

Journal of Social Science and Human Research Studies

Volume : 02 Issue : 09 September-2026

Generative Engine Optimization in Pharmaceutical Marketing: Benchmarking AI Search Visibility Across Commercial Therapeutics

Isaac Awulu

Arizona State University, Arizona, United States

Valentina Ochuko Obukadata

University of New Haven, Connecticut, United States

Imoiboho Williams

German Development Cooperation (GIZ), African Union, Addis Ababa, Ethiopia

Published Date: September 05, 2026**Article DOI: 10.65150/EP-jsshrs/V2E9/2026-04**

ABSTRACT : Background. Generative search interfaces have become a primary channel through which patients and healthcare professionals obtain drug information. Unlike the ranked-link paradigm that governed twenty years of pharmaceutical search marketing, generative engines synthesize a single answer, name a small number of brands, and cite a narrow set of sources. A therapeutic brand that does not appear in that answer is functionally absent from the channel, regardless of its organic search position.

Objective. This paper has three aims: to characterize the 2026 state of AI search visibility in health and pharmaceutical contexts using published empirical evidence; to specify a reproducible benchmarking framework for measuring the AI search visibility of commercial therapeutics across engines, languages, and therapeutic areas; and to identify the regulatory and methodological constraints that distinguish pharmaceutical Generative Engine Optimization (GEO) from GEO in unregulated commercial categories.

Methods. This is a synthesis and framework paper. No primary data collection was performed. Empirical figures are drawn from published academic literature, industry citation studies, and vendor benchmark indices, each attributed in the reference list. Because pharmaceutical GEO has no mature native literature, the measurement apparatus is assembled from six adjacent fields, and Section 2.4 states those debts explicitly so that the strength of each transfer can be judged. The proposed benchmark protocol is specified at a level of detail sufficient for independent replication.

Key findings from the reviewed evidence. (1) Authoritative medical sources account for a minority of AI health citations; one large-scale analysis of roughly 825,000 citations found that about 17.3 percent went to government, academic, medical-publisher, or hospital sources, while roughly three quarters went to commercial and platform domains. (2) Engines diverge sharply. The same analysis found ChatGPT directed about 24.1 percent of health citations to authoritative medical sources against about 15.1 percent for Google AI Mode, and Google's search-integrated products cited social and user-generated content four to five times more often than ChatGPT. (3) Brand-level visibility is engine-specific and language-specific to a degree that invalidates single-number reporting; one published pharmaceutical index reported a 33-percentage-point Answer Rate gap for a single dermatology brand between two engines in the same week and query set. (4) Generative engines are non-deterministic, and published bootstrap analyses show that apparent differences in citation share of fewer than roughly five to seven percentage points frequently fall inside the noise floor of single-run measurement.

Conclusion. AI search visibility is measurable for pharmaceutical brands, but only as a distribution rather than a point estimate, and only when stratified by engine, language, audience, and query intent. Because generative answers routinely surface safety and off-label content outside any promotional review process, pharmaceutical GEO measurement should be treated as a pharmacovigilance and medical-legal-regulatory monitoring function as much as a marketing one.

KEYWORDS: generative engine optimization, answer engine optimization, AI search visibility, pharmaceutical marketing, share of voice, large language models, fair balance, promotional compliance, retrieval-augmented generation

1. INTRODUCTION

1.1 The displacement of the ranked link

For two decades, pharmaceutical digital strategy rested on a stable assumption: a person with a health question typed it into a search engine, received a ranked list of documents, and chose one. Brand visibility was therefore a function of position, and position was the object of optimization.

That assumption no longer holds for a growing share of health queries. Generative engines return a synthesized answer with a small number of embedded citations. The user reads the answer. In many sessions no link is clicked at all. Analysis of pharmaceutical keyword sets reported by CMI Media Group found AI Overviews present on 52 percent of pharma-related searches, and that organic click-through rate on the same query fell from 25.8 percent to 7.4 percent when an AI Overview was displayed, a decline of roughly 71 percent (CMI Media Group, 2026). Gartner's widely cited projection anticipates a 25 percent decline in traditional search engine volume by 2026, with organic search traffic potentially falling substantially further (Spectrum Science, 2026).

The consequence for a commercial therapeutic is structural rather than incremental. A brand can hold the top organic position for a disease-state query and still be entirely absent from the synthesized answer that most users read. Ranking and answer inclusion are distinct outcomes produced by different mechanisms.

The shift is better understood as an interface change than as a channel change. Ogundairo and Ayivi-Donkor (2024b) argue that the movement toward autonomous conversational agents requires a corresponding change in product marketing architecture rather than a reallocation of spend within the existing architecture, since the surface on which a product is represented is no longer one the brand renders. The revenue consequence tracks the funnel-leakage pattern described by Sanni and Wedraogo (2024), in which value is lost at transition points that conventional channel reporting does not instrument.

1.2 Scale of the health channel

The volume involved is not marginal. Spectrum Science's analysis places ChatGPT alone at roughly 230 million health-related questions per week (Spectrum Science, 2026, as reported in PharmaGEO, 2026). On the professional side, IQVIA's March 2026 analysis found that 54 percent of healthcare professionals now use generative AI in clinical contexts, including drug information lookup, trial data summarization, and guideline checking (IQVIA, 2026, as reported in PharmaGEO, 2026). Specialist clinical AI platforms report query volumes in the millions per month from verified physicians.

Ilodigwe and Adesemoye (2025a) characterize this class of adoption as artificial intelligence functioning as a force multiplier for clinical decision-making, with the accompanying observation that implementation risk scales with the speed of uptake. Where the tool consulted is a general-purpose engine rather than a validated clinical system, the drug information it returns is assembled from an uncontrolled corpus.

A counterweight is warranted before overstating the shift. Conventional search still handles vastly more query volume in absolute terms, and one analysis estimates Google's volume at roughly 373 times ChatGPT's (Codeless, 2026, as reported in PharmaGEO, 2026). GEO does not replace search engine optimization. It adds a second, structurally different visibility surface governed by different mechanics, and brands that treat the two as competing budget lines tend to underfund both.

1.3 Why pharmaceutical GEO is a distinct problem

Most published GEO research and nearly all commercial GEO tooling was developed for unregulated consumer categories: retail, software, hospitality, consumer packaged goods. Pharmaceutical marketing differs on four dimensions that make direct transfer of those methods unsafe.

Promotional speech is regulated. In the United States, prescription drug promotion is governed by 21 CFR Part 202, which requires that risk information be presented with prominence and readability reasonably comparable to efficacy information. As of mid-2026 the FDA had not issued a rule specific to AI-generated promotional content; existing promotional labeling requirements and Office of Prescription Drug Promotion enforcement patterns apply regardless of what produced the draft (XDS Health, 2026a; XDS Health, 2026b). The regulatory question in a generative channel is therefore not whether the rules apply but who is accountable when the sponsor did not author the answer.

The channel is a safety information surface, not only a marketing surface. Generative answers about therapeutics routinely incorporate boxed warnings, contraindications, and risk evaluation and mitigation strategy content pulled from regulatory labels and reframed by the engine. Published index data indicates that regulatory label documents rank among the most-cited sources in metabolic and obesity answers (PharmaGEO, 2026). Visibility measurement in this channel therefore has pharmacovigilance implications absent from consumer GEO.

Off-label surfacing occurs without sponsor action. The same index reports that engines answering obesity queries return mentions of products approved only for type 2 diabetes. This is engine behavior rather than sponsor promotion, but it constitutes an unprompted indication expansion reaching prescribers and patients outside any medical-legal-regulatory review (PharmaGEO, 2026). No conventional promotional control addresses it.

The audience is bifurcated. A single therapeutic must be measured separately for patient-intent and professional-intent queries, which draw on different source pools, use different terminology (brand versus international nonproprietary name), and carry different regulatory expectations. Consumer GEO frameworks assume a single audience. Atima et al. (2022) show that predictive segmentation in regulated professional service settings resolves targeting inefficiencies that generic segmentation models leave in place, a finding that transfers directly to the patient and prescriber split in this channel.

These four constraints together explain why measurement instruments imported from unregulated categories underperform here. Sanni et al. (2020a) document the specific difficulty of measuring marketing return on investment in regulated services, where the permissible claim set and the observable outcome set are both narrowed by regulation, and Sanni and Attah (2023c) describe the parallel problem of constructing market research frameworks for highly regulated sectors where standard instruments are not usable without modification.

1.4 Contribution

This paper contributes:

1. A synthesis of the 2026 empirical evidence on AI search citation behavior in health, with attention to cross-engine divergence.
2. A specified benchmarking protocol, referred to here as the Therapeutic Visibility Benchmark, covering unit of analysis, prompt taxonomy, metric definitions, sampling design, and uncertainty quantification.
3. A pharmaceutical-specific metric layer covering fair balance surfacing, indication fidelity, and source authority, absent from existing commercial GEO measurement.
4. A discussion of the boundary between legitimate optimization and manipulation in a channel that distributes safety-critical information.

2. BACKGROUND AND RELATED WORK

2.1 Academic foundations of GEO

The field was named and formalized by Aggarwal and colleagues, whose work introduced GEO-Bench, a benchmark of roughly 10,000 diverse user queries, and a black-box optimization framework for improving a source's presence in generated answers (Aggarwal et al., 2024). Two of their findings remain load-bearing for pharmaceutical application.

First, they modeled a generative engine as a composition of retrieval and synthesis rather than as a ranking system, and proposed impression metrics suited to citations dispersed throughout a text rather than listed in rank order. Their objective measures combine the volume of answer text attributable to a source with a position weighting that discounts later placement, an approach subsequently formalized as Position-Adjusted Word Count and adopted across follow-on work (Fan et al., 2026; IF-GEO, 2026).

Second, their intervention experiments found that classical SEO signals, keyword density in particular, showed minimal influence on citation probability, while content-level characteristics associated with epistemic authority (presence of statistics, explicit source citations, and quotations from credible authorities) improved visibility by up to 40 percent. Their analysis also reported that a disproportionate share of citations, approximately 44 percent, derived from the first 30 percent of source content, which supports a front-loading rule for claim placement.

Subsequent work has extended the framework to e-commerce testbeds (Bagga et al., 2025), cooperative multi-source rewriting (Fan et al., 2026), and multi-objective feature-level page generation. A critical survey of the 2023 to 2026 literature characterizes the discipline as young, its measurement contested, and its long-run efficacy unproven, and notes that the no-interference assumption underlying most reported gains is violated in practice: normalized visibility is zero-sum, so when one source gains share another loses it, and when all competitors optimize simultaneously each actor's treatment alters the others' outcomes.

2.2 The measurement reliability problem

The most consequential recent methodological contribution for benchmarking is Sielinski's statistical framework for generative search measurement (Sielinski, 2026). The core argument is that generative engines are non-deterministic by design, that identical queries submitted at different times produce different responses and different cited sources, and that visibility metrics must therefore be treated as sample estimators of an underlying response distribution rather than as fixed system properties.

The empirical results are directly relevant to any pharmaceutical benchmarking programme:

- **Response-level instability.** Across repeated runs of the same query, median domain-level Jaccard overlap between cited source sets was approximately 0.29 to 0.31 for Gemini, 0.33 to 0.40 for SearchGPT, and 0.50 for Perplexity. Repeated identical queries frequently return substantially different evidence bases.
- **Citation volume varies by roughly sixfold across platforms.** Median citations per response ranged from about 5 to 7 for SearchGPT to about 36 to 40 for Gemini, with Perplexity intermediate. Raw citation counts are therefore not comparable across engines; normalized share is required.
- **Confidence intervals are wide relative to observed differences.** Bootstrap 95 percent intervals on citation share commonly spanned 3 to 7 percentage points, and overlapping intervals were the norm for domains differing by less than roughly 5 to 7 points. The paper's worked example shows two domains apparently separated by 3.5 points whose intervals overlap so substantially that the ranking is statistically indistinguishable from a tie.
- **Minimum sample sizes are engine-specific.** To achieve a 95 percent interval spanning five percentage points on citation share, the study estimated roughly 40 to 50 queries for Gemini, about 100 for Perplexity, and 150 or more for SearchGPT, the last reflecting within-sample non-stationarity that additional sampling does not smoothly resolve.
- **Rank orderings are unstable across the full distribution,** not only among closely ranked leaders, and short-window stability does not imply stability across longer windows.
- **A content-checksum control confirmed the variability is structural.** Cited page content was largely stable across the observation window, so the observed oscillation reflects engine retrieval behavior rather than changing source material.

The implication for pharmaceutical benchmarking is unambiguous. A brand visibility score reported without an interval, or derived from a single measurement occasion, is not a finding. It is a draw from a distribution. This is a specific instance of the general problem Wedraogo and Sanni (2024) address for cross-channel campaign performance, where forecasting accuracy degrades unless model uncertainty is carried explicitly rather than collapsed into a point value.

2.3 Commercial measurement practice and its gaps

A substantial vendor ecosystem now offers AI visibility tracking, and a parallel practitioner literature has emerged specifically addressing pharmaceutical application (Pharma Marketing Network, 2026). Common practice defines share of voice as the proportion of AI answers mentioning a brand relative to total brand mentions for a target query set, tracked per engine and over time (LLM Pulse, 2026b). Practitioner guidance converges on a small set of stable recommendations: measure per engine rather than as a blended figure, because the same brand can hold materially different share across engines in the same week; treat trend and relative momentum as more informative than absolute level; and separate mention rate (does the brand appear at all) from share of voice (how does its appearance compare to named competitors).

Reported category benchmarks in this literature should be treated cautiously. Figures such as a 30 percent share-of-voice target for category parity, or 15 percent as leadership in fragmented markets, are vendor-derived, rarely accompanied by uncertainty estimates, and drawn overwhelmingly from unregulated consumer and B2B categories. They are not validated for prescription therapeutics, where the competitive set is small, the terminology is bimodal (brand name versus molecule), and a meaningful fraction of answers legitimately name no brand at all because the clinically correct answer is a drug class.

Two structural weaknesses in current practice deserve explicit naming, because both have well-documented analogues in conventional marketing measurement. The first is attribution. Sanni et al. (2022a) review the bias introduced by attribution modeling across multi-touch journeys and find that method choice, rather than underlying performance, frequently determines which channel appears to be working. AI answer visibility is a

further touchpoint with no click, no identifier, and no timestamp, so it is the hardest case of the problem they describe rather than an exception to it. The second is executive reporting. Sanni and Atima (2021b) document how business intelligence dashboard design creates or closes visibility gaps in strategic marketing governance, which is precisely the failure mode of a single blended AI visibility figure presented to leadership without its interval.

Cross-engine behavioral differences reported in large vendor analyses are more robust and more useful. One February 2026 analysis of over 2.4 million responses reported wide variation in how often models name brands at all and whether they attach external links, with some models exhibiting very high brand-mention rates while including no external links, and others including links in a majority of responses while naming brands less frequently. Any composite visibility index must therefore be constructed so that a model's citation architecture does not masquerade as a brand's performance.

2.4 Adjacent literatures recruited by this framework

Pharmaceutical GEO has no mature native literature. The academic work summarized above is three years old, concerns unregulated categories, and contains no pharmaceutical application. The framework in Section 4 therefore borrows its measurement apparatus from six adjacent fields that have each solved some component of the problem in a different setting. This subsection states those debts explicitly, so that a reader can judge where the transfer is well founded and where it is an analogy under test.

2.4.1 Marketing measurement under regulatory constraint

The closest available analogue to pharmaceutical marketing measurement is analytics practice in other regulated professional sectors, where permissible claims are restricted, outcomes are partly unobservable, and compliance sits upstream of execution rather than downstream. A sustained body of work in this area addresses return-on-investment measurement, positioning, and go-to-market design under those constraints (Sanni et al., 2020a; Sanni et al., 2020b; Sanni et al., 2020c; Sanni and Atima, 2021a; Sanni and Attah, 2023c; Sanni and Attah, 2023a), and extends to demand and revenue forecasting in volatile service organizations (Sanni, 2026a; Sanni, 2026b).

Three threads within it bear directly on the metric design in Section 4. The first is attribution. Sanni et al. (2022a) review how method selection introduces bias in multi-touch journeys, and Sanni (2026d) document the resulting revenue attribution conflicts in regulated environments. The second is lifecycle instrumentation and automation control, where the relevant question is whether a marketing system's decisions remain inspectable after the fact (Sanni et al., 2025; Sanni et al., 2022b; Sanni et al., 2023; Sanni et al., 2024; Sanni et al., 2026). The third is the executive reporting layer, where Sanni and Atima (2021b) and Sanni (2026c) both find that compression of regulated-environment complexity into single indicators is the point at which governance visibility is lost. Related work on customer lifecycle modeling, telemetry-linked value measurement, and the interface shift implied by autonomous agents completes the picture (Ogundairo and Ayivi-Donkor, 2023a; Ogundairo and Ayivi-Donkor, 2023b; Ogundairo and Ayivi-Donkor, 2024a; Ogundairo and Ayivi-Donkor, 2024b), as does the funnel-leakage analysis in Sanni and Wedraogo (2024) and the segmentation work in Atima et al. (2022). The uncertainty treatment in Wedraogo and Sanni (2024) and the content optimization findings in Sanni and Attah (2023b) are used directly in Sections 4.4 and 5 respectively.

2.4.2 Health system data architecture and evidence infrastructure

Whether a therapeutic is retrievable at all is an architecture question before it is a content question. Work on health system transformation treats interoperability, data quality, and analytic capability as governance outcomes rather than technical ones, and traces the path from scientific evidence to delivered care (Ilodigwe and Adesemoye, 2021; Ilodigwe and Adesemoye, 2019b; Ilodigwe and Adesemoye, 2024a; Ilodigwe and Adesemoye, 2024b; Ilodigwe and Adesemoye, 2025b). The financing, access, and incentive literature in the same programme establishes why an information asymmetry in this channel has downstream consequences for utilization (Ilodigwe and Adesemoye, 2019a; Ilodigwe and Adesemoye, 2020; Ilodigwe and Adesemoye, 2026a; Ilodigwe and Adesemoye, 2026b), and Ilodigwe and Adesemoye (2025a) addresses the specific case of artificial intelligence in clinical decision-making.

A parallel body of work on decision support and performance measurement supplies the instrumentation vocabulary: outcome measurement at scale, predictive risk detection, and lifecycle performance management (Adelanwa et al., n.d.; Adelanwa et al., 2024; Adelanwa et al., 2023a; Adelanwa et al., 2023b), alongside audience modeling and traffic forecasting methods that transfer to prompt-set construction and cadence design (Basnet et al., 2021; Basnet et al., 2023; Basnet et al., 2024; Oghenemaiga et al., 2024). Resilience-oriented logistics work in the same group (Anene and Clement, 2022; Anene and Clement, 2024) is relevant to the access and availability questions in Section 4.2.

Physical infrastructure research completes this layer. Ogbete and Aminu-Ibrahim (2024) examines how infrastructure investment translates into measurable population health outcomes, a question structurally identical to the unmeasured-construct problem acknowledged in Section 7, and related work addresses diagnostic access expansion, accountability models, and system resilience (Aminu-Ibrahim and Ogbete, 2023; Aminu-Ibrahim et al., 2020; Aminu-Ibrahim et al., 2024; Aminu-Ibrahim et al., 2025b; Ogbete et al., 2026).

2.4.3 Governance, auditability, and privacy of AI systems in health

The governance layer of the framework draws on three overlapping literatures. The first concerns AI oversight in healthcare organizations, covering comparative accountability structures, organizational readiness, and privacy-preserving data governance (Afrihyia et al., 2024a; Afrihyia et al., 2025; Afrihyia et al., 2024b), with applied work on detection systems and equity-oriented service models (Afrihyia et al., 2023; Akinse et al., 2023a; Akinse et al., 2023b).

The second concerns the economics and control architecture of digital health systems: AI governance linked to privacy compliance and financial sustainability, audit quality, security investment, and the incident economics that determine whether a control is maintained (Ekechi et al., 2026; Ozowara et al., 2025a; Ozowara et al., 2022; Ozowara et al., 2025b; Adebayo et al., 2025; Adebayo et al., 2023), together with compliance architecture for distributed healthcare networks and identity and exchange frameworks (Ogunwola and Alozie, 2026; Kumuyi et al., 2023; Kumuyi et al., 2024b; Ekechi et al., 2022; Kumuyi et al., 2024a).

The third concerns data governance and regulatory traceability in enterprise settings, which supplies the operational vocabulary for the reporting standard in Section 4.7: data sensitivity classification, traceability mechanisms, continuous auditing, lakehouse governance, and comparative data

protection regimes across jurisdictions (Aliliele et al., 2024a; Aliliele et al., 2025a; Aliliele et al., 2025b; Mbonu et al., 2019; Aliliele et al., 2024b; Mbonu et al., 2022a; Mbonu et al., 2018; Mbonu et al., 2020; Mbonu et al., 2022; Aliliele et al., 2023a).

Legal scholarship frames the limits of all three. Annan (2024) sets out the tension between algorithmic accountability and trade secret protection, Annan (2023) the authorship problem for AI-generated works, Annan (2025a) automated decision-making under anti-discrimination law, and Annan (2022) cross-border platform privacy governance; Annan (2025b) treats compliance itself as a competitive position rather than a cost.

2.4.4 Pharmaceutical regulatory science, quality, and medicine access

The compliance metrics in Layer 3 assume an organization capable of acting on what they report. Regulatory science work on readiness maturity, accelerated approval trade-offs, cross-border harmonization, shortage response, and supply chain governance defines that capability (Eze et al., 2022b; Eze et al., 2024b; Eze et al., 2024a; Eze et al., 2022a; Eze et al., 2023).

The access literature establishes what an informational gap costs in practice. Work on affordability, last-mile distribution, continuity of care for chronic therapies, and manufacturing quality as a determinant of supply reliability (Asiedu and Asiedu, 2023b; Asiedu and Asiedu, 2022; Asiedu and Asiedu, 2024a; Asiedu and Asiedu, 2024b; Asiedu, 2025; Asiedu, 2026) is complemented by clinical and diagnostic work on interaction screening and diagnostic integration (Asiedu, 2022; Asiedu, 2023; Asiedu, 2024; Asiedu and Asiedu, 2023a), which defines the safety content that Section 4.3 measures for presence.

2.4.5 Language, literacy, and equity in health information

Section 3.3 reports that a brand's Answer Rate can shift by more than thirty points between languages. Interpreting that finding requires a literature on language rather than on marketing. Work on terminological equivalence in administrative and official translation shows that domain terminology does not map cleanly across languages even when both are well resourced (Kpoka, 2023; Kpoka, 2024), and work on multilingual identity, intersectionality, and cross-cultural communication addresses how language shapes access to institutional knowledge (Lilian et al., 2024; Lilian et al., 2025a; Lilian et al., 2025b; Lilian et al., 2020).

The procedural safeguards and accessibility literature contributes the equity framing used in Section 6.4: documented rights of access, compliance frameworks for service delivery, risk assessment for accountability systems, and technology-mediated accessibility (Oloto and Yeboah, 2022; Oloto and Yeboah, 2024; Oloto and Yeboah, 2025; Yeboah et al., 2024). Applied work on digital equity, data governance, and trust in health-adjacent learning and telehealth systems addresses the same question from the technology side (Orise and Niniola, 2024; Orise and Niniola, 2025; Orise and Niniola, 2023; Orise and Niniola, 2020; Orise and Niniola, 2022; Orise and Niniola, 2026; Akintola et al., 2025), and Fawehinmi et al. (2026) and Onwubuya et al. (2026) examine personalization and language acquisition in AI-mediated learning, which is the mechanism by which a generative system's register choices affect comprehension.

2.4.6 Forecasting and decision systems under uncertainty

Section 4.4 treats visibility as a distribution rather than a value. That position is standard in operational forecasting, where demand, cash flow, and cost models are built with explicit uncertainty and evaluated on interval calibration rather than point accuracy (Tonoyan et al., 2021a; Tonoyan et al., 2021b; Tonoyan et al., 2022c; Tonoyan et al., 2022b). Adjacent work on real-time performance tracking, valuation, geospatial market entry, routing, vendor risk, and brand portfolio architecture (Tonoyan et al., 2024b; Tonoyan et al., 2023; Tonoyan et al., 2022a; Tonoyan et al., 2025; Tonoyan et al., 2024a; Tonoyan et al., 2026) supplies the pattern for a monitoring programme that runs continuously rather than as a periodic study.

The transfer here is methodological, not substantive: none of this work concerns generative search. What it establishes is that a stochastic measurement problem with a commercial decision attached is a solved genre, and that the solution involves stating uncertainty rather than eliminating it.

3. THE 2026 AI SEARCH LANDSCAPE IN HEALTH

3.1 What generative engines actually cite for health questions

The largest published citation analysis specific to health queries examined 824,997 citations drawn from AI answers to generic, non-branded United States health questions, collected between mid-April and mid-July 2026 across ChatGPT, Google AI Mode, Google AI Overviews, Gemini, and Perplexity (LLM Pulse, 2026a). Pages owned by tracked brands were excluded so that the results reflect sources the engines reach for independently.

Source-type distribution across all engines:

Source type	Share of citations	Mean citation position
Commercial and platform	75.4%	7.25
Academic and journals	8.3%	6.56
User-generated and social	6.2%	8.53
Medical publishers	4.6%	5.71
Government and institutional	3.4%	5.51
News media	1.2%	6.83
Hospitals and providers	1.0%	7.38

Grouping government, academic, medical-publisher, and hospital sources together yields 17.3 percent of all health citations. The remaining four fifths flow to a long tail of commercial sites, applications, insurers, comparison sites, and platform domains.

Several findings from this dataset carry direct implications for therapeutic benchmarking:

The single most-cited domain is a research archive, not a health brand. PubMed Central accounted for approximately 33,669 citations, about 4.1 percent of the total, ahead of YouTube, google.com, and Reddit. For AI health answers, the primary literature outweighs the consumer health publishers with the strongest brand recognition.

Named consumer health brands are individually marginal. Healthline, Mayo Clinic, Cleveland Clinic properties, WebMD, and Medical News Today each held well under 1 percent of citations, all of them behind YouTube and Reddit. Domain authority alone does not secure visibility; the contest is at the level of individual pages and studies.

User-generated content is a top-tier health source. YouTube ranked second overall and Reddit fourth. Social and user-generated content collectively drew more citations than government sites and hospitals combined.

Position encodes an authority hierarchy. Government sources and medical references occupied the earliest citation slots on average (5.51 and 5.71 respectively), while social and user-generated content occupied the latest (8.53). Engines appear to treat authoritative sources as the lead and community content as supporting material. This is why position-weighted metrics, rather than raw mention counts, better approximate influence on the reader.

3.2 Engine divergence

The same study's per-engine breakdown is the clearest available evidence that a single blended visibility figure is uninterpretable:

Engine	Citations sampled	Authoritative medical	Social / UGC	Commercial / platform
ChatGPT	157,708	24.1%	1.9%	72.9%
Perplexity	43,334	18.4%	3.3%	77.7%
Gemini	50,595	16.8%	1.2%	80.9%
Google AI Overviews	249,784	15.8%	8.5%	74.4%
Google AI Mode	323,576	15.1%	7.6%	76.1%
All engines	824,997	17.3%	6.2%	75.4%

Two patterns dominate. ChatGPT is the most authority-first engine, sending nearly a quarter of its health citations to medical, academic, or government sources and rarely citing community platforms. Google's two search-integrated products sit at the opposite pole, citing social and user-generated content four to five times more often. Gemini, Google's standalone assistant, behaves differently from Google's search AI and cites social content least of any engine measured.

A methodological caveat applies: the underlying question set reflects the health topics that participating brands chose to monitor, so it is a large real-world sample rather than a census of health queries.

3.3 Brand-level divergence in therapeutic areas

Engine divergence at the source-type level translates into brand-level divergence of a magnitude that defeats naive reporting. The May 2026 PharmaGEO public index measured Answer Rate and Share of Voice for publicly named brands across four therapeutic areas (atopic dermatitis, obesity, psoriasis, and lung cancer) on three engines and in three languages (PharmaGEO, 2026). Two reported findings are instructive.

Engine gap. In atopic dermatitis, one branded biologic held an Answer Rate of 41.4 percent on one engine, a top-four category position, while scoring 8.2 percent on another engine in the same week using the same query set, a rank-eight niche presence. The 33.2-point gap between two engines measuring the same brand in the same therapeutic area implies there is no such thing as a single AI visibility score for a pharmaceutical brand.

Language gap. In the same therapeutic area, another branded biologic held an English-language Answer Rate of 13.8 percent (rank nine), which rose to 48.8 percent (rank four) in French and 51.9 percent (rank four) in Spanish. The index attributes this to approval sequencing and clinical coverage differences across geographies, which are reflected in the non-English source pool. A brand's AI visibility is therefore partly governed by the language of its approval geography rather than by its media investment.

These findings should be read with their provenance in mind: the index is vendor-produced and not peer-reviewed. They are nevertheless directionally consistent with the peer-reviewed evidence on engine-specific retrieval behavior, and the magnitude of the reported gaps is large relative to any plausible measurement noise floor.

3.4 Structural asymmetry: what pharma holds and does not use

The source archetypes that generative engines weight most heavily for health questions are precisely the content types the pharmaceutical industry already produces at scale: peer-reviewed publications, regulatory labels, medical society guidance, patient information leaflets, and real-world evidence registries. No other commercial sector holds a comparable volume of high-authority, retrieval-eligible material.

The constraint is infrastructural rather than editorial. Much of this material is distributed as PDF documents that retrieval systems parse inconsistently, or is scattered across third-party properties with no clean association to the brand entity. Ildigwe and Adesemoye (2021) frame the equivalent problem at health system scale as a data backbone question, arguing that interoperability, data quality, and analytic capability are governance and architecture outcomes rather than content outcomes. The same diagnosis applies here: the sponsor's disadvantage in the answer

layer is an architecture failure, not a shortage of evidence. The practical asymmetry is that a regulator-authored label document is frequently more retrievable, and more cited, than the sponsor's own structured content for the same product.

A qualitative mapping of retrievability against sponsor control:

Content type	Retrievability by current engines	Degree of sponsor control
Structured HTML on brand and HCP portals	High	Full
Medical society guideline pages	Very high	Influence only
PMC-indexed peer-reviewed publications	High	Partial, via publication strategy
Regulatory label documents	Very high	Indirect, via label update process
PDF publications and patient leaflets	Low to medium	Full but structurally underused
Patient forums and consumer health sites	Medium, varies by engine	None

Retrievability here is a qualitative judgement about current parsing behavior, not a measured score, and it shifts as engines improve.

4. A BENCHMARKING FRAMEWORK FOR COMMERCIAL THERAPEUTICS

This section specifies a protocol, referred to as the Therapeutic Visibility Benchmark (TVB), designed to be replicable by an independent team. It is presented as a specification rather than as results.

4.1 Unit of analysis

The primary unit is the **brand-engine-language-audience cell**. A single commercial therapeutic marketed in three languages and measured on four engines across two audience segments generates 24 distinct measurement cells, each of which must be estimated separately. Aggregation across cells is permitted only for reporting, never for inference.

Sanni et al. (2020b) note that differentiation in regulated professional markets depends on consistent, data-anchored positioning rather than on message volume, which is the reason entity consistency appears as a measured variable in Section 4.3 rather than as an assumption.

A secondary unit, the **molecule**, must be tracked in parallel. Generative engines frequently answer in international nonproprietary names rather than brand names, particularly for professional-intent queries and for products with generic competition. A benchmark that counts only brand-name mentions will systematically understate visibility for mature products and overstate the visibility gap between originator and generic-era assets. Brand-name and molecule-name mentions should be recorded as separate variables and reported both separately and combined.

4.2 Prompt universe and taxonomy

Prompt design is the single largest source of construct validity risk. A prompt set assembled around brand names will trivially return high brand visibility and measure nothing. The TVB prompt universe is stratified across six intent classes, with a target of at least 200 prompts per therapeutic area per language.

Intent class	Description	Target share	Example form
Disease-state, unbranded	Condition, symptoms, natural history, prognosis	20%	"What causes moderate to severe atopic dermatitis"
Treatment-selection, unbranded	Options, first line, sequencing, class comparison	25%	"What are the treatment options for moderate to severe plaque psoriasis"
Comparative	Named or implicit head-to-head	15%	"How do IL-23 inhibitors compare to TNF inhibitors"
Branded	Product named in the prompt	15%	"How does [brand] work"
Safety and tolerability	Adverse events, contraindications, interactions, monitoring	15%	"What are the serious risks of [class]"
Access and practical	Cost, coverage, administration, adherence	10%	"Is [class] covered by insurance"

Two design rules apply. First, prompts must be generated independently of the brands being measured, ideally by adapting real query logs, clinical FAQ inventories, and patient advocacy question sets, rather than by asking a language model to generate them. Sielinski (2026) explicitly flags model-generated query sets as a confound, since the generating model may carry implicit topic-to-source associations that shape the resulting distribution. Second, each prompt must be tagged for audience (patient-intent versus professional-intent), because vocabulary register alone materially changes the retrieved source pool. The segmentation and forecasting approach described by Basnet et al. (2021) provides a workable basis for constructing and weighting these audience strata, and Atima et al. (2022) supply the regulated-sector variant.

The access and practical intent class deserves particular attention in therapeutic benchmarking. Asiedu and Asiedu (2023b) document how supply chain inefficiency and affordability constraints shape what patients can actually obtain, which means that answers to cost and coverage questions carry real behavioral consequence even though they contain no efficacy claim and attract little promotional scrutiny.

4.3 Metric specification

The TVB metric set has three layers. Layers 1 and 2 are adapted from existing GEO and commercial practice. Layer 3 is specific to regulated therapeutics and, to the author's knowledge, is not implemented in any current commercial tool.

Layer 1: Visibility

Let P be the prompt set for a cell, R the set of responses collected, and $N = |R|$.

- **Answer Rate (AR).** The proportion of responses in which the brand is named at all. $AR = |\{r \in R : \text{brand named in } r\}| / N$. This is the primary presence metric and the closest analogue to the practitioner term "mention rate."
- **Share of Voice (SoV).** The brand's share of total named-brand surface within the competitive set, computed over responses rather than over prompts, so that a response naming three competitors is not scored as a full presence for each. $SoV = (\text{brand mentions}) / (\text{total mentions across the defined competitive set})$.
- **Citation Share.** The proportion of all citations in the sample that resolve to sponsor-owned domains. Reported separately from mention metrics, because a brand can be named frequently while owning no cited source.
- **Citation Prevalence.** The proportion of responses containing at least one citation to a sponsor-owned domain. Prevalence and share capture different dimensions (breadth versus depth) and are only moderately correlated; both should be reported.
- **Position-Adjusted Presence.** A position-weighted analogue of Aggarwal et al.'s Position-Adjusted Word Count, applying exponential decay to mentions and citations occurring later in an answer. Justified by the observed authority-position gradient in Section 3.1.

Normalization requirement. Because median citations per response differ by roughly sixfold across engines, raw citation counts must never be compared cross-engine. All cross-engine comparison uses share and prevalence.

Layer 2: Representation

- **Indication Fidelity.** The proportion of brand mentions in which the stated indication matches the approved label for the market being measured. Its complement is the **Off-Label Mention Rate**, a metric with no analogue in consumer GEO and direct regulatory significance.
- **Claim Accuracy.** The proportion of substantive claims about the brand that are consistent with the label and the underlying evidence base, adjudicated by a qualified reviewer against a pre-registered rubric. Requires human adjudication with a reported inter-rater reliability statistic.
- **Sentiment and Framing.** Classification of each mention as positive, neutral, or negative, plus a framing tag (first-line, alternative, last-resort, cautionary). Framing frequently carries more commercial signal than sentiment.
- **Entity Consistency.** Whether the brand, molecule, manufacturer, and class are correctly and consistently associated across responses. Entity fragmentation dilutes recognition across the whole channel.

Layer 3: Compliance and safety (pharmaceutical-specific)

- **Risk Information Surfacing Rate.** The proportion of responses containing an efficacy claim about the brand that also present material risk information. This operationalizes the fair balance concept from 21 CFR 202.1 as a measurable channel property. It is not a compliance judgement about the sponsor, who did not author the answer, but it is a direct measure of the channel's safety-information behavior and therefore a defensible pharmacovigilance signal.
- **Boxed Warning Presence.** For products carrying a boxed warning, the proportion of substantive responses in which it appears. Reported separately given its clinical weight.
- **Source Authority Ratio.** The proportion of citations in brand-relevant responses resolving to government, regulatory, academic, or medical-society sources. Section 3.1 establishes an all-health baseline of approximately 17.3 percent, permitting comparison of a therapeutic area against the general health corpus.
- **Hallucinated Attribution Rate.** The proportion of citations that do not support the claim to which they are attached, following the citation-faithfulness tradition established by Liu et al. (2023) for generative search engines and by Gao et al. (2023) for attributed generation. This requires manual verification on a sampled subset and is the most labor-intensive metric in the set. It is also the one with the clearest patient-safety implication.
- **Interaction and Contraindication Coverage.** For prompts that describe a concomitant medication, comorbidity, or pregnancy context, the proportion of responses that surface the relevant interaction or contraindication. Asiedu (2022) sets out a conceptual framework for systematic drug interaction screening in antenatal care and shows that unscreened interaction risk concentrates precisely in the populations least likely to receive specialist review. That is also the population most likely to be asking a general-purpose engine, which makes this a high-yield measurement target rather than an edge case.

Layer 2 and Layer 3 both depend on adjudication against an evidence base rather than against a keyword list. Ilodigwe and Adesemoye (2024a) describe the analogous requirement in real-world evidence programmes, where the value of a data-driven decision system rests on the quality of the reference standard against which outputs are scored.

4.4 Sampling design and uncertainty quantification

This is the component most often omitted in commercial practice and the one that determines whether results support inference.

Sample size. Following Sielinski (2026), the minimum prompt count per cell required for a 95 percent confidence interval spanning five percentage points on citation share is engine-dependent: on the order of 40 to 50 for engines with high per-response citation volume, approximately 100 for intermediate engines, and 150 or more for engines exhibiting within-sample non-stationarity. The TVB specifies 200 prompts per cell as a uniform floor, which satisfies the most demanding case observed.

Repeated measurement. Every cell is measured on at least nine occasions before any inference is drawn. Single-occasion measurement is prohibited. The evidence for this is direct: in the reference study, one domain produced a point estimate of 3.2 percent citation share on the first sampling occasion with a bootstrap interval of 2.4 to 4.2 percent, while the cross-sample mean across all subsequent occasions was 0.5 percent. A practitioner acting on the single sample would have concluded the domain was a leading source, a conclusion every later sample contradicted.

Interval estimation. Bootstrap resampling at the response level, with 1,000 replicates, recomputing both numerator and denominator on each replicate so that the ratio structure is preserved. All reported metrics carry 95 percent intervals. Any comparison whose intervals overlap is reported as statistically indistinguishable, not as a ranking.

Cadence and drift. Weekly collection, with quarterly span comparison against the earliest window. Consecutive-pair stability does not imply span stability; the reference study documented cases where every adjacent pair showed high rank correlation while the first-to-last comparison showed substantial cumulative drift. Both analyses are required.

Fixed-n commitment. Sample size is committed in advance. Using running interval width as a stopping criterion is unsafe, because interval width can narrow and then widen again as prompts are added under non-stationarity, causing a practitioner who stops at a favorable point to report an interval that understates true uncertainty.

Controls. Content checksums on cited URLs at each collection, so that observed variability can be attributed to engine behavior rather than to source-page edits. Prompt order randomization across occasions, to guard against within-sample non-stationarity driven by query clustering.

The general principle is the one Wedraogo and Sanni (2024) apply to campaign forecasting: where the measured system is stochastic, the deliverable is a distribution with stated uncertainty, and any process that reports a single number has discarded the information that determines whether the number is actionable.

4.5 Composite index construction

A single headline figure is frequently demanded by brand leadership and is frequently misleading. If one is constructed, the TVB specifies the following constraints:

1. The composite is computed **within cell** and never across engines, because engine citation architecture (mention propensity, link inclusion, citation volume) varies enough that a blended figure measures the engine mix as much as the brand.
2. Layer 3 compliance metrics enter as **flags rather than as weighted components**. Averaging a fair-balance measure into a marketing score creates an incentive structure in which safety-surfacing failures can be offset by visibility gains, which is indefensible in a regulated category.
3. Weights are pre-registered before data collection.
4. The composite is reported with its interval, and never reported alone.

Sanni (2026c) identifies the governance, risk, and compliance decision-making gap that arises when executive-facing marketing analytics compress regulated-environment complexity into a single indicator, and Sanni and Atima (2021b) describe the dashboard design conditions under which that compression misleads. Both argue for surfacing the constraint alongside the metric rather than behind it, which is the reasoning behind constraint 2 above.

4.6 Therapeutic area selection

For a benchmark intended to characterize the commercial therapeutics landscape, therapeutic area selection should span the structural variation that plausibly moderates AI visibility. The 2026 revenue landscape offers a natural sampling frame: oncology remains the largest therapy area by value, the metabolic and obesity category is the fastest-growing, and immunology contains several of the largest individual assets. Reported 2025 full-year figures place pembrolizumab at approximately 31.7 billion dollars, the tirzepatide franchise at roughly 36.5 billion dollars combined, and the semaglutide franchise at roughly 31.8 billion dollars, with dupilumab and the IL-23 class representing the immunology tier (BioSpace, 2026; Drug Discovery and Development, 2026; Genetic Engineering & Biotechnology News, 2026).

The TVB recommends stratifying across at least five areas selected on the following axes rather than on revenue alone:

Axis	Rationale	Contrasting areas
Competitive concentration	Zero-sum share dynamics differ between duopoly and fragmented categories	Obesity (concentrated) vs. plaque psoriasis (fragmented)
Audience mix	Patient-led versus prescriber-led information seeking	Obesity and dermatology (patient-led) vs. oncology (prescriber-led)
Consumer search intensity	Volume of unbranded lay-language queries	Metabolic and dermatology (high) vs. rare disease (low)
Direct-to-consumer advertising exposure	Prior media saturation shapes the mention corpus	US DTC-advertised brands vs. non-DTC markets
Off-label pressure	Where class use exceeds labelled indication	Metabolic (high) vs. oncology (protocol-bounded)
Generic and biosimilar status	Determines whether the answer names brands or molecules	Post-LOE assets vs. protected assets

4.6.1 A worked therapeutic area: hypertension

Hypertension is a useful area to specify concretely, because it exhibits an unusual combination of properties: very high patient-side search volume in lay language, a large and evolving guideline base, a mix of decades-old generics and recently approved agents, and substantial geographic variation in both prevalence and treatment pattern.

The evidence base a benchmark would adjudicate against illustrates what Layer 2 requires. Guideline and therapy reviews establish the label-consistent claim set for the class (Abah et al., 2025a; Abah et al., 2025b; David et al., 2025). Risk stratification work, including the triglyceride-glucose index literature and its cardio-renal extensions (Okwah, 2022; Amadi et al., 2025; Amadi et al., 2026a), defines the population subdivisions in which a treatment recommendation is or is not appropriate, which is precisely the distinction an Indication Fidelity score must capture. Epidemiological work on the contemporary burden of hypertension in specific populations (Amadi et al. (2026b)) shows why an English-language benchmark measured against United States prevalence assumptions would mis-specify the competitive set for other markets.

Newer agents make the temporal problem visible. Yusuf et al. (2026) report a phenotype-stratified network meta-analysis of aldosterone synthase inhibitors in resistant hypertension. A benchmark run before such evidence enters the retrievable corpus and one run after would measure different competitive sets for the same prompts, which is why Section 4.4 requires span comparison rather than only consecutive-pair stability. Adjacent work on inflammatory biomarkers, nurse-led intervention models, and multisite programme implementation (Okwah et al., 2026; Obi et al., 2025; Omo Enabulele et al., 2025) indicates the breadth of source material an engine may draw on when answering a question that appears narrowly pharmacological.

4.7 Reporting standard

Every published TVB result should state: the exact prompt set (published in full), collection dates and number of occasions, engine versions or access method, language and market, the competitive set definition, all metric definitions with formulas, sample sizes per cell, interval estimation method, and the adjudication rubric and inter-rater reliability for any human-scored metric. Absent these, a reported visibility figure is not replicable and should not enter commercial decision-making.

Two constraints on this standard should be acknowledged rather than assumed away. First, the measured systems are proprietary and their retrieval mechanics are not disclosed, so replication extends to the measurement procedure and not to the system under measurement. Annan (2024) examines this tension directly, arguing that algorithmic accountability and trade secret protection impose competing demands that cannot be resolved by transparency norms applied to the observer alone. Second, the measurement pipeline itself must be auditable if its outputs are to carry regulatory weight. Sanni et al. (2025) demonstrate a process mining approach to auditable marketing automation lifecycle control that offers a workable template: the record of what was collected, when, and under what configuration is treated as a first-class artifact rather than as operational exhaust.

5. DETERMINANTS OF VISIBILITY: WHAT THE EVIDENCE SUGGESTS MOVES THE NUMBER

The benchmark measures an outcome. This section summarizes what the published evidence indicates about its causes, with the caveat from Section 2.1 that intervention effects reported without confidence intervals may not exceed the measurement noise floor.

Epistemic authority markers. Aggarwal et al. found that adding statistics, explicit source citations, and quotations from credible authorities improved visibility by up to 40 percent, while keyword density had minimal effect. This is the strongest available signal, and it maps unusually well onto pharmaceutical content, which is already dense with citable evidence.

Claim placement. Approximately 44 percent of citations in the same study derived from the first 30 percent of source content. Medical content that buries the conclusion beneath methodology sections is structurally disadvantaged. Leading with the claim sentence is a low-cost intervention.

Content optimization as a measured discipline. Sanni and Attah (2023b) develop AI-assisted content optimization models aimed at conversion bottlenecks and report that the gains come from systematic structural revision rather than from volume. The mechanism they describe, in which explicit and well-formed claim structure outperforms additional content, is consistent with the GEO intervention findings above and suggests that the marginal return on a new asset is lower than the return on restructuring an existing one.

Format and parseability. PDF distribution is the most consistently identified structural barrier in the practitioner literature. Converting mechanism-of-action explainers, clinical summaries, and patient FAQs to structured HTML with explicit claim sentences is repeatedly identified as the highest-yield infrastructure change.

Mention consensus across independent sources. An analysis of roughly 75,000 brands reported a correlation of approximately $r = 0.664$ between web brand-mention volume and AI Overview citation (Ahrefs, as reported in PharmaGEO, 2026). The proposed mechanism is that retrieval systems use agreement across high-authority sources as a reliability proxy. A correlation of this size explains under half the variance and should not be read causally, but it is consistent with the practitioner observation that a single brand property is insufficient and that presence across society guidance, indexed literature, and specialist hubs matters more.

Language and approval geography. The PharmaGEO index findings in Section 3.3 suggest that the non-English source pool reflects approval sequencing, and that a brand concentrating GEO investment in English content may be competing in the wrong corpus for a substantial share of its European prescriber base.

Engine-specific source preference. Given that ChatGPT directs roughly 24 percent of health citations to authoritative medical sources while Google's search-integrated products cite community content four to five times more often, the content strategy that maximizes visibility on one engine is not the strategy that maximizes it on another. This is a strategic fork, not a tuning parameter, and it is one reason cross-engine composite scores are unsafe.

6. REGULATORY, ETHICAL, AND GOVERNANCE CONSIDERATIONS

6.1 The accountability gap

The governing framework for prescription drug promotion assumes a sponsor who authors, reviews, and disseminates a message. Generative answers break each assumption. The sponsor authors source content but not the answer; the answer is assembled at query time from sources the sponsor does not control; and dissemination is triggered by the user rather than the sponsor.

As of mid-2026, the FDA had not issued a rule specific to AI-generated promotional content, and the prevailing analysis is that existing requirements apply irrespective of what produced the material (XDS Health, 2026a). The compliance analysis for sponsor-operated patient and prescriber chatbots (XDS Health, 2026c) is the nearest existing guidance, though it addresses a surface the sponsor controls rather than one it does not. Enforcement activity in 2025 and 2026 has extended to influencer content, earned media, and patient testimonials, and Office of Prescription Drug Promotion letters in this period have focused on net impression and on risk information presented without adequate concurrency or prominence rather than on channel novelty. The reasonable inference for brand teams is that the agency is applying established doctrine to new interfaces rather than waiting to construct a separate one.

The accountability question has a close analogue in the wider governance literature. Afrihyia et al. (2024a) compare oversight structures for AI-driven healthcare management across the United States and developing health systems and find that accountability is most often lost at the boundary between the organization deploying a system and the organization that built it. The answer layer places a third party in that gap: the engine operator, who is neither the sponsor nor the regulator. A related question concerns authorship. Annan (2023) analyzes ownership and authorship in AI-generated works and shows that existing doctrine assumes a human author whose contribution can be identified, an assumption that a synthesized answer assembled from many sources does not satisfy. Akintola et al. (2025) reach a compatible conclusion for AI-powered remote monitoring, where user trust, privacy, and regulatory governance are found to be interdependent rather than separable compliance workstreams.

Two boundaries remain genuinely unsettled and should be flagged rather than papered over. First, whether sponsor-authored source content optimized for retrieval, and subsequently synthesized into an answer lacking fair balance, creates sponsor exposure. Second, whether a sponsor's knowledge of systematic off-label surfacing in a channel creates an affirmative obligation to act. Neither has been resolved by guidance or enforcement as of this writing, and practitioners should treat confident claims in either direction with skepticism.

6.2 Off-label surfacing as a measurable systemic property

The reported behavior of engines returning diabetes-indicated products in response to obesity queries (PharmaGEO, 2026) is a systemic property of the retrieval and synthesis process, reflecting the composition of the underlying corpus, rather than evidence of sponsor promotion. It is nonetheless consequential: each such mention functions as an indication expansion delivered outside promotional review, to prescribers and patients alike.

This is the strongest argument for locating pharmaceutical GEO measurement partly within medical affairs and pharmacovigilance rather than wholly within commercial. The Off-Label Mention Rate specified in Layer 3 is a monitoring instrument, not a marketing metric, and organizations should determine in advance what internal escalation a material reading triggers.

The organizational capability this requires is a regulatory readiness question rather than a tooling question. Eze et al. (2022b) propose a regulatory readiness maturity model for pharmaceutical companies, and its logic applies here: an organization that cannot describe its current state of monitoring, escalation, and remediation for this channel has no basis for asserting that its exposure is managed. The interaction screening framework in Asiedu (2022) supplies the clinical counterpart, defining what a systematic rather than incidental review of interaction risk requires.

6.3 Where optimization becomes manipulation

Consumer GEO literature treats visibility maximization as an unqualified goal. In a channel that distributes safety-critical information, it is not. A defensible boundary can be drawn on the basis of whether an intervention improves the answer's accuracy. Publishing structured, well-cited, accurate mechanism-of-action content in parseable HTML improves both visibility and answer quality. Seeding content designed to suppress risk information, to displace competitor mentions in comparative answers without evidentiary basis, or to shift the corpus toward unapproved uses degrades the channel while improving the metric.

The distinction matters more here than elsewhere because normalized visibility is zero-sum. When one brand gains share of a synthesized answer, another loses it, and the displaced content may be the risk information or the clinically superior alternative. The adversarial-manipulation literature on generative engines is explicit that model outputs can be systematically influenced; the difference between that work and legitimate optimization is intent and truthfulness, not technique.

Five governance recommendations follow:

1. **Route AI visibility content through existing medical-legal-regulatory review**, applying the same standard to retrieval-targeted content as to any other promotional or educational asset. Sanni et al. (2022b) describe adaptive control models for AI-driven marketing automation operating under financial compliance constraints, and the control architecture they specify, in which the compliance boundary is enforced in the system rather than at review, is the applicable pattern.
2. **Report Layer 3 metrics to medical and safety functions, not only to commercial**, with pre-defined escalation thresholds. Ekechi et al. (2026) set out a conceptual framework linking AI governance, data privacy compliance, and financial sustainability in digital health, and their argument that these cannot be governed in separate reporting lines applies to the split between visibility and safety metrics proposed here.
3. **Publish methodology**. Because the field's measurement is contested and vendor-dominated, a sponsor asserting a visibility position should be able to show the prompt set, cadence, and intervals behind it.
4. **Require explainability in any automated component**. Where classification, scoring, or content generation is automated within the measurement or optimization pipeline, the basis for each decision must be recoverable. Sanni et al. (2026) address this directly for marketing automation architectures in healthcare and financial applications, where an unexplainable classification is not usable as evidence in a regulatory conversation.

5. **Assess organizational readiness before scaling.** Afrihyia et al. (2025) find that readiness for generative AI integration in healthcare operations is a management capability rather than a technology procurement, and that capability gaps surface after deployment rather than during selection. Sanni and Atima (2021a) reach a parallel conclusion for analytics-driven go-to-market execution under compliance and sustainability complexity.

Where measurement or content workflows span organizational boundaries, for example between a sponsor, its medical communications agency, and a measurement vendor, Sanni et al. (2023) describe a lifecycle-aware architecture using federated learning that permits cross-organizational coordination without pooling the underlying data, which is a relevant pattern where the inputs include commercially sensitive or personally identifiable material.

6.4 Health equity

Two findings in the reviewed evidence carry equity implications that a purely commercial framing misses. First, the language effect documented in Section 3.3 means that the quality and completeness of AI-delivered treatment information varies systematically by query language, and non-English speakers may receive a materially different picture of available options. Second, the finding that community platforms rank among the most-cited health sources on Google's search-integrated products, while government and clinical sources dominate on ChatGPT, means that information quality in this channel now varies by which product a person happens to use, a variable strongly correlated with access, device, and default settings.

Both effects compound existing access inequities rather than operating independently of them. Asiedu and Asiedu (2022) document the barriers in last-mile drug distribution to underserved communities, and Asiedu and Asiedu (2024a) set out a continuity-of-care framework for equitable access to chronic therapies; in both cases the populations facing the largest structural barriers are those for whom an informational gap is least recoverable. Orise and Niniola (2024) report the same pattern in access to digital health education for underrepresented learners. Ilodigwe and Adesemoye (2026b) argue for treating health equity as a design parameter of system transformation rather than as a downstream measurement, which is the position adopted here: an equity term belongs in the benchmark specification, not in its discussion section.

There is also a discrimination-law dimension that pharmaceutical teams are unlikely to encounter through their usual promotional review. Annan (2025a) examines automated decision-making and anti-discrimination compliance under United States law and finds that disparate outcomes can arise without any discriminatory input, purely from differences in the training and retrieval corpus. A channel whose answer quality varies systematically by query language falls within the shape of that problem, whether or not any single actor is responsible for it.

7. LIMITATIONS

No primary data. This paper synthesizes published sources and specifies a protocol. It reports no original measurement, and the framework in Section 4 is untested at scale.

Source quality is uneven. The academic literature on GEO metrics and measurement uncertainty is peer-reviewed or preprint-with-methods. Several of the most pharmaceutically specific figures, including the therapeutic-area Answer Rate gaps and the AI Overview click-through decline, originate with vendors who sell measurement or optimization services and have a commercial interest in the magnitude of the effect. These are reported with attribution and should be read as directional evidence requiring independent replication rather than as established values.

Rapid obsolescence. Every figure here is a snapshot of engine behavior in the first three quarters of 2026. Engine architectures, retrieval corpora, and citation interfaces change on a timescale of months. The framework is intended to outlast the figures; the figures are not intended to be durable.

Geographic and linguistic narrowness. The largest citation dataset reviewed covers United States health queries. The regulatory discussion is United States-centric. EU pharmacovigilance obligations, the near-total prohibition on direct-to-consumer prescription advertising outside the United States and New Zealand, and non-Latin-script query behavior are all material and are not addressed. The gap is most consequential for markets undergoing regulatory convergence: Eze et al. (2024a) describe a cross-border harmonization model for African pharmaceutical markets in which approval status, and therefore the composition of the retrievable corpus, differs across jurisdictions that share a language. Annan (2022) raises the parallel governance problem for cross-border digital platforms, where a single system serves users under incompatible regulatory regimes. Asiedu (2026) situates both within a systems-level view of access, affordability, and supply chain resilience, which is the frame in which an information-availability metric becomes interpretable in a resource-constrained setting.

Attribution to commercial outcome is unestablished. Sanni (2026d) documents how revenue attribution conflicts arise in regulated professional services even for channels that are fully instrumented. The answer layer is not instrumented at all: there is no click, no session identifier, and no reliable referral signal for the majority of exposures. Any claim that a change in Answer Rate produced a change in commercial outcome should therefore be treated as a hypothesis rather than as a measurement.

The zero-sum problem is unresolved. If all competitors in a therapeutic area optimize simultaneously, the interference assumption underlying reported intervention effects fails, and the equilibrium outcome of an industry-wide GEO arms race in a safety-critical channel has not been modeled.

Unmeasured constructs. The framework measures presence, representation, and compliance surfacing. It does not measure whether appearing in an AI answer changes prescribing behavior, patient request behavior, or clinical outcomes. The commercial and clinical value of an Answer Rate point is currently unknown.

8. CONCLUSION

Three conclusions follow from the reviewed evidence.

First, the channel is real and materially different. With AI Overviews reported on roughly half of pharmaceutical searches, organic click-through on those queries falling by about 71 percent, and a majority of healthcare professionals reporting generative AI use in clinical contexts, the answer layer is not an emerging channel to be piloted. It is an operating information environment for prescription therapeutics that no sponsor authored and none controls.

Second, current measurement practice is inadequate to the decisions it supports. The dominant commercial pattern, a single blended visibility score derived from one measurement occasion, is refuted by the evidence on three independent grounds: engines diverge by more than 30 percentage points on the same brand in the same week; language flips category rank; and bootstrap intervals on citation share commonly span 3 to 7 points, rendering apparent differences below that threshold statistically indistinguishable from noise. Any benchmark intended to support investment decisions must be stratified by engine, language, and audience, repeated across occasions, and reported with intervals.

Third, pharmaceutical GEO cannot be governed as a marketing discipline alone. The answer layer distributes boxed warnings, contraindications, and off-label indication associations without promotional review. A measurement framework that captures only share of voice measures the least consequential property of the channel. The compliance layer specified here, covering risk information surfacing, indication fidelity, source authority, and attribution faithfulness, is where the distinctive pharmaceutical value of benchmarking lies, and it is precisely the layer that current commercial tooling does not provide.

The brands that establish accurate, well-sourced, retrievable representations of their products now will shape how their categories are described in this channel for years. The correction problem is asymmetric: an inaccurate default answer becomes harder to displace as it propagates through the corpus that trains and grounds successive engine generations. That asymmetry, more than any near-term traffic argument, is the case for treating this measurement seriously.

REFERENCES

- 1) Abah, M. A., Okwah, M. N., Oladosu, M. A., Yohanna, N. R., & Blessing, C. T. (2025a). Advances in the management of hypertension: A review of current guidelines, pharmacological therapies, and lifestyle interventions. *Drug and Pharmaceutical Science Archives*, 5(4), 35-41. <https://doi.org/10.47587/DAP.2025.5401>
- 2) Abah, M. A., Okwah, M. N., Oladosu, M. A., Yohanna, N. R., Blessing, C. T., & Vanessa, O. I. (2025b). The impact of lifestyle modifications on cardiovascular health: A review of diet, exercise, and stress reduction intervention. *Drug and Pharmaceutical Science Archives*, 5(3), 26-34. <https://doi.org/10.47587/DAP.2025.5301>
- 3) Adebayo, A., Adepoju, P. A., & Ozowara, D. E. (2023). Conceptual model for cyber risk pricing and insurance structuring in health technology platforms. *Gyanshauryam, International Scientific Refereed Research Journal*, 6(6).
- 4) Adebayo, A., Anunagba, C. O., & Ozowara, D. E. (2025). A review of zero trust security models and cost effectiveness in healthcare infrastructure. *International Journal of Advanced Multidisciplinary Research and Studies*, 5(6), 2269-2283. <https://doi.org/10.62225/2583049X.2025.5.6.6051>
- 5) Adelanwa, A., Basnet, A., & Anene, U. N. (2023a). Data driven digital transformation models for lifecycle performance management in infrastructure delivery. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), 2646-2662. <https://doi.org/10.62225/2583049X.2023.3.6.5968>
- 6) Adelanwa, A., Basnet, A., & Anene, U. N. (2023b). Predictive analytics models for financial risk detection and fraud prevention in public systems. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6). <https://doi.org/10.62225/2583049X.2023.3.6.5969>
- 7) Adelanwa, A., Basnet, A., & Anene, U. N. (2024). Performance intelligence models for optimization and outcome measurement in large scale public services. *Shodhshauryam, International Scientific Refereed Research Journal*, 6(1).
- 8) Adelanwa, A., Basnet, A., & Anene, U. N. (n.d.). Advanced AI-based decision support systems for healthcare operations and resource planning. *Gyanshauryam, International Scientific Refereed Research Journal*, 7(1).
- 9) Afrihyia, E., Akinse, S. G., & Ojukwu, P. U. (2023). Development and implementation of an AI-driven early detection system for non-communicable diseases. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), 2680-2691.
- 10) Afrihyia, E., Akinse, S. G., & Ojukwu, P. U. (2024a). Comparative governance of AI-driven healthcare management: Executive oversight, regulatory structures, and accountability in the United States and developing countries. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3087-3102. <https://doi.org/10.62225/2583049X.2024.4.6.5920>
- 11) Afrihyia, E., Akinse, S. G., & Ojukwu, P. U. (2025). Organizational readiness for generative AI integration in healthcare operations: Comparative management capabilities between the U.S. and low- and middle-income countries. *Iconic Research and Engineering Journals*, 8(10), 1673-1697. <https://doi.org/10.64388/IREV8I10-1714670>
- 12) Afrihyia, E., Ojukwu, P. U., & Akinse, S. G. (2024b). Privacy-preserving health data governance models: A comparative review of blockchain and cryptographic strategies in U.S. and developing healthcare systems. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3071-3086. <https://doi.org/10.62225/2583049X.2024.4.6.5919>
- 13) Aggarwal, P., Murahari, V., Rajpurohit, T., Kalyan, A., Narasimhan, K., and Deshpande, A. (2024). GEO: Generative Engine Optimization. *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD '24)*, pp. 5–16. ACM. [arXiv:2311.09735](https://arxiv.org/abs/2311.09735). <https://arxiv.org/abs/2311.09735>
- 14) Akinse, S. G., Ekechi, N. V., & Ohanebo, C. G. (2023a). Integrative maternal health model: Addressing socioeconomic inequities in rural populations. *International Journal of Advanced Multidisciplinary Research and Studies*, 3(6), 2832-2838.
- 15) Akinse, S. G., Ohanebo, C. G., & Ekechi, N. V. (2023b). A trauma-informed educational framework for adolescents in underserved communities. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 9(10), 434-443. <https://doi.org/10.32628/CSEIT2361073>
- 16) Akintola, A. S., Fawehinmi, Y. O., Aiyenitaju, O., Chinemerem, B., & Emmanuella, O. (2025). Digital transformation in telehealth: A systematic review of user trust, privacy, and regulatory governance in AI-powered remote monitoring systems. *Journal of Scientific Research and Reports*, 31(11), 163-182. <https://doi.org/10.9734/jsrr/2025/v31i113658>

- 17) Aliliele, C., Mbonu, I. S., & Iwuanyanwu, U. (2023a). A conceptual framework for continuous cloud misconfiguration monitoring and enterprise risk mitigation strategies. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 9(10), 373-394. <https://doi.org/10.32628/CSEIT2361071>
- 18) Aliliele, C., Mbonu, I. S., & Iwuanyanwu, U. (2024a). A conceptual framework for enterprise data sensitivity classification and regulatory traceability mechanisms. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3103-3124. <https://doi.org/10.62225/2583049X.2024.4.6.5991>
- 19) Aliliele, C., Mbonu, I. S., & Iwuanyanwu, U. (2024b). Advances in HIPAA compliant data architecture and secure analytics frameworks for community healthcare organizations. *Shodhshauryam, International Scientific Refereed Research Journal*, 7(2), 277-324. <https://doi.org/10.32628/SHISRRJ2472163>
- 20) Aliliele, C., Mbonu, I. S., Uzoka, E., & Iwuanyanwu, U. (2025a). A review of AI assisted continuous auditing systems in technology risk and cybersecurity oversight. *Gyanshauryam, International Scientific Refereed Research Journal*, 8(4), 210-250. <https://doi.org/10.32628/GISRRJ258369>
- 21) Aliliele, C., Mbonu, I. S., Uzoka, E., & Iwuanyanwu, U. (2025b). Advances in data lakehouse governance architectures for enterprise data loss prevention and compliance assurance. *Shodhshauryam, International Scientific Refereed Research Journal*, 8(4), 171-213. <https://doi.org/10.32628/SHISRRJ258474>
- 22) Amadi, C. E., Akanbi, P. O., Okwah, M. N., Mbakwem, A. C., & Ajuluchukwu, J. N. (2026a). Triglyceride-glucose index and reduced kidney function in hypertensive Nigerian adults: Independent association within the cardio-renal continuum. *BMC Cardiovascular Disorders*. Advance online publication. <https://doi.org/10.1186/s12872-026-05985-5>
- 23) Amadi, C. E., Ale, O. K., Okorafor, U. C., Okwah, M. N., Akanbi, P. O., Mbakwem, A. C., & Ajuluchukwu, J. N. (2026b). Unmasking the contemporary burden of hypertension in Lagos, Nigeria: Insights from an opportunistic screening. *BMC Public Health*. Advance online publication. <https://doi.org/10.1186/s12889-026-26310-x>
- 24) Amadi, C. E., Okorafor, U. C., Okorafor, C. I., Okwah, M. N., Akanbi, P. O., Okam, O. V., Udenze, C. I., Mbakwem, A. C., & Ajuluchukwu, J. N. (2025). Association between triglyceride-glucose index and estimated 10-year risk of cardiovascular disease in Nigerian non-diabetic hypertensives: A cross-sectional study. *PLOS Global Public Health*, 5(7), e0004760. <https://doi.org/10.1371/journal.pgph.0004760>
- 25) Aminu-Ibrahim, A. Y., & Ogbete, J. C. (2023). Healthcare infrastructure as a public health intervention using evidence from large laboratory networks. *Shodhshauryam, International Scientific Refereed Research Journal*, 6(1), 256-286. <https://doi.org/10.32628/SHISRRJ23678>
- 26) Aminu-Ibrahim, A. Y., Ogbete, J. C., & Ambali, K. B. (2020). Infrastructure driven expansion of diagnostic access across underserved and rural healthcare regions. *International Journal of Multidisciplinary Research and Growth Evaluation*, 1(5), 691-706. <https://doi.org/10.54660/IJMRGE.2020.1.5.691-706>
- 27) Aminu-Ibrahim, A. Y., Ogbete, J. C., & Ambali, K. B. (2024). Governance and accountability models for public private partnerships in healthcare infrastructure development. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 2943-2960. <https://doi.org/10.62225/2583049X.2024.4.6.5699>
- 28) Aminu-Ibrahim, A. Y., Ogbete, J. C., & Iwuanyanwu, O. C. (2025b). Infrastructure resilience planning for national diagnostic systems under public health stress conditions. *Gyanshauryam, International Scientific Refereed Research Journal*, 8(1), 340-381. <https://doi.org/10.32628/GISRRJ2582311>
- 29) Anene, U. N., & Clement, T. (2022). A resilient logistics framework for humanitarian supply chains: Integrating predictive analytics, IoT, and localized distribution to strengthen emergency response systems. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 8(5), 398-424.
- 30) Anene, U. N., & Clement, T. (2024). Localized supply chain solutions for sustainable community development: A strategic model for economic revitalization and regional resilience. *International Journal of Scientific Research in Science and Technology*, 11(5).
- 31) Annan, A. O. (2022). Data privacy governance models for cross-border digital platforms. *International Journal of Multidisciplinary Research and Growth Evaluation*, 3(6), 949-964.
- 32) Annan, A. O. (2023). Intellectual property ownership and authorship in AI-generated works. *Gyanshauryam, International Scientific Refereed Research Journal*, 6(3), 425-450.
- 33) Annan, A. O. (2024). Algorithmic accountability and trade secret protection in artificial intelligence. *Shodhshauryam, International Scientific Refereed Research Journal*, 7(5), 315-347.
- 34) Annan, A. O. (2025a). Automated decision-making and anti-discrimination compliance under U.S. law. *International Journal of Advanced Multidisciplinary Research and Studies*, 5(2), 2522-2540.
- 35) Annan, A. O. (2025b). Cybersecurity compliance as a source of competitive advantage in technology markets. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 11(5), 456-487.
- 36) Asiedu, C. S. (2022). Drug interaction risk in antenatal care: A conceptual framework for systematic interaction screening. *Shodhshauryam, International Scientific Refereed Research Journal*, 5(3), 465-484.
- 37) Asiedu, C. S. (2023). Diagnostic integration in infectious disease management: A review of microbiology, virology, and serology workflows. *International Journal of Medical Evaluation and Physical Report*, 7(4), 173-195. <https://doi.org/10.56201/ijmepr.v7.no4.2023.pg173.195>
- 38) Asiedu, C. S. (2024). Droplet digital PCR for HIV viral reservoir quantification: A review of diagnostic accuracy and clinical utility. *International Journal of Medical Evaluation and Physical Report*, 8(6), 253-273. <https://doi.org/10.56201/ijmepr.v8.no6.2024.pg253.273>

- 39) Asiedu, C. S. (2025). A lean process-improvement framework for pharmaceutical inventory and distribution management. *International Journal of Health and Pharmaceutical Research*, 10(12), 225-252. <https://doi.org/10.56201/ijhpr.vol.10.no12.2025.pg225.252>
- 40) Asiedu, C. S. (2026). Toward a systems-level framework for healthcare access, affordability, and supply chain resilience. *International Journal of Health and Pharmaceutical Research*, 11(5), 128-156. <https://doi.org/10.56201/ijhpr.vol.11.no5.2026.pg128.156>
- 41) Asiedu, C. S., & Asiedu, A. A. (2022). Last-mile drug distribution to underserved communities: A review of barriers and intervention models. *Gyanshauryam, International Scientific Refereed Research Journal*, 5(6), 439-466.
- 42) Asiedu, C. S., & Asiedu, A. A. (2023a). Comparative bioavailability of fresh versus dried botanical compounds: A review and cost-effectiveness modelling framework. *Research Journal of Pure Science and Technology*, 6(3), 226-244. <https://doi.org/10.56201/rjpst.v6.no3.2023.pg226.244>
- 43) Asiedu, C. S., & Asiedu, A. A. (2023b). Pharmaceutical supply chain inefficiencies and medication affordability in resource-constrained health systems: A narrative review. *International Journal of Health and Pharmaceutical Research*, 8(4), 192-222. <https://doi.org/10.56201/ijhpr.v8.no4.2023.pg192.222>
- 44) Asiedu, C. S., & Asiedu, A. A. (2024a). Equitable access to chronic therapies in hospital pharmacy settings: A conceptual framework for continuity of care. *International Journal of Medical Evaluation and Physical Report*, 8(6), 274-298. <https://doi.org/10.56201/ijmepr.v8.no6.2024.pg274.298>
- 45) Asiedu, C. S., & Asiedu, A. A. (2024b). Reconceptualizing quality assurance in pharmaceutical manufacturing as a determinant of supply reliability. *International Journal of Health and Pharmaceutical Research*, 9(5), 148-173. <https://doi.org/10.56201/ijhpr.v9.no5.2024.pg148.173>
- 46) Atima, M. E., Sanni, J. O., & Attah, A. (2022). Predictive audience segmentation models resolving targeting inefficiencies in regulated professional service enterprises. *Shodhshauryam, International Scientific Refereed Research Journal*, 5(1), 271-303. <https://doi.org/10.32628/SHISRJR247134>
- 47) Bagga, P. S., Farias, V. F., Korkotashvili, T., Peng, T., and Wu, Y. (2025). E-GEO: A Testbed for Generative Engine Optimization in E-Commerce. *arXiv:2511.20867*. <https://arxiv.org/pdf/2511.20867>
- 48) Basnet, A., Oghenemaiga, E., & Anene, U. N. (2021). Audience segmentation and forecasting models for enhancing targeted digital marketing effectiveness. *Shodhshauryam, International Scientific Refereed Research Journal*, 4(5), 279-305.
- 49) Basnet, A., Oghenemaiga, E., & Anene, U. N. (2023). Real time analytics and monitoring systems for livestream media performance using Google platforms. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 10(1).
- 50) Basnet, A., Oghenemaiga, E., & Anene, U. N. (2024). Attribution, revenue, and yield optimization models using Google Ad Manager reporting and forecasting tools. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6). <https://doi.org/10.62225/2583049X.2024.4.6.5667>
- 51) BioSpace (2026). Keytruda Hangs On to Best Seller Crown as GLP-1s Gain Ground. <https://www.biospace.com/business/keytruda-hangs-on-to-best-seller-crown-as-glp-1s-gain-ground>
- 52) CMI Media Group (2026). The Crucial Role of GEO and SEO in the Age of AI. <https://cmimediagroup.com/resources/the-crucial-role-of-geo-and-seo-in-the-age-of-ai/>
- 53) David, F. T., Abah, M. A., Oladosu, M. A., Agida, O. D., Clive, N. O., Adebukola, A. M., Nnabuko, O. M., & Maryanne, O. U. (2025). The impact of plant-based diets on cardiovascular health: A comprehensive review. *Plant Science Archives*, 10(4), 51-59. <https://doi.org/10.51470/PSA.2025.10.4.51>
- 54) Drug Discovery and Development (2026). Pharma 50: Keytruda holds the top brand slot, but tirzepatide and semaglutide have already passed it at the molecule level in FY2025. <https://www.drugdiscoverytrends.com/pharma-50-keytruda-holds-the-top-brand-slot-but-tirzepatide-and-semaglutide-have-already-passed-it-at-the-molecule-level-in-fy2025/>
- 55) Ekechi, N. V., Anunagba, C. O., & Ozowara, D. E. (2022). Advances in financial forensics techniques for detecting cyber-enabled fraud in telemedicine services. *Shodhshauryam, International Scientific Refereed Research Journal*, 5(1), 428-450.
- 56) Ekechi, N. V., Ozowara, D. E., & Anunagba, C. O. (2026). Conceptual framework for AI governance, data privacy compliance, and financial sustainability in digital health. *Computer Science and IT Research Journal*, 7(4), 275-299.
- 57) Eze, F. I., Akinleye, O. K., & Anene, U. N. (2024a). Development of a cross-border regulatory harmonization model for African pharmaceutical markets: The ARCH-Model framework. *International Journal of Health and Pharmaceutical Research*, 9(5), 148-186. <https://doi.org/10.56201/ijhpr.v9.no5.2024.pg148.186>
- 58) Eze, F. I., Akinleye, O. K., & Anene, U. N. (2024b). Development of a risk-based accelerated approval decision model balancing speed, safety, and evidence quality in pharmaceutical regulatory science. *Research Journal of Pure Science and Technology*, 7(6), 120-161. <https://doi.org/10.56201/rjpst.v7.no6.2024.pg120.161>
- 59) Eze, F. I., Anene, U. N., & Akinleye, O. K. (2022a). Creation of a drug shortage early warning and regulatory response model for essential medicines in global health systems. *Research Journal of Pure Science and Technology*, 5(1), 57-95. <https://doi.org/10.56201/rjpst.v5.no1.2022.pg57.95>
- 60) Eze, F. I., Anene, U. N., & Akinleye, O. K. (2022b). Design of an integrated regulatory readiness maturity model for pharmaceutical companies seeking global market entry. *International Journal of Health and Pharmaceutical Research*, 7(1), 77-112. <https://doi.org/10.56201/ijhpr.v7.no1.2022.pg77.112>
- 61) Eze, F. I., Anene, U. N., & Akinleye, O. K. (2023). Pharmaceutical supply chain governance, environmental risk management, and sustainable access to essential medicines in emerging economies: The PSCG-ERM framework. *International Journal of Health and Pharmaceutical Research*, 8(4), 192-241. <https://doi.org/10.56201/ijhpr.v8.no4.2023.pg192.241>

- 62) Fan, Z., Ghaddar, B., Wang, X., Xing, L., Zhang, Y., and Zhou, Z. (2026). What Generative Search Engines Like and How to Optimize Web Content Cooperatively. arXiv:2510.11438. <https://arxiv.org/pdf/2510.11438>
- 63) Fawehinmi, Y. O., Okereke, G. M. T., Williams, P., Chijioke-Churuba, J. N., & Asare-Badu, B. (2026). AI-human collaboration in education and psychology: Personalised learning, language acquisition and student support systems. *Journal of Political Science and International Relationship*, 3(1), 54-64. <https://doi.org/10.54536/jpsir.v3i1.6998>
- 64) Gao, T., Yen, H., Yu, J., and Chen, D. (2023). Enabling Large Language Models to Generate Text with Citations. *Proceedings of EMNLP 2023*, pp. 6465–6488. arXiv:2305.14627
- 65) Genetic Engineering & Biotechnology News (2026). Top 10 Best-Selling Drugs 2026. <https://www.genengnews.com/industry-news/top-10-best-selling-drugs-2026/>
- 66) IF-GEO (2026). Conflict-Aware Instruction Fusion for Multi-Query Generative Engine Optimization. arXiv:2601.13938. <https://arxiv.org/pdf/2601.13938>
- 67) Iloigwe, L., & Adesemoye, A. C. (2019a). A critical review of health insurance financing mechanisms and provider payment reform strategies for balancing coverage expansion and financial protection in emerging markets. *International Journal of Health and Pharmaceutical Research*, 5(2), 31-55. <https://doi.org/10.56201/ijhpr.vol.5.no2.2019.pg31.55>
- 68) Iloigwe, L., & Adesemoye, A. C. (2019b). From molecules to health systems: A conceptual model explaining why scientific innovation fails to improve patient outcomes without effective healthcare delivery infrastructure. *International Journal of Health and Pharmaceutical Research*, 5(2), 56-79. <https://doi.org/10.56201/ijhpr.vol.5.no2.2019.pg56.79>
- 69) Iloigwe, L., & Adesemoye, A. C. (2020). Advances in pharmacy-led delivery and access models for strengthening primary healthcare systems and expanding medicine availability in emerging markets. *International Journal of Health and Pharmaceutical Research*, 5(3), 78-96. <https://doi.org/10.56201/ijhpr.vol.5.no3.2020.pg78.96>
- 70) Iloigwe, L., & Adesemoye, A. C. (2021). The data backbone of health system transformation: A governance and architecture framework for interoperability, data quality, and advanced analytics at national scale. *International Journal of Health and Pharmaceutical Research*, 6(2), 52-76. <https://doi.org/10.56201/ijhpr.vol.6.no2.2021.pg52.76>
- 71) Iloigwe, L., & Adesemoye, A. C. (2024a). Advances in real-world evidence, predictive analytics, and data-driven decision making for improving health system performance and patient outcomes. *International Journal of Health and Pharmaceutical Research*, 9(5), 174-197. <https://doi.org/10.56201/ijhpr.v9.no5.2024.pg174.197>
- 72) Iloigwe, L., & Adesemoye, A. C. (2024b). Beyond technology: A strategic framework for building sustainable and resilient health systems through organizational strategy, operational excellence, and workforce capability. *International Journal of Medical Evaluation and Physical Report*, 8(6), 299-319. <https://doi.org/10.56201/ijmepr.v8.no6.2024.pg299.319>
- 73) Iloigwe, L., & Adesemoye, A. C. (2025a). Advances, risks, and implementation challenges of artificial intelligence as a force multiplier for clinical decision making and health system efficiency. *International Journal of Medical Evaluation and Physical Report*, 9(7), 165-185. <https://doi.org/10.56201/ijmepr.v9.no7.2025.pg165.185>
- 74) Iloigwe, L., & Adesemoye, A. C. (2025b). Executing healthcare transformation at scale: A strategic change model derived from large programs across payers, providers, and integrated health systems. *International Journal of Health and Pharmaceutical Research*, 10(12), 253-272. <https://doi.org/10.56201/ijhpr.vol.10.no12.2025.pg253.272>
- 75) Iloigwe, L., & Adesemoye, A. C. (2026a). A conceptual framework review of incentive alignment mechanisms between payers and providers in value-based care contracting and their effects on cost and quality outcomes. *Journal of Public Administration and Social Welfare Research*, 11(1), 111-133. <https://doi.org/10.56201/jpaswr.v11.no1.2026.pg111.133>
- 76) Iloigwe, L., & Adesemoye, A. C. (2026b). Reimagining population health: An integrated conceptual model combining value-based care, digital health innovation, and health equity for sustainable system transformation. *International Journal of Health and Pharmaceutical Research*, 11(5), 157-176. <https://doi.org/10.56201/ijhpr.vol.11.no5.2026.pg157.176>
- 77) IQVIA (2026). The Evolution of Pharma Engagement as AI Becomes a Front Door to Medical Information. <https://www.iqvia.com/locations/emea/blogs/2026/03/the-evolution-of-pharma-engagement-as-ai-becomes-a-front-door-to-medical-information>
- 78) Kpoka, M. I. (2023). Comparative analysis of administrative French in Francophone African states and its impact on translation practice. *Journal of Public Administration and Social Welfare Research*, 8(2), 149-163. <https://doi.org/10.56201/jpaswr.v8.no2.2023.pg149.163>
- 79) Kpoka, M. I. (2024). Terminological challenges in the translation of civil registry and administrative documents. *International Journal of Social Sciences and Management Research*, 10(11), 552-578. <https://doi.org/10.56201/ijssmr.v10.no11.2024.pg.552.578>
- 80) Kumuyi, O., Akeju, B., Uzoka, E., & Ozowara, D. E. (2023). Framework for blockchain-based cross-border data exchange and regulatory transparency. *International Journal of Multidisciplinary Futuristic Development*, 4(1), 99-113.
- 81) Kumuyi, O., Uzoka, E., Akeju, B., & Ozowara, D. E. (2024a). Architecture for machine learning-enabled predictive energy management using IoT sensor networks. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3138-3149. <https://doi.org/10.62225/2583049X.2024.4.6.6004>
- 82) Kumuyi, O., Uzoka, E., Akeju, B., & Ozowara, D. E. (2024b). Framework for privacy-focused digital identity verification supporting financial inclusion in Africa. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3150-3162. <https://doi.org/10.62225/2583049X.2024.4.6.6005>
- 83) Lilian, I. N., Liadi, K. O., & Yeboah, T. J. (2024). Multilingual ecofeminism and environmental justice: Rethinking gender, language, and ecology in global contexts. *Journal of Humanities and Social Policy*, 10(6), 195-222. <https://doi.org/10.56201/jhsp.v10.no6.2024.pg195.222>

- 84) Lilian, I. N., Liadi, K. O., & Yeboah, T. J. (2025a). Intersectionality, language, and identity: Multilingual perspectives on inclusion and social justice. *Iconic Research and Engineering Journals*, 8(9), 1761-1788. <https://doi.org/10.64388/IREV8I9-1713695>
- 85) Lilian, I. N., Liadi, K. O., & Yeboah, T. J. (2025b). Language and collaboration across the humanities and social sciences: Addressing global challenges through multidisciplinary dialogue. *Iconic Research and Engineering Journals*, 9(1), 2024-2050. <https://doi.org/10.64388/IREV9I1-1713696>
- 86) Lilian, I. N., Liadi, K. O., Yeboah, T. J., & Apelehin, A. A. (2020). Understanding cross-cultural communication: Identity, diversity, and global interaction. *Gyanshauryam, International Scientific Refereed Research Journal*, 3(4), 153-176. <https://doi.org/10.32628/GISRRJ21352>
- 87) Liu, N. F., Zhang, T., and Liang, P. (2023). Evaluating Verifiability in Generative Search Engines. *Findings of EMNLP 2023*, pp. 7001–7025. arXiv:2304.09848
- 88) LLM Pulse (2026b). Share of Voice in AI Search: How to Calculate It in 2026. <https://llmpulse.ai/blog/share-of-voice-ai-search/>
- 89) LLM Pulse (2026a). What Sources Does AI Trust for Health Questions? 825,000 Citations Analyzed. <https://llmpulse.ai/blog/ai-health-sources-analyzed/>
- 90) Mbonu, I. S., Aliliele, C., Iwuanyanwu, U., & Oluoha, O. M. (2018). A conceptual framework for legal and ethical risk modeling in enterprise data protection governance systems. *Iconic Research and Engineering Journals*, 2(2), 207-226. <https://doi.org/10.64388/IREV2I2-1714911>
- 91) Mbonu, I. S., Aliliele, C., Iwuanyanwu, U., & Uzoka, E. (2020). A review of identity and access management integration strategies in hybrid and multi cloud environments. *International Journal of Multidisciplinary Research and Growth Evaluation*, 1(5), 795-810. <https://doi.org/10.54660/IJMRGE.2020.1.5.795-810>
- 92) Mbonu, I. S., Aliliele, C., Iwuanyanwu, U., & Uzoka, E. (2022). A conceptual framework for AI enabled IT general controls and SOX audit automation processes. *Gyanshauryam, International Scientific Refereed Research Journal*, 5(5), 384-414. <https://doi.org/10.32628/GISRRJ2256239>
- 93) Mbonu, I. S., Aliliele, C., Uzoka, E., & Oluoha, O. M. (2019). A review of comparative data protection regulations and secure cloud implementation strategies across jurisdictions. *Iconic Research and Engineering Journals*, 2(9), 482-501. <https://doi.org/10.64388/IREV2I9-1714912>
- 94) Mbonu, I. S., Iwuanyanwu, U., Aliliele, C., & Uzoka, E. (2022a). A review of data protection impact assessment models in multi cloud financial infrastructure systems. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 8(1), 589-623. <https://doi.org/10.32628/CSEIT25442>
- 95) Obi, L. I., Dogbanya, G., Awah, L. C., Tawose, O. M., Ubani, C., Ayo-ige, A. B., Odedele, M., & Edefeade, O. (2025). A systematic analysis of effectiveness of nurse-led dementia care interventions on health outcomes among community-dwelling older adults. *Journal of Pharma Insights and Research*, 3(5), 24-33. <https://doi.org/10.69613/3f5cq717>
- 96) Ogbete, J. C., & Aminu-Ibrahim, A. Y. (2024). Translating healthcare infrastructure investment into measurable population health and diagnostic outcomes. *International Journal of Scientific Research in Humanities and Social Sciences*, 1(2), 955-985. <https://doi.org/10.32628/IJSRSSH242775>
- 97) Ogbete, J. C., Aminu-Ibrahim, A. Y., & Iwuanyanwu, O. C. (2026). Scaling molecular diagnostic facilities through standardized infrastructure and operational design models. *World Scientific News*, 212, 220-258. <https://doi.org/10.65770/ZUIY6924>
- 98) Oghenemaiga, E., Basnet, A., & Anene, U. N. (2024). Predictive analytics applications for forecasting traffic patterns across owned and operated media platforms. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6). <https://doi.org/10.62225/2583049X.2024.4.6.5668>
- 99) Ogundairo, K. M., & Ayivi-Donkor, S. S. (2023a). Machine learning approaches to customer acquisition and retention: A framework for predictive lifecycle management. *IIARD International Journal of Economics and Business Management*, 9(10), 210-246. <https://doi.org/10.56201/ijebm.v9.no10.2023.pg210.246>
- 100) Ogundairo, K. M., & Ayivi-Donkor, S. S. (2023b). Telemetry-driven value realization: Infusing economic evaluation algorithms directly into enterprise data pipelines. *International Journal of Economics and Financial Management*, 8(8), 195-232. <https://doi.org/10.56201/ijefm.v8.no8.2023.pg195.232>
- 101) Ogundairo, K. M., & Ayivi-Donkor, S. S. (2024a). Bridging the "adoption gap": Why supply chain AI fails without product marketing principles. *International Journal of Marketing and Communication Studies*, 8(5), 169-205. <https://doi.org/10.56201/ijmcs.v8.no5.2024.pg169.205>
- 102) Ogundairo, K. M., & Ayivi-Donkor, S. S. (2024b). The vanishing interface: Why the coming era of autonomous agents demands a paradigm shift in product marketing architecture. *IIARD International Journal of Economics and Business Management*, 10(11), 350-385. <https://doi.org/10.56201/ijebm.v10.no11.2024.pg350.385>
- 103) Ogunwola, T. A., & Alozie, C. (2026). Architecting secure and compliant distributed healthcare networks: Operational approaches to health insurance portability and accountability act and health information trust alliance alignment. *International Journal of Computer Applications*, 187(111), 1-6.
- 104) Okwah, M. N. (2022). Integrating triglyceride-glucose index and echocardiographic parameters for improved cardiovascular risk stratification in Sub-Saharan Africa. *International Journal of Cardiology*, 1(2), 1-16. https://doi.org/10.34218/IJC_01_02_001
- 105) Okwah, M. N., Abah, M. A., Oladosu, M. A., Blessing, T. B., Anku, A. I., & Agida, O. D. (2026). Anti-inflammatory effects of omega-3 fatty acids: Evidence from circulating biomarkers and inflammatory gene-expression studies. *African Research Reports*, 2(4), 365-382. <https://doi.org/10.65221/0167>

- 106) Oloto, A. M., & Yeboah, T. J. (2022). A comprehensive review of IDEA compliance frameworks in modern special education service delivery systems. *Gyanshauryam, International Scientific Refereed Research Journal*, 5(6), 323-348.
- 107) Oloto, A. M., & Yeboah, T. J. (2024). A critical review of procedural safeguards and regulatory compliance in special education programs. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 3125-3137.
- 108) Oloto, A. M., & Yeboah, T. J. (2025). Developing conceptual risk assessment models for special education documentation and accountability systems. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 11(4), 691-734. <https://doi.org/10.32628/CSEIT251116283>
- 109) Omo Enabulele, A. B., Oyebamiji, I., Ikubanni, O., Okwah, M. N., Babalola, W. S., Azeez, O. M., & Olaniyi, A. O. (2025). A coordinated project-management approach to multisite implementation of motor-rehabilitation programs for children with autism spectrum disorder in United States healthcare systems: A narrative review. *Cureus*, 17(12), e100513. <https://doi.org/10.7759/cureus.100513>
- 110) Onwubuya, L. I., Fawehinmi, Y. O., Sunday, D., Igwenagu, M. O., Adeniran, A. O., & Agyapong, E. A. (2026). Adaptive learning systems powered by artificial intelligence for STEM education improvement. *Asian Journal of Education and Social Studies*, 52(6), 553-566. <https://doi.org/10.9734/ajess/2026/v52i63115>
- 111) Orise, C., & Niniola, S. (2020). Data-informed digital learning platforms: A conceptual framework for evidence-based educational technology in healthcare and STEM environments. *International Journal of Multidisciplinary Research and Growth Evaluation*, 1(5), 1078-1089. <https://doi.org/10.54660/IJMRGE.2020.1.5.1078-1089>
- 112) Orise, C., & Niniola, S. (2022). Real-time analytics for educational excellence: A conceptual framework for performance dashboards in health professional training. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 8(2), 802-815.
- 113) Orise, C., & Niniola, S. (2023). Privacy, safeguarding, and ethical data practices in healthcare EdTech: A conceptual model for compliance and trust. *Shodhshauryam, International Scientific Refereed Research Journal*, 6(1), 507-529.
- 114) Orise, C., & Niniola, S. (2024). Digital equity in healthcare education: A systematic review of access frameworks and interventions for underrepresented learners in clinical training. *International Journal of Scientific Research in Humanities and Social Sciences*, 1(2), 1233-1246.
- 115) Orise, C., & Niniola, S. (2025). Data governance in digital health learning systems: A systematic review of implementation frameworks and best practices. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 11(5), 520-538.
- 116) Orise, C., & Niniola, S. (2026). From classroom to clinic: Digital transformation in health professional education and training delivery. *International Journal of Health and Pharmaceutical Research*, 11(6), 52-64. <https://doi.org/10.56201/ijhpr.vol.11.no6.2026.pg52.64>
- 117) Ozowara, D. E., Adebayo, A., & Anunagba, C. O. (2022). A systematic review of cybersecurity investments and their impact on healthcare financial performance. *Shodhshauryam, International Scientific Refereed Research Journal*, 5(1), 404-427.
- 118) Ozowara, D. E., Adebayo, A., & Anunagba, C. O. (2025a). Advanced conceptual model for strengthening audit quality using data analytics across financial institutions. *International Journal of Advanced Multidisciplinary Research and Studies*, 5(6), 2246-2268. <https://doi.org/10.62225/2583049X.2025.5.6.6050>
- 119) Ozowara, D. E., Anunagba, C. O., & Adepoju, P. A. (2025b). A review of ransomware economics and financial resilience strategies in hospital networks. *International Journal of Advanced Multidisciplinary Research and Studies*, 5(6), 2284-2298. <https://doi.org/10.62225/2583049X.2025.5.6.6052>
- 120) Pharma Marketing Network (2026). Generative Engine Optimization in Pharma Explained. <https://www.pharma-mkting.com/articles/generative-engine-optimization-in-pharma/>
- 121) PharmaGEO (2026). Why Pharma Needs GEO in 2026: Generative Engine Optimization for Life Sciences. <https://pharma-geo.com/articles/why-pharma-needs-geo-2026>
- 122) Sanni, J. O. (2026a). Artificial intelligence driven demand forecasting models addressing volatility in complex service organizations. *Computer Science & IT Research Journal*, 7(1), 1-26. <https://doi.org/10.51594/csitrj.v7i1.2170>
- 123) Sanni, J. O. (2026b). Integrated advertising and marketing framework for U.S. retail growth. *International Journal of Marketing and Communication Studies*, 10(1), 12-36. <https://doi.org/10.56201/ijmcs.v10.no1.2026.pg12.36>
- 124) Sanni, J. O. (2026c). Marketing analytics frameworks addressing governance risk compliance decision making gaps for executives. *International Journal of Marketing and Communication Studies*, 10(1), 37-64. <https://doi.org/10.56201/ijmcs.v10.no1.2026.pg37.64>
- 125) Sanni, J. O. (2026d). Predictive marketing frameworks resolving revenue attribution conflicts in regulated professional services environments. *World Scientific News*, 212, 52-81. <https://doi.org/10.65770/QPHC9342>
- 126) Sanni, J. O., & Atima, M. E. (2021a). Analytics driven go to market frameworks addressing compliance sustainability complexity. *International Journal of Multidisciplinary Research and Growth Evaluation*, 2(6), 647-660.
- 127) Sanni, J. O., & Atima, M. E. (2021b). Business intelligence dashboard frameworks resolving executive visibility gaps in strategic marketing governance. *International Journal of Multidisciplinary Research and Growth Evaluation*, 2(6), 633-646. <https://doi.org/10.54660/IJMRGE.2021.2.6.633-646>
- 128) Sanni, J. O., & Attah, A. (2023a). A comprehensive framework for digital transformation in capital markets: Solving operational challenges and enhancing stakeholder engagement. *Gyanshauryam, International Scientific Refereed Research Journal*, 6(6), 275-302. <https://doi.org/10.32628/GISRRJ236640>
- 129) Sanni, J. O., & Attah, A. (2023b). Artificial intelligence assisted content optimization models resolving digital conversion performance bottlenecks issues. *Gyanshauryam, International Scientific Refereed Research Journal*, 6(6), 303-336. <https://doi.org/10.32628/GISRRJ236641>

- 130) Sanni, J. O., & Attah, A. (2023c). Market research frameworks addressing entry barriers within highly regulated industrial service sectors. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 10(2), 904-930. <https://doi.org/10.32628/CSEIT2342439>
- 131) Sanni, J. O., & Wedraogo, L. (2024). Data-centric funnel optimization frameworks resolving revenue leakage in high-value services. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 2859-2874. <https://doi.org/10.62225/2583049X.2024.4.6.5641>
- 132) Sanni, J. O., Ajiga, D., & Atima, M. E. (2020a). Analytical models addressing measurement challenges of marketing return on investment in regulated services. *International Journal of Multidisciplinary Research and Growth Evaluation*, 1(5), 636-648. <https://doi.org/10.54660/IJMRGE.2020.1.5.636-648>
- 133) Sanni, J. O., Ajiga, D., & Atima, M. E. (2020b). Data driven brand positioning frameworks resolving differentiation challenges in regulated professional markets. *International Journal of Multidisciplinary Research and Growth Evaluation*, 1(5), 649-660. <https://doi.org/10.54660/IJMRGE.2020.1.5.649-660>
- 134) Sanni, J. O., Ajiga, D., & Atima, M. E. (2020c). Systematic review of product management strategies in mobile network rollouts across emerging markets. *International Journal of Multidisciplinary Research and Growth Evaluation*, 1(5), 661-673. <https://doi.org/10.54660/IJMRGE.2020.1.5.661-673>
- 135) Sanni, J. O., Atima, M. E., & Attah, A. (2022a). Systematic review of attribution modeling methods resolving bias in multi-touch journeys. *Shodhshauryam, International Scientific Refereed Research Journal*, 5(1), 304-334. <https://doi.org/10.32628/SHISRRJ247135>
- 136) Sanni, J. O., Iwuanyanwu, U. A., & Essien, M. A. (2025). Problem-oriented process mining for auditable marketing automation lifecycle control. *International Journal of Advanced Multidisciplinary Research and Studies*, 5(6), 1933-1947. <https://doi.org/10.62225/2583049X.2025.5.6.5650>
- 137) Sanni, J. O., Iwuanyanwu, U. A., & Essien, M. A. (2026). Designing explainable AI based marketing automation architectures for healthcare and financial applications. *World Scientific News*, 213, 119-148. <https://doi.org/10.65770/MJFK8593>
- 138) Sanni, J. O., Iwuanyanwu, U. A., Essien, M. A., & Attah, A. (2023). Lifecycle-aware marketing automation using federated learning for secure cross-organizational data management. *Gyanshauryam, International Scientific Refereed Research Journal*, 6(6), 337-364. <https://doi.org/10.32628/GISRRJ236642>
- 139) Sanni, J. O., Iwuanyanwu, U. A., Essien, M. A., & Wedraogo, L. (2024). Integrating blockchain-enabled smart contracts for transparent and verifiable marketing workflows. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 2891-2906. <https://doi.org/10.62225/2583049X.2024.4.6.5646>
- 140) Sanni, J. O., Iwuanyanwu, U. A., Essien, M. A., Atima, M. E., & Attah, A. (2022b). Adaptive control models for AI-driven marketing automation in financial compliance environments. *Shodhshauryam, International Scientific Refereed Research Journal*, 5(1), 243-270. <https://doi.org/10.32628/SHISRRJ247133>
- 141) Sielinski, R. (2026). Quantifying Uncertainty in AI Visibility: A Statistical Framework for Generative Search Measurement. *arXiv:2603.08924*. <https://arxiv.org/pdf/2603.08924>
- 142) Spectrum Science (2026). 2026 Predictions for Healthcare and Pharmaceutical Marketing. <https://www.spectrumscience.com/perspectives/2026-pharma-marketing-predictions/>
- 143) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2021a). Advances in demand forecasting: Machine learning algorithms for revenue projection and financial planning accuracy. *Gyanshauryam, International Scientific Refereed Research Journal*, 4(1), 282-317.
- 144) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2021b). Predictive cash flow modeling in retail: A review of advanced forecasting and working capital optimization. *Shodhshauryam, International Scientific Refereed Research Journal*, 4(3), 366-397.
- 145) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2022a). Advances in geospatial analysis: Predictive analytics methods for store location selection and market entry return on investment. *International Journal of Marketing and Communication Studies*, 6(2), 78-105. <https://doi.org/10.56201/ijmcs.v6.no2.2022.pg78.105>
- 146) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2022b). Algorithmic process optimization and cost reduction: A review of simulation modeling and financial impact assessment. *Journal of Accounting and Financial Management*, 8(8), 139-169. <https://doi.org/10.56201/jafm.vol.8.no8.2022.p139.169>
- 147) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2022c). Conceptual model for predictive cost optimization: Machine learning applications in business transformation analysis. *International Journal of Economics and Business Management*, 8(6), 68-102. <https://doi.org/10.56201/ijebm.v8.no6.2022.pg68.102>
- 148) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2023). A conceptual framework for retail asset valuation: Integrating quantitative pricing and multi-unit financial performance prediction. *Journal of Accounting and Financial Management*, 9(12), 218-250. <https://doi.org/10.56201/jafm.v9.no12.2023.pg218.250>
- 149) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2024a). Advances in supply chain resilience: Predictive models for vendor risk assessment and procurement cost optimization. *International Journal of Social Sciences and Management Research*, 10(11), 525-551. <https://doi.org/10.56201/ijssmr.v10.no11.2024.pg525.551>
- 150) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2024b). Real-time KPI tracking systems: A review of automated performance monitoring and data-driven decision making. *World Journal of Innovation and Modern Technology*, 8(6), 247-281. <https://doi.org/10.56201/wjimt.v8.no6.2024.pg247.281>
- 151) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2025). A conceptual model for real-time route optimization: Algorithmic pathways to delivery profitability and logistics cost reduction. *World Journal of Innovation and Modern Technology*, 9(12), 308-364. <https://doi.org/10.56201/wjimt.v9.no12.2025.pg308.364>

- 152) Tonoyan, A., Dada, O., & Ayivi-Donkor, S. S. (2026). Conceptual framework for financial brand architecture optimization: Portfolio analysis and economic value maximization. *Gulf Journal of Advance Business Research*, 4(3), 145-164.
- 153) United States Code of Federal Regulations. 21 CFR Part 202, Prescription Drug Advertising; 21 CFR 202.1.
- 154) Wedraogo, L., & Sanni, J. O. (2024). Machine learning models addressing uncertainty in cross-channel campaign performance forecasting accuracy. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 2875-2890. <https://doi.org/10.62225/2583049X.2024.4.6.5648>
- 155) XDS Health (2026a). Fair Balance in AI-Generated Pharma Content: How MLR Teams Review LLM Outputs in 2026. <https://blog.madebyxds.com/fair-balance-ai-generated-pharma-content-2026>
- 156) XDS Health (2026b). AI-Generated Pharma Content and FDA Compliance. <https://blog.madebyxds.com/ai-generated-pharma-content-fda-compliance-2026>
- 157) XDS Health (2026c). Pharma Chatbot Compliance: FDA Patient and HCP Bot Guide. <https://blog.madebyxds.com/pharma-chatbot-compliance-patient-hcp-bots>
- 158) Yeboah, T. J., Bobga, M. A., Boakye, K., & Ogbona, C. S. (2024). Technology integration in special education: Enhancing student engagement and accessibility. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 2978-2993.
- 159) Yusuf, I. A., Badero, O. J., Ihesiulo, A., Okwah, M. N., Oshodin, A., & Asikong, E. (2026). Aldosterone synthase inhibitors in uncontrolled and resistant hypertension: A phenotype-stratified systematic review and network meta-analysis of randomized trials. *PLOS ONE*, 21(6), e0349932. <https://doi.org/10.1371/journal.pone.0349932>