; and Walters admitted to no actual damages at deposition. The ruling gave AI developers a partial reprieve, but the key legal questions it surfaced remain unsettled.

Legal scholars Lyrissa Lidsky and Andrew Daves of the University of Florida, writing in the *Journal of Free Speech Law* in 2025, have proposed a framework for addressing AI defamation by hallucination in reasoning models.⁵⁴ Their central argument: hallucinated defamatory statements should be treated as "inevitable errors," analogous to the honest mistakes courts have learned to tolerate from journalists in the interest of vigorous public discourse. They argue that AI developers should face liability standards applicable to distributors rather than publishers, but with statutory duties to warn users of verification needs and to preserve mechanisms for correcting the record.

Whether courts adopt this type of framework, or a stricter one, stands to substantially shape the liability environment for AI-generated defamation through the remainder of the decade.

11 AI-Optimized Reputation Management

The strategic recommendations in our 2025 paper centered on monitoring AI profiles, curating authoritative content, integrating AI tools, and implementing GEO. Those priorities remain valid. But what has since emerged is a clearer hierarchy of what to do first, greater tactical specificity at

each layer, and two new categories (agents and reasoning models) that require distinct preparation. The following is an overview of an AI-optimized reputation management strategy.

Six-Step AI Reputation Management Framework – Quick Reference

- 1. Conduct systematic AI reputation audits:** query all major AI platforms regularly and establish a baseline for how each system currently represents you.
- 2. Build a GEO content strategy:** consolidate entity information, create structured extractable content, and prioritize earned media placements in authoritative publications.
- 3. Prepare for agents:** structure product and service data for machine readability and establish clear transaction boundaries defining what agents are authorized to do on your behalf.
- 4. Update and maintain authoritative reference content:** particularly Wikipedia, ChatGPT's single most-cited source, and other high-authority reference pages.
- 5. Establish AI crisis protocols:** define correction workflows for hallucinations, LLM-groomed misinformation, and outdated AI representations before a crisis occurs.
- 6. Diversify platform presence:** distribute authoritative content across multiple high-credibility platforms, since AI systems treat multi-source corroboration as a credibility signal.

1. Conduct systematic AI reputation audits

The basic recommendation from our 2025 paper (querying major AI platforms to see how they represent you) now requires more rigor. Standard queries (“Who is [name]?” “What is [company] known for?”) should be supplemented with more specific relevant queries, adversarial queries (“What are the main criticisms of [company]?” “Has [executive] faced any controversies?”), and reasoning-model queries that expose the source evaluation process (“What sources support and challenge what is known about [company]?” “Evaluate the evidence for [company]’s claims about [topic]”). Log results over time, across platforms, and note which specific claims and sources recur.

For organizations with consumer brand exposure, agent-specific auditing stands to become more relevant. Test what Claude or ChatGPT’s agent mode recommends when asked to research your company for a meeting, compare you against competitors, or book a service you provide. Agent recommendations may differ substantially from what passive queries return, because agents can draw on different source combinations, deploy different tools, and apply different relevance criteria.

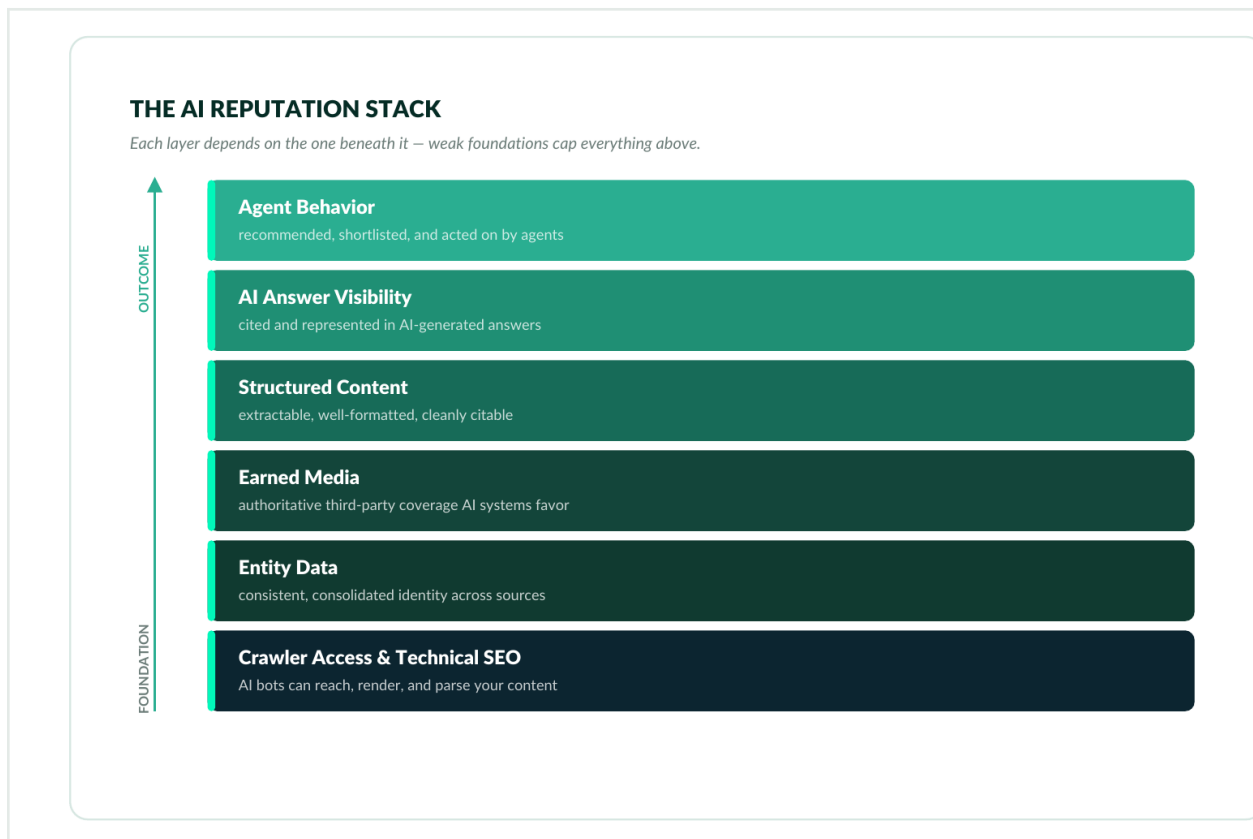
2. Build a GEO content strategy

Entity consolidation comes first: audit how you or your organization, executives, and key products are described across Wikipedia, LinkedIn, a company website, press coverage, and community platforms like Reddit. Resolve inconsistencies. The goal is a coherent, consistent entity description

that AI systems can draw on reliably: not keyword-rich copy but accurate, well-sourced, clearly structured information about who you are and what you do.

Earned media investment should be reframed as a GEO priority. Because AI systems favor third-party coverage over brand-owned content, media placements, industry publication features, and authoritative mentions carry direct AI visibility value. Track which placements produce AI citations, not just traffic.

The AI Reputation Stack. Each layer depends on the one beneath it—weak foundations cap everything above.



Structure existing content for extractability: add authoritative citations, integrate statistics with clear attribution, use FAQ formats for key topic areas, and ensure that the most important claims about your organization appear in self-contained paragraphs that AI systems can quote cleanly. The Princeton GEO research found that these structural investments can boost AI visibility by up to 40%, and the leverage is highest for organizations that currently rank lower in AI citation frequency.¹⁷

3. Prepare for agents

Organize product and service data for machine readability: consistent schema markup, explicit specifications, clear pricing and availability, stated return and privacy policies. Establish what HBR

describes as “transaction boundaries,” explicit parameters that define what agents are authorized to do when interacting with your brand. Clear transaction boundaries reduce the risk of agent misrepresentation.³⁴

Monitor how your brand appears in agent-generated recommendations by testing agent workflows directly. Where you find errors or omissions, correct the underlying source material exactly as you would for any AI representation issue. And consider pre-emptive stakeholder communication about how your brand participates in agentic workflows, particularly in sectors where sensitive transactions are common.

4. Update and maintain authoritative reference content

Wikipedia accounts for 47.9% of citations among ChatGPT’s 10 most-cited domains.¹³ A clean, accurate, appropriately detailed Wikipedia article is the highest-leverage single content investment for organizations with substantial ChatGPT exposure. Beyond Wikipedia, the authoritative reference layer includes official biographies on organizational websites, professional profiles on LinkedIn and other platforms, and entries in industry databases and recognized directories. These sources carry disproportionate weight in AI representations. Maintaining them is active reputation management, not optional upkeep.

5. Establish AI crisis protocols

When AI systems misrepresent you, whether through hallucination, LLM-groomed disinformation, or outdated source material, the response pathway differs from traditional PR crisis management. The path to correction runs through underlying sources, not through the AI platform. Identify the source material driving the error. Correct it where possible. Publish well-sourced counter-content if the originating source can’t be directly corrected. Submit factual corrections through AI developers’ official processes. Monitor for subsequent improvement. AI systems update their representations on varying timescales; correction may take weeks or months to propagate.

Pre-emptive inoculation against deepfake fraud remains essential, particularly for finance and executive teams. Clear internal protocols specifying that wire transfers or sensitive communications will never be authorized through AI-generated videoconferences or voice calls, communicated to employees and key partners, reduce the attack surface for Arup-style fraud even as deepfake technology continues to improve.

6. Diversify platform presence

Platform diversity means establishing the multi-source presence that AI systems treat as a credibility signal, not spreading effort across every channel. An authoritative company website, substantive LinkedIn activity for executives, a maintained Wikipedia presence, accurate and

non-promotional third-party coverage, and authentic community engagement on platforms like Reddit collectively build the kind of cross-platform entity recognition that AI systems draw on when forming representations.

12 Conclusion

In 2025, we described AI search as creating parallel optimization demands for traditional search and AI search. The verdict of 2026 is clearer. These demands are no longer parallel, and the trajectory is toward more relevance for AI search as the first layer of reputational exposure.

ChatGPT and Claude have billions of monthly visits. Google AI Overviews appear in one-quarter of all searches. Agents are completing research and purchases without users visiting brand websites. The infrastructure of AI-mediated information access is already operating at scale, not building toward it.

What is still transitional, and where the competitive advantage will be won or lost, is the gap between awareness and execution. Most individuals and organizations understand what is happening but have not done the operational work to address it. For those who close that gap now, the benefit is real: the individuals and organizations that establish strong entity recognition, consistent earned media presence, and agent-ready data infrastructure today will be far better positioned as AI-mediated discovery deepens, because AI systems favor well-established entities over newly introduced ones.

The agents question deserves particular emphasis at close. Every other development in this paper operates on the same basic logic as traditional reputation management: a human, at some point, decides what to believe or what to do. Agents change that logic. When AI acts on behalf of users, the reputation management challenge is no longer persuading a human who has received AI-mediated information. It is ensuring that the AI itself has accurate, complete, and well-structured information before it acts. That is a harder problem, with fewer established tools, and it is one that is already operating in the market, with major players like Google doubling down on agent capabilities.

At Status Labs, we have spent the past three years actively adapting our practice to AI: building GEO capabilities, developing monitoring infrastructure, and helping clients navigate the specific challenges of AI-mediated reputation. The pace of change in AI means that any assessment of the state of the industry like this one will require updating. What we can say with confidence is that those best positioned for what comes next are those that treat AI reputation management as an

ongoing operational discipline. Not a one-time project, not a checkbox: a standing commitment to understanding and actively shaping how AI systems directly influence reputation.

13 Frequently Asked Questions

What is generative engine optimization (GEO)?

Generative engine optimization (GEO) is the practice of structuring content, citations, and digital presence to ensure that AI search systems — including ChatGPT, Google AI Overviews and Gemini, and Claude — accurately represent and cite an organization or individual when responding to relevant queries. Where traditional SEO targets keywords and ranking algorithms, GEO targets queries and the retrieval and synthesis systems used by AI platforms that frequently answer those queries without directing users to any website.

How is GEO different from traditional SEO?

Traditional SEO optimizes for ranking in a list of links. GEO optimizes for accurate, favorable representation in an AI-generated answer. The two share some foundations—authoritative content, strong backlink profiles, and technical site health all help with both—but GEO diverges in its emphasis on earned media, structured extractability, and entity consolidation.

How do AI search platforms like ChatGPT, Claude, and Google AI Overviews decide what to cite?

Each platform uses different signals. Google AI Overviews draw almost exclusively from content already ranking in Google's top-10 organic results (the large majority of citations), making traditional SEO a primary lever. Claude, particularly in Research mode, prioritizes source authority and depth over recency, drawing from established publications, reference sources, and well-structured web content, and tending to synthesize across a broader range of source types than platforms optimized for speed. ChatGPT draws heavily on Wikipedia (47.9% of its 10 most-cited sources) and authoritative reference sources baked into its training data. A coherent GEO strategy accounts for all three rather than optimizing for a single platform.

What is the zero-click paradox, and why does it matter for brands?

The zero-click paradox is the disconnect between search volume (which continues to grow) and website traffic (which is declining). Approximately 60% of searches now end without a click to any external website, because AI systems and featured snippets resolve queries on the results page. For brands, visibility and traffic have decoupled: a brand can be prominently featured in AI-generated answers without generating any referral traffic. That visibility still carries

commercial value. AI-referred traffic, when it occurs, converts at markedly higher rates than organic search traffic.

How does AI search affect online reputation management?

AI search introduces reputation risks that traditional online reputation management does not address. AI systems can generate hallucinated or outdated characterizations that spread across multiple platforms and are difficult to correct. Adversarial actors can also deliberately seed AI training and retrieval data with false narratives (LLM grooming), shaping what AI systems say about a target without any of the traditional signals — a viral post, a negative article — that would alert a reputation manager. Effective AI reputation management requires systematic monitoring of AI outputs across all major platforms, not just traditional search results.

How should brands prepare for AI agents?

AI agents — systems that take autonomous actions on behalf of users, including making purchases and selecting service providers — present a distinct reputation challenge. When an agent selects a vendor, it delivers a recommendation with no human review between the AI's judgment and the transaction. Brands should structure product and service data for machine readability (consistent schema markup, explicit specifications, accurate availability data), establish transaction boundaries defining what agents are authorized to do on their behalf, and audit agent interactions regularly to detect misrepresentation.

What regulations currently govern AI-generated content and AI reputation risks?

The regulatory landscape is sharply divided by jurisdiction. In the United States, the Trump administration has adopted a permissive federal posture, while states including Colorado and Texas have enacted AI governance legislation covering high-risk automated decision systems. The EU AI Act imposes transparency and accuracy requirements on AI systems with consequential applications. On the defamation front, the first notable U.S. ruling involving AI-generated false statements (*Walters v. OpenAI*, Georgia, May 2025) found for OpenAI, dismissing the claim while leaving the broader question of liability for AI-generated falsehoods unsettled—a legal standard that will continue to develop.

14 Endnotes

1. Reid, Elizabeth, "A New Era for AI Search," Google The Keyword, May 19, 2026.
<https://blog.google/products-and-platforms/products/search/search-io-2026/>
2. Southern, Matt G., "Google's New Search Box Hands Queries to AI Agents, I/O Reveals," Search Engine Journal, May 19, 2026.

- <https://www.searchenginejournal.com/google-adds-ai-agents-to-search-redesigns-search-box-at-i-o/575311/>
3. Silberling, Amanda, "ChatGPT Users Send 2.5 Billion Prompts a Day," TechCrunch, July 21, 2025. <https://techcrunch.com/2025/07/21/chatgpt-users-send-2-5-billion-prompts-a-day/>; Malik, Aisha, "ChatGPT reaches 900M weekly active users," TechCrunch, February 27, 2026. <https://techcrunch.com/2026/02/27/chatgpt-reaches-900m-weekly-active-users/>
 4. Perez, Sarah, "Google's AI Overviews Have 2B Monthly Users, AI Mode 100M in the US and India," TechCrunch, July 23, 2025. <https://techcrunch.com/2025/07/23/googles-ai-overviews-have-2b-monthly-users-ai-mode-100m-in-the-us-and-india/>
 5. "Perplexity received 780 million queries last month, CEO says," TechCrunch, June 5, 2025. <https://techcrunch.com/2025/06/05/perplexity-received-780-million-queries-last-month-ceo-says/>; "Perplexity reportedly raised \$200M at \$20B valuation," TechCrunch, September 10, 2025. <https://techcrunch.com/2025/09/10/perplexity-reportedly-raised-200m-at-20b-valuation/>
 6. Similarweb, "AI Global: Global Sector Trends on Generative AI," January 2, 2026. <https://www.similarweb.com/corp/wp-content/uploads/2026/01/attachment-Global-AI-Tracker-6.pdf>
 7. Zeff, Maxwell, "OpenAI Launches Operator, an AI Agent That Performs Tasks Autonomously," TechCrunch, January 23, 2025. <https://techcrunch.com/2025/01/23/openai-launches-operator-an-ai-agent-that-performs-tasks-autonomously/>
 8. Teekah, Etan, "DeepSeek," Britannica, May 28, 2026. <https://www.britannica.com/money/DeepSeek>
 9. Roose, Kevin, "Not a Coder? With A.I., Just Having an Idea Can Be Enough," The New York Times, February 27, 2025. <https://www.nytimes.com/2025/02/27/technology/personaltech/vibecoding-ai-software-programming.html>
 10. Brynjolfsson, Erik, "AI Changed Work Forever in 2025," TIME, January 2, 2026. <https://time.com/7342494/ai-changed-work-forever/>
 11. Goodwin, Danny, "Google AI Overviews Surged in 2025, Then Pulled Back: Data," Search Engine Land, December 16, 2025. <https://searchengineland.com/google-ai-overviews-surge-pullback-data-466314>
 12. Thorson, Tanya, "AI Is Rewriting Visibility in the Zero-Click Search Era," Martech, November 4, 2025. <https://martech.org/ai-is-rewriting-visibility-in-the-zero-click-search-era/>
 13. Lafferty, Nick, "AI Platform Citation Patterns: How ChatGPT, Google AI Overviews, and Perplexity Source Information," Profound, June 5, 2025. <https://www.tryprofound.com/blog/ai-platform-citation-patterns>
 14. Hernandez, Jenna, "How GEO Will Replace Traditional SEO in 2026," Status Labs. <https://statuslabs.com/blog/how-geo-will-replace-traditional-seo-in-2026>
 15. Google Search Central, "AI Optimization Guide." <https://developers.google.com/search/docs/fundamentals/ai-optimization-guide>
 16. Schwartz, Barry, "Google Says No AI System Currently Uses LLMs.txt," Search Engine Roundtable, June 17, 2025. <https://www.seroundtable.com/google-ai-llms-txt-39607.html> (0.1% AI-crawler-traffic figure: OtterlyAI, otterly.ai/blog/the-llms-txt-experiment).
 17. Aggarwal, Pranjal, Vishvak Murahari, Tanmay Rajpurohit, Ashwin Kalyan, Karthik Narasimhan, and Ameet Deshpande, "GEO: Generative Engine Optimization," KDD 2024. <https://arxiv.org/abs/2311.09735>
 18. Chen, Mahe, Xiaoxuan Wang, Kaiwen Chen, and Nick Koudas, "Generative Engine Optimization: How to Dominate AI Search," 2025. <https://arxiv.org/abs/2509.08919>

19. Blyskal, Josh, "10-Step Framework for Generative Engine Optimization," Profound, July 1, 2025. <https://www.tryprofound.com/resources/articles/generative-engine-optimization-geo-guide-2025>
20. Bandom, Russell, "OpenAI Launches GPT-5.4 with Pro and Thinking Versions," TechCrunch, March 5, 2026. <https://techcrunch.com/2026/03/05/openai-launches-gpt-5-4-with-pro-and-thinking-versions/>
21. Tully, Tim, Joff Redfern, Deedy Das, and Derek Xiao, "2025: The State of Generative AI in the Enterprise," Menlo Ventures, December 2025 (survey of 495 U.S. enterprise AI decision-makers; Anthropic 40% of enterprise LLM API market share, OpenAI 27%, Google 21%). <https://menlovc.com/perspective/2025-the-state-of-generative-ai-in-the-enterprise/>
22. Anthropic, "Claude takes Research to new places," April 15, 2025. <https://www.anthropic.com/news/research>; and "Claude can now connect to your world," May 2025. <https://www.anthropic.com/news/integrations>.
23. Bort, Julie, "Google Launched Its Deepest AI Research Agent Yet — on the Same Day OpenAI Dropped GPT-5.2," TechCrunch, December 11, 2025. <https://techcrunch.com/2025/12/11/google-launched-its-deepest-ai-research-agent-yet-on-the-same-day-openai-dropped-gpt-5-2/>
24. Optimizely, "AI Discovery Has Rewritten the Rules of Online Shopping and Research — Most Brands Are Underprepared," PRNewswire, November 12, 2025. <https://www.prnewswire.com/news-releases/ai-discovery-has-rewritten-the-rules-of-online-shopping-and-research-most-brands-are-underprepared-302613059.html>
25. Microsoft Defender Security Research, "Manipulating AI Memory for Profit: The Rise of AI Recommendation Poisoning," February 10, 2026. <https://www.microsoft.com/en-us/security/blog/2026/02/10/ai-recommendation-poisoning/>
26. Hernandez, Jenna, "Generative Engine Optimization (GEO): The Future of Search and Digital Authority," Status Labs. <https://statuslabs.com/blog/generative-engine-optimization-geo>
27. Baraishuk, Dmitry, "OpenAI's ChatGPT Agent Outperforms the Model Alone," Belitsoft, July 17, 2025. <https://belitsoft.com/news/chatgpt-agent-openai-20250717>
28. Ford, Paul, "The A.I. Disruption We've Been Waiting for Has Arrived," The New York Times, February 18, 2026. <https://www.nytimes.com/2026/02/18/opinion/ai-software.html>
29. Heaven, Will Douglas, "The Great AI Hype Correction of 2025," MIT Technology Review, December 15, 2025. <https://www.technologyreview.com/2025/12/15/1129174/the-great-ai-hype-correction-of-2025/>
30. Bandom, Russell, "OpenAI's Codex Is Part of a New Cohort of Agentic Coding Tools," TechCrunch, May 20, 2025. <https://techcrunch.com/2025/05/20/openais-codex-is-part-of-a-new-cohort-of-agentic-coding-tools/>
31. Kahn, Jeremy, "OpenAI Reports Codex Usage Is Surging, Says It Plans to Make Codex Heart of Wider Agent Push," Fortune, March 4, 2026. <https://fortune.com/2026/03/04/openai-codex-growth-enterprise-ai-agents/>
32. Acar, Oguz A. and David A. Schweidel, "Preparing Your Brand for Agentic AI," Harvard Business Review, March–April 2026. <https://hbr.org/2026/03/preparing-your-brand-for-agentic-ai>; see also Kearney, "Agentic Commerce: From Brand Loyalty to Bot Logic," August 2025. <https://www.kearney.com/documents/d/asset-library-291362522/agentic-commerce-retail-disruption-pdf>
33. Schumacher, Katharina and Roger Roberts with Katharina Giebe, "The Agentic Commerce Opportunity: How AI Agents Are Ushering in a New Era for Consumers and Merchants," McKinsey & Company, October 17, 2025. <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-agentic-commerce-opportunity-how-ai-agents-are-ushering-in-a-new-era-for-consumers-and-merchants>

34. Furman, Ali, Ege Gürdeniz, Rima Safari, and Remzi Ural, "How Brands Can Adapt When AI Agents Do the Shopping," Harvard Business Review, February 19, 2026.
<https://hbr.org/2026/02/how-brands-can-adapt-when-ai-agents-do-the-shopping>
35. DeepSeek-AI, "DeepSeek-R1: Incentivizing Reasoning Capability in LLMs via Reinforcement Learning," Nature, vol. 645, pp. 633–638 (2025). <https://arxiv.org/abs/2501.12948>
36. Chua, James and Owain Evans, "Are DeepSeek R1 and Other Reasoning Models More Faithful?"
<https://arxiv.org/abs/2501.08156>
37. Shojaei, Parshin, Iman Mirzadeh, Keivan Alizadeh, Maxwell Horton, Samy Bengio, and Mehrdad Farajtabar, "The Illusion of Thinking: Understanding the Strengths and Limitations of Reasoning Models via the Lens of Problem Complexity," Apple Machine Learning Research, June 2025.
<https://machinelearning.apple.com/research/illusion-of-thinking>
38. Ha, Anthony, "OpenAI Unveils a New ChatGPT Agent for 'Deep Research,'" TechCrunch, February 2, 2025.
<https://techcrunch.com/2025/02/02/openai-unveils-a-new-chatgpt-agent-for-deep-research/>
39. Fishkin, Rand, "2024 Zero-Click Search Study," SparkToro (analysis of Datas clickstream data), 2024.
<https://sparktoro.com/blog/2024-zero-click-search-study-for-every-1000-us-google-searches-only-374-clicks-go-to-the-open-web-in-the-eu-its-360/>; Bain & Company, "Goodbye Clicks, Hello AI: Zero-Click Search Redefines Marketing," February 19, 2025.
<https://www.bain.com/insights/goodbye-clicks-hello-ai-zero-click-search-redefines-marketing/>
40. Harsel, Luke, "Google AI Mode's Early Adoption and SEO Impact," Semrush, July 30, 2025.
<https://www.semrush.com/blog/google-ai-mode-seo-impact/>
41. "Google users are less likely to click on links when an AI summary appears in the results," Pew Research Center, July 22, 2025.
<https://www.pewresearch.org/short-reads/2025/07/22/google-users-are-less-likely-to-click-on-links-when-an-ai-summary-appears-in-the-results/>
42. Handley, Rachel, "We Studied the Impact of AI Search on SEO Traffic. Here's What We Learned," Semrush, July 21, 2025. <https://www.semrush.com/blog/ai-search-seo-traffic-study/>
43. Lu, Yaxi, et al., "Proactive Agent: Shifting LLM Agents from Reactive Responses to Active Assistance," January 2025. <https://openreview.net/forum?id=sRIU6k2TcU>
44. Kalai, Adam Tauman, Ofir Nachum, Santosh S. Vempala, and Edwin Zhang, "Evaluating Large Language Models for Accuracy Incentivizes Hallucinations," Nature, April 22, 2026.
<https://doi.org/10.1038/s41586-026-10549-w>
45. "Deloitte to refund government after using AI in \$440k report," Accounting Times, October 2025.
<https://www.accountingtimes.com.au/technology/deloitte-to-refund-government-after-using-ai-in-440k-report/>; see also Paoli, Nino, "Deloitte Was Caught Using AI in a Government Report After a Researcher Flagged Hallucinations," Fortune, October 7, 2025.
<https://fortune.com/2025/10/07/deloitte-ai-australia-government-report-hallucinations-technology-290000-refund/>
46. Huschens, Martin, Martin Briesch, Dominik Sobania, and Franz Rothlauf, "Do You Trust ChatGPT? Perceived Credibility of Human and AI-Generated Content," arXiv, 2023. <https://arxiv.org/abs/2309.02524>
47. Bastian, Matthias, "Russian Fake News Network Floods Western AI Chatbots with Millions of Propaganda Articles," The Decoder, March 8, 2025.
<https://the-decoder.com/russian-fake-news-network-floods-western-ai-chatbots-with-millions-of-propaganda-articles/>

48. White House, “Removing Barriers to American Leadership in Artificial Intelligence,” Executive Order, January 23, 2025.
<https://www.whitehouse.gov/presidential-actions/2025/01/removing-barriers-to-american-leadership-in-artificial-intelligence/>
49. White House, “Ensuring a National Policy Framework for Artificial Intelligence,” Executive Order, December 2025.
<https://www.whitehouse.gov/presidential-actions/2025/12/eliminating-state-law-obstruction-of-national-artificial-intelligence-policy/>
50. White House, “Preventing Woke AI in the Federal Government,” Executive Order 14319, July 2025.
<https://www.whitehouse.gov/presidential-actions/2025/07/preventing-woke-ai-in-the-federal-government/>;
Office of Management and Budget, “Increasing Public Trust in Artificial Intelligence Through Unbiased AI Principles,” M-26-04, December 11, 2025.
<https://www.whitehouse.gov/wp-content/uploads/2025/12/M-26-04-Increasing-Public-Trust-in-Artificial-Intelligence-Through-Unbiased-AI-Principles-1.pdf>
51. European Commission, “AI Act: Regulatory Framework for AI.”
<https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>
52. Texas Legislature Online, “H.B. No. 149: Texas Responsible Artificial Intelligence Governance Act,” 2025.
<https://capitol.texas.gov/tlodocs/89R/billtext/pdf/HB00149F.pdf>
53. Scarcella, Mike, “OpenAI Defeats Radio Host’s Lawsuit Over Allegations Invented by ChatGPT,” Reuters, May 19, 2025.
<https://www.reuters.com/legal/litigation/openai-defeats-radio-hosts-lawsuit-over-allegations-invented-by-chatgpt-2025-05-19/>
54. Lidsky, Lyrissa and Andrew Daves, “Inevitable Errors: Defamation by Hallucination in AI Reasoning Models,” Journal of Free Speech Law, 2025. <https://www.journaloffreespeechlaw.org/lidskydaves.pdf>