

# Digital Distrust: Patient-Centred Communication and the Calibration of Epistemic Reliance in the Era of Online Self-Diagnosis and AI Symptom Checkers

Nonye Tochi Aghanya<sup>1\*</sup> and Osinakachi Akuma Kalu<sup>2</sup>

<sup>1</sup>Independent clinician-scholar, Virginia, United States.

<sup>2</sup>Founder & Board Chair, TAFFD's and Afrolongevity.

## \*Corresponding Authors

**Nonye Tochi Aghanya, MSc, FNP-C,**  
Independent clinician-scholar, Virginia, United States.  
ORCID: <https://orcid.org/0009-0004-9426-8012>

**Submitted:** 6 Aug 2026; **Accepted:** 27 Aug 2026; **Published:** 3 Sept 2026

**Citation:** Aghanya, N. T. & Kalu, O. A. (2026). Digital Distrust: Patient-Centred Communication and the Calibration of Epistemic Reliance in the Era of Online Self-Diagnosis and AI Symptom Checkers. *J Medical Case Repo* 8(5):1-13.  
DOI : <https://doi.org/10.47485/2767-5416.1162>

## Abstract

**Background:** Patients are coming to consultations informed by search engines, social media and generative AI. Clinicians may characterize this as distrust of clinicians. Actualities are messier, and our language for talking about how communication breaks down is limited.

**Purpose:** Operationalize the concept of digital distrust, disambiguate it from related concepts with which it is often confused, and report a communication framework grounded in theory for clinical practice.

**Methods:** Methodology is an interpretive conceptual integration, described per guidelines for conceptual articles and interpretive description. We performed an iterative, theory-driven integration of peer-reviewed scholarly literature from digital health, health communication, judgement and decision-making, and bioethics scholarship, interpreted by the first author's practice-grounded understanding developed across over thirty years of clinical nursing practice. No primary data were generated or analyzed.

**Findings:** Five themes emerged from our analysis. Pre-informed consultations are normalized. Conversational fluency is becoming dissociated from clinical credibility. Checking behavior is varied depending on information source, with conversational artificial intelligence (AI) less often double-checked than search engines. Rather than merely possessing divergent epistemic authorities, patients navigate them. Patients also hold back information to preserve the relationship. The interpersonal exchange is where trust is appraised or misappraised.

**Implications:** We reconceptualize patients' digital distrust as an epistemic recalibration gone awry at the margins of an otherwise benign revolution in information access, rather than a generalized loss of faith in clinicians. We propose five working categories of digitally mediated stance that patients take during consultations, as well as the Digital Calibration Protocol (DCP), a four-step communication maneuver (Acknowledge – Clarify – Educate – Reframe).

**Conclusion:** Blanket acceptance or rejection of patients' digital content is indefensible. Patient-centered communication provides a mechanism for recalibrating patients' reliance. Here, we introduce the DCP as a hypothesis-generating framework in need of empirical testing.

**Keywords:** patient-centered communication; artificial intelligence; symptom checkers; epistemic trust; shared decision-making; health information seeking; large language models

## Introduction

A patient sits down and before anything else, says, “I looked it up”. Her phrasing gives the clinician insight into what kind of consultation this will be: whether she brings information to work through together or an answer against which the next twenty minutes will be judged.

Clinicians know the difference. Little in the communication literature will help them name it. The phenomenon is not new. Large language models are not intelligent and generative assistants (LLM-IGAs) or chatbots. In mid-2024, a nationally

representative Australian survey found roughly 1 in 10 adults had used a general-purpose conversational AI assistant for a health question in the past six months. Use was higher among those with low health literacy and those born in a non-English speaking country (Ayre et al., 2025). Intentions data looking at willingness to use conversational agents for self-diagnosis is similarly directional (Shahsavari & Choudhury, 2023). These systems build on three decades of online health information seeking which research has repeatedly investigated for impacts on clinical relationships. On balance, past findings suggest

information seeking is more helpful than harmful to clinical communication (Luo et al., 2022; Tan & Goonawardene, 2017). Patients largely want clinicians to help them understand what they learn online, rather than be usurped by it (Sommerhalder et al., 2009).

Comments to the contrary have abounded over the past year. Scholars and journalists have suggested AI symptom checkers and chatbots are creating widespread distrust of clinicians. The research does not support this. Previous literature, mixed data signals from LLM-IGAs, and reasonable base rates counsel skepticism towards broad claims of a crisis. If you have encountered COVID-19 patients with full-blown Capgras delusions about their doctors online, you have not looked very hard.

Something else is emerging. Observations from decision sciences highlight how an interface can affect judgements about credibility and intent (Christensen, 2023). LLM-IGAs list the correct diagnosis first in 19–38% of standardized case vignettes depending on the system being tested (Wallace et al., 2022). While large language models achieve state-of-the-art scores of 84–90% on medical licensing-style benchmarks, these rocket ship metrics mask considerably weaker performance on clinically focused tasks produced by healthcare practitioners (45–69%) and decrease further on patient safety- critical assessments (Gong et al., 2025). Users rarely double-check answers provided by chatbots, while they often double-check search engine results (Yun & Bickmore, 2025). Clinicians do not need to assume patients trust machine-generated information implicitly. They need only recognize that patients may judge it less skeptically than they do — while its reliability is lower.

One patient at a time. Here is the problem concisely stated in a recent mental-health themed commentary: Chatbots cannot reliably distinguish what symptoms patients report as concerning from those a provider would treat as urgent, and patients can arrive at an appointment already framed by a diagnosis (Kaneda, 2026).

What follows are three questions and attempts to answer them. What is digital distrust? How can clinicians recognize it? And how should clinicians respond? At the end of this essay is a definition of digital distrust, a provisional taxonomy of orientations clinicians might hold toward digitally informed patients, and the Digital Calibration Protocol (DCP). We present no validation data for these three, only that each can be defensible and tested. That is a different and prior standard.

## Conceptual Background

### Digital health information seeking and the consultation

Empirical research on health information seeking on the internet does not portray a broken relationship. Tan and Goonawardene's (2017) review found that patients' web searching most commonly supplemented clinical consultations instead of replacing them. Luo et al.'s (2022) meta-analysis of 53 studies found that 58% described positive effects on the

physician–patient relationship, 26% described no effect, and 15% described harmful effects. Through these studies, patients' online information seeking did not undermine physicians' roles when information was brought up in the clinical encounter. Importantly, the studies on which that conclusion is based are skewed toward high-income countries. There are too few to cover the globe systematically, and too few from outside North America to claim representativeness. Furthermore, the percentages I quote describe the spread of conclusions drawn by individual studies, not the magnitude of any particular effect.

The best account of why remains Sommerhalder et al.'s (2009) qualitative work. Study participants reported being confused by information they found online, and turning to their physician to make sense of it. These physicians appreciated having that starting point for a conversation with patients, but were troubled when patients became attached to misinformation. Patients sometimes elected not to mention the information they found online, to avoid undermining their physician's authority or because they were short on time.

I mention that last point because it has larger implications than may be obvious. If a physician never knows what information a patient has encountered, that information cannot be corrected or clarified. How wrong these tools can be helps illustrate the danger. First passes at diagnosis are correct only 19–37.9% of the time, and within three guesses only 33–58% of the time (Wallace et al., 2022). Triage recommendations are somewhat more accurate, but those numbers come from only two sources and range from 48.8–90.1%. (Semigran et al., 2015) found similar performance several years ago. They are not perfectly inaccurate—they help many patients—but are far more wrong than you would realize from using them.

### AI-Mediated Health Information

Generative systems alter that encounter in four ways: they generate rather than retrieve answers; they generate in fluent prose; they personalize their responses to the wording of a question; and, as far as consumers are concerned, they typically don't display source material.

Much here should sound familiar. Fluency isn't incidental. Chatbot-generated responses have been found to participants' questions about mental health to be perceived as more empathetic than human-generated ones, using validated measures (Howcroft et al., 2025). In meta-analysis across 13 studies, the standardized mean difference was 0.87 (95% CI 0.54–1.20), to simplify somewhat, approximately two points on a typical ten-point rating scale. The finding replicates in earlier work comparing responses to patient questions (Ayers et al., 2023). Both studies carry important limitations: most studies in the meta-analysis used non-validated rating scales, and both analyzed text responses only. This is a finding about perceived warmth of writing, not healthcare. The consistency is still suggestive, and fluency of processing is known to affect judgements of credibility: information that can be processed easily is evaluated more positively on criteria unrelated to processing difficulty (Alter & Oppenheimer, 2009).

Experimental studies press the case. Mayer and colleagues simulated a medical interview (Mayer et al., 2026) and found participants reported lower subjective stress under the belief that they were interacting with AI versus the human-physician condition. Participants also rated an AI bot as similarly credible to that of a human clinician, yet recalled the least information from the chatbot.

That combination of results — lower subjective stress and higher (or in this case equivalent) credibility alongside poorer learning — is cause for concern. The interaction feels easier and no less legitimate but contains less useful information.

Clinical validity does not appear to have kept pace. Examining results across 39 benchmarks, Gong and colleagues find that language models perform significantly worse on clinical tasks compared to exam-style questions, with gaps largest for items related to safety (Gong et al., 2025). Systems-level safeguards against health misinformation generation have been inconsistently adopted by, and disappeared from, major models between releases (Menz et al., 2024).

Patient interpretation complicates whatever conclusion we might draw about indiscriminate credulity. Wardle and colleagues tested participant reactions to AI-generated health information (Wardle et al., 2025). Users reported using AI tools while explicitly questioning them, using outputs as information to be built upon rather than stated truths, and criticizing chimp for not citing sources. Yun and Bickmore similarly report high use of search engines for health information, relatively low use of conversational systems, but better cross-checking of search results against knowledgeable sources than chatbots (Yun & Bickmore, 2025). Reported rate of adhering to chatbot recommendations was lower.

Taken together, I don't think these findings fit the frame they're normally called upon to justify. Consumers aren't hyping the benefits; they're unevenly skeptical. And that skepticism isn't linked to reliability but presentation format.

### Trust, Mistrust and Epistemic Authority

Precision here requires some pedantry, because "digital distrust" has sometimes been invoked casually. **Medical mistrust** is an enduring attitude about health-care entities and personnel, rooted in structural/historical context and measured with psychometrically validated tools (Williamson & Bigman, 2018). It tends to be stable, it sticks to organizations, and it is often justified. Digital distrust is not those things. **Epistemic trust**, as theorized by Fonagy and Allison (2014), is the openness to consider information received from another person as applicable and universal; epistemic vigilance refers to a set of filtering mechanisms to avoid accepting misinformation (Sperber et al., 2010). Neither of these latter terms implies something that should be heightened. They are positions on continuums to be adjusted as needed. This is the space our metaphorical child seeks a place in.

**Trust in AI** is meaningfully divisive. Hatherley (2020) writes that AI systems are systems we can rely on, but not systems we

can trust. The qualities that make an algorithm trustworthy, for Hatherley, include intentions and purposes that AI lacks. I lean on the distinction between relying on and trusting someone throughout. **Algorithm aversion & algorithm appreciation** refer to broad predispositions to discount algorithmic guidance (Dietvorst et al., 2015; Logg et al., 2019). The related resistance to medical AI seems to be driven by patients' perceptions that the algorithm won't be able to identify what's unique about them (uniqueness neglect; Longoni et al., 2019). **Automation bias** captures the over-use of decision support tools, but that work mostly addresses how physicians interact with software (Goddard et al., 2012). There's simply no consensus term for what happens when a patient with distrust based in prior digital experiences interacts with a human doctor.

### Cyberchondria, health literacy and interpretive overload

Cyberchondria – online searching that amplifies health concern – was defined by White & Horvitz (2009), and linked to health anxiety meta-analytically by McMullan et al. (2019). Cyberchondria is a process of anxiety amplification. A patient may experience digital distrust calmly – with no anxiety whatsoever – if they are confident in their position; by contrast, an anxious patient may comply without hesitation. These are different constructs that don't always co-occur.

eHealth literacy refers to the ability to find, understand, and use online health information (Norman & Skinner, 2006). It's relevant to who is susceptible to calibration inaccuracy, but it isn't a behavior. Uncertainty management theory is more actionable: After reviewing existing theories, Brashers (2001) concluded that people address uncertainty by either seeking or avoiding information, depending on whether uncertainty has positive or negative implications for them. Digital seeking and digital avoidance are both strategies for managing uncertainty, which is why they can be driven by the same mechanism.

The circumstances influencing access also vary. In online health information access in Africa, researchers conducted a systematic review of studies from Africa and found that device ownership, cost of internet connectivity, e-health literacy, and living in rural areas were the most common barriers identified (Chereka et al., 2025).

### A Proposed Definition

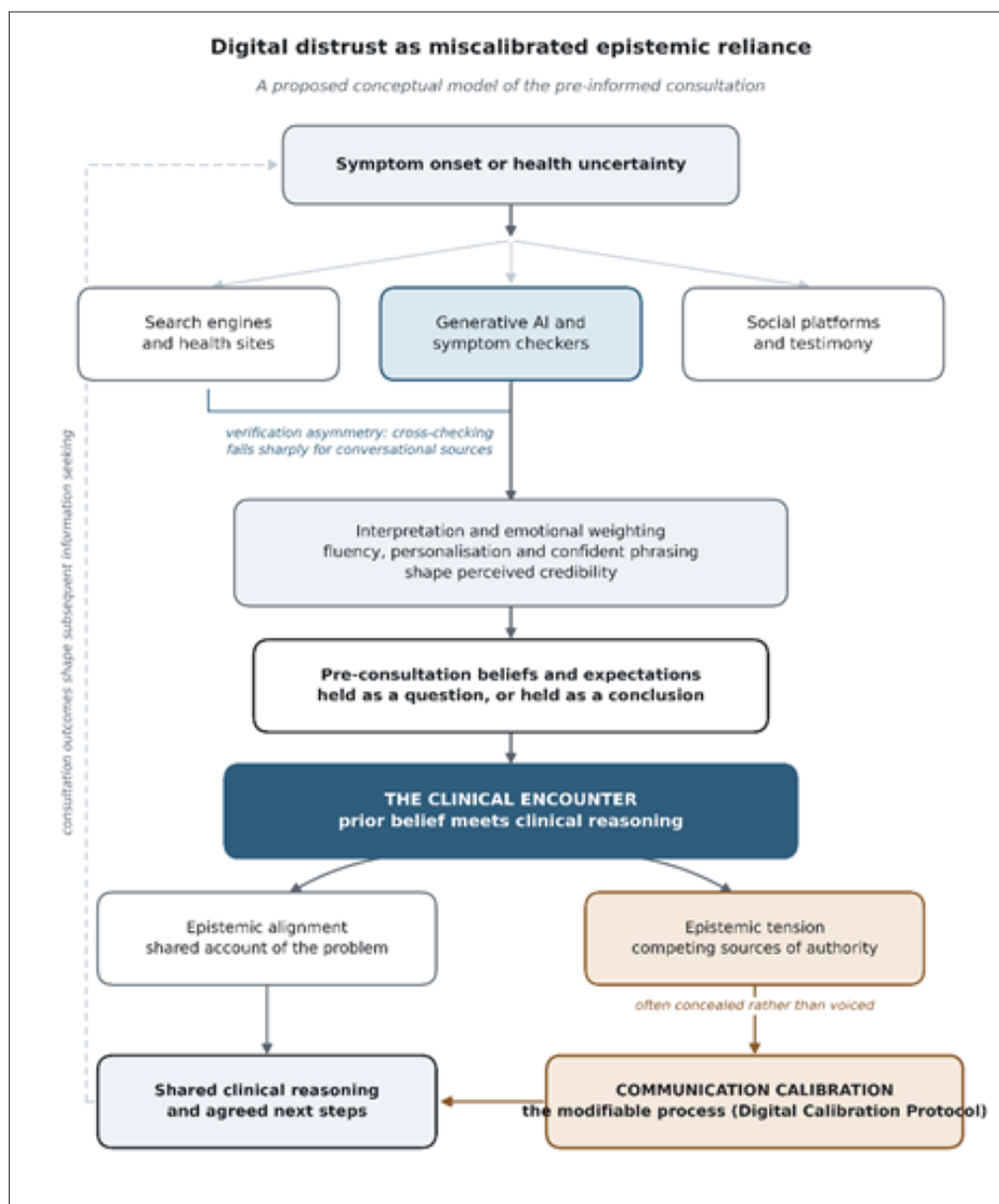
We propose that digital distrust can be operationalized as follows: We define digital distrust as *an interactional state of the clinical encounter wherein past digital experience with medical information has conferred upon a nonhuman agent epistemic authority that is considered (often implicitly) alongside and sometimes above the clinician's clinical reasoning in a way that interferes with collaborative sensemaking around history-taking, diagnosis, and/or treatment planning*.

Three points of clarification. First, digital distrust is a property of an encounter, not a patient or clinician trait. Second, digital distrust is orthogonal to interpersonal trust. A patient could very well hold great esteem for his clinician and Google as a medical authority at the same time. Third, clinical correctness does not determine digital distrust. It is possible for a patient to

access and recount the correct information and still defer to it in a manner that is incongruent with shared clinical decision-making. This suggests a fourth clarification: Digital distrust is continuous, not binary.

The mechanism responsible for digital distrust, we argue, is miscalibrated epistemic reliance. Reliance on a source that is disproportionate to the demonstrated reliability of that source for the task at hand. From this follows our paper's key conceptual distinction: digital engagement versus digital overreliance. Digital engagement encompasses preparation for the clinical encounter, accessing and understanding medical vocabulary, sensemaking orientation, describing symptoms,

and remembering your question once you've entered the clinical space. Digital overreliance is substituting an eloquent but uninformed opinion for clinical judgment. Digital engagement should be encouraged. Only digital overreliance should be addressed by the clinician. Treating the former as the latter is a communication issue. Figure 1 displays our proposed model. Information seeking is prompted by uncertainty and subsequent output is integrated given its perceived correctness and affective valence. Pre-consultation beliefs form that eventually participate in the clinical encounter that will end in either consonance or discord. Communication calibration is the moment of intervention.



**Figure 1:** Proposed conceptual model of digital distrust as miscalibrated epistemic reliance. Uncertainty prompts information seeking across search, social and AI-mediated channels; acquired information is interpreted and emotionally weighted; pre-consultation beliefs form and are carried into the clinical encounter, where they meet clinical reasoning. The encounter resolves towards either epistemic alignment or epistemic tension. Communication calibration is the modifiable process by which tension is converted into shared clinical reasoning. The dashed return path indicates that consultation outcomes shape subsequent information-seeking behavior. Original figure created by the authors.

---

## Methodological Approach

### Design

This study is an interpretive conceptual synthesis. Based on the typology of conceptual article design proposed by Jaakkola (2020), this article fits closest with theory synthesis: drawing evidence and concepts from previously separate bodies of literature to create an account of a phenomenon not well-described by any one of them. Interpretive description informed analytic stance. It was chosen for its orientation towards knowledge that is usable in applied practice (Thorne, 2016).

This design was chosen because the research question posed prior to this one is conceptual, rather than empirical. Until the construct of digital distrust has been delineated from those with which it is currently confounded, a survey instrument cannot be created nor can an intervention be trialed. Trying to measure it before it is properly conceptualized will simply produce precise estimates of the wrong thing.

This study does not report interviews, focus groups, surveys, observations or consultation recordings. We assembled no dataset and make no claims to thematic saturation or enrollment of participants. What was collected was literature, all available articles and books containing data on the topic of interest. No patients or clinicians were involved in this study.

### Material and Literature Identification

Three strands of source material were consulted: scholarly peer-reviewed literature, prioritizing systematic reviews/meta-analyses over primary research articles and conceptual scholarship; authoritative institutional guidance issued on topics related to governance and safety; and practitioner experiential knowledge from the first author's professional practice. A systematic search was not performed, nor does this article purport to report one. Searching was purposive and theory-driven, performed across PubMed, Scopus and Google Scholar, supplemented by citation searching, using a combination of keywords for digital/AI-mediated health information and keywords for trust/communication and the clinical encounter. Literature published from 2020 through to 2026 was prioritized, though seminal sources were retained where necessary for conceptual clarity. Sources were included or excluded based on whether they related directly to the mechanisms under review, rather than containing only thematic relevance. All references were validated with their primary source to ensure they indeed substantiate the claim for which they were used.

### Analysis

Analysis was conducted in four stages, iteratively rather than sequentially. Candidate constructs were lifted from the literature, and their definitions compared. Constructs were then tested against their boundaries by asking what each adjacent construct explains and does not explain, producing the definitional analysis in Section 2.5. Patterns were grouped into themes reported in Section 4, using Braun and Clarke's (2021) reflexive orientation to theme development, in which themes

are constructed analytically rather than extracted. The themes were then interrogated for their communicative implications, producing the protocol in Section 5.

### Reflexivity and Positionality

The first author practices as a family nurse practitioner and has spent more than three decades providing clinical nursing care in various health systems. She has published one blog series describing clinician–patient communication strategies directly to practitioners and patients-facing audiences (Aghanya, 2021a, 2021b, 2021c), one commentary on communication as a quality care lever (Aghanya, 2023), and begun publishing two frameworks of clinical communication (Aghanya, 2025a, 2025b). The orientations presented below are distilled from clinical encounters she has seen multiple times over that span. This yields both the lens through which this paper was constructed, and its primary risk of bias: nursing experience improves clinicians' ability to recognize clinical patterns, but that experience is not evidence, and suffers from recall bias towards memorable patients.

We used three safeguards. Anecdote was allowed to provide fodder for constructs, but never to validate them; all assertions in Section 4 are based on our review of published evidence, and we flag limited data where present. The details of any identifiable patient encounter – not that many make it into the final draft! – appear nowhere in this manuscript. And we include an assessment of the evidence strength for each proposed orientation in Table 1, even when that assessment is unfavorable.

The second author provided input on theoretical framing and ethics from expertise in future-of-health and longevity economics, as well as health-policy research in Nigeria (Kalu, 2024). They also prompted questions and healthy skepticism during each stage of construct development.

### Trustworthiness

We have tried to address quality via source triangulation (four disciplinary literatures); clearly differentiating evidence from interpretation; seeking out negative cases, which in our case led to changes in our arguments (section 2.2); considering alternative theoretical explanations (section 2.3); and acknowledging the limits of what a conceptual synthesis can show. These align with sincerity, credibility and meaningful coherence in Tracy's proposed criteria (2010). We lay no claim to forms of rigour requiring primary data.

### Findings from the Interpretive Synthesis

#### The pre-informed consultation is ordinary

The informed patient who has already looked can now be considered the rule rather than exception in environments of reliable connectivity (Ayre et al., 2025; Shahsavari & Choudhury, 2023; Tan & Goonawardene, 2017). The clinically consequential variable isn't whether patients have searched, but what searching has generated: a question, or an endpoint.

That said, there is a boundary condition. Access to digital medical knowledge still has structural inequalities (Chereka et al., 2025), and universalizing the pre-informed patient would misrepresent much of global clinical practice.

### Fluency has Decoupled from Clinical Validity

Diagnostic abilities of symptom checkers are poor, showing low first-diagnosis accuracy and high variance across tools given the same inputs (Semigran et al., 2015; Wallace et al., 2022). Language models excel at quizzes compared to clinical tasks designed to assess real-world performance, falling most noticeably on questions where a wrong answer could be safety-critical (Gong et al., 2025). Toolkits to improve performance on these questions are unevenly effective across models themselves (Menz et al., 2024).

At the same time, outputs have become much more human-like. Conversations with chatbots are judged as more empathetic than those with real humans on meta-analysis (Ayers et al., 2023; Howcroft et al., 2025), synthetic doctors are judged as equally trustworthy to human clinicians in mock medical visits (Mayer et al., 2026), and simple linguistic fluency has been shown to consistently overinflate ratings of credibility (Alter & Oppenheimer, 2009).

Interpretation: the easy-to-observe outputs a patient receives — coherence, confidence, even simulated attentiveness — have become divorced from the quality that actually matters: clinical accuracy. The patient isn't being irrational; they're reading the signals available to them, and those signals no longer provide information about accuracy.

### Verification Behavior is Asymmetric

Yun and Bickmore's (2025) observation is the evidentiary centerpiece of this synthesis: You searched the Web for health information 49.5% of the time and double-checked it using another source, but only did so 19.4% of the time for conversational AI. The comparison is based on a small sample—only 63 of 297 participants had interacted with a conversational agent—and it needs to be replicated before its magnitude is taken as established. But the direction agrees with Wardle et al. (2025), who found participants in their study described similar thinking, noting distrust of AI responses but also reporting not being able to see sources as preventing them from following through on their distrust.

Interpretation: An interactive results page offers itself as an argument to be evaluated; a chatbot answer offers itself as a solution. The former solicits judgment; the latter forecloses it. This — rather than some blanket skepticism toward technology — is how we think people over- (and under-) trust algorithms, and it may also be why studies examining algorithm aversion/appropriation have found mixed results (Dietvorst et al. 2015; Logg et al. 2019; Longoni et al. 2019): global attitudes are being assessed where interface-specific phenomena may be at play.

### Competing Authorities are Negotiated, not simply held

Patients edit their disclosures. Sommerhalder et al. (2009)

interviewed patients who were concealing online information that contradicted their doctor, specifically so as not to damage the physician's authority. Levy et al. (2018) noted that withholding clinically significant information from physicians is standard practice, not an anomaly. Past rejection informs future action; research on gendered invalidation of chronic pain illustrates how patients who have been belittled about their illness will seek support elsewhere (Samulowitz et al., 2018). Luo et al. (2022) demonstrated clinician response style to be predictive of relational resilience.

Emotion factors into the decision. Anxiety can guide what a patient can voice and process, rather than simply existing as a sentiment that overlays the interaction, and fear cues may be acted on as medical data that must be navigated prior to exchanging information (Aghanya, 2025b). That patients experience less anxiety when they think they are speaking with AI instead of a doctor (Mayer et al., 2026) may explain why they are more comfortable with a computer than with their physician when it comes to the question, they'd least like to ask.

Interpretation: the issue isn't always that people trust digital information over the physician. Often, it's that they don't trust the physician enough to share the digital information. How do you know if the patient agrees with your treatment plan if they're too afraid to let you know they read it shouldn't be approved? Interpretation: digital distrust often hides in plain sight when it's at work. If a patient doesn't bring it up, it doesn't mean it isn't there.

### Communication is where reliance is calibrated

Providers' openness to discussing internet information depends more on their own abilities to appraise evidence, as well as their perceptions of its usefulness, rather than on the amount of time they have (Paige et al., 2023). Patient-centered communication impacts outcomes through mechanisms we can point to, like better information sharing and increased patient agency (Epstein & Street, 2007; Street et al., 2009), and shared decision-making provides a tested framework for discussing choices when faced with uncertainty (Elwyn et al., 2012). Narrative competence provides the contextual lens: allowing oneself to hear someone's story — which may be in part crafted from internet content — as something to understand before fixing (Charon, 2001).

Interpretation: calibration is not something we tack on to the visit. It is the visit, informed by attention to something that is now newly relevant.

### A provisional taxonomy of consultation orientations

The preceding themes are descriptions of mechanisms. In practice clinicians meet the mechanisms as patterns of presentation that they recognize. Table 1 lists five orientations extracted from the synthesis. Each is offered with a suggested communicative response, and an explicit grading of the evidential basis. Three caveats on their use. Firstly, they are orientations to a clinical situation, not types of patient. A patient may display more than one at the same time, and may

move between them during a consultation. The labels describe a position about information; they are never a judgement about the person. Secondly, none has been empirically validated.

**Table 1:** Digitally mediated consultation orientations and corresponding communication strategies (provisional conceptual taxonomy).

| Orientation | Core feature                                                                                                 | Possible mechanism                                                                                                                                                                                                                  | Clinical communication challenge                                                         | Suggested response                                                                                             | Evidence status |
|-------------|--------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------|-----------------|
| Anchored    | A prior digital conclusion functions as the reference standard against which clinical assessment is measured | Anchoring on a fluent, confident output; uniqueness neglect may operate in the opposite direction and is not resolved (Alter & Oppenheimer, 2009; Longoni et al., 2019)                                                             | Clinical reasoning is heard as a competing opinion rather than a different kind of claim | Clarify what the source actually stated versus what was inferred; contextualize rather than contradict         | Low             |
| Aggregative | Extensive digital material is presented; volume substitutes for synthesis                                    | Information seeking as uncertainty management (Brashers, 2001); accumulation as a proxy for control                                                                                                                                 | The presenting problem is obscured by the material assembled around it                   | Acknowledge the effort, then ask which single finding worried them most                                        | Low             |
| Depleted    | Prolonged, contradictory searching has produced fatigue and disengagement                                    | Information avoidance following overload (Brashers, 2001); skepticism about all sources, including clinical ones (Wardle et al., 2025)                                                                                              | Apparent indifference may be read as non-engagement rather than exhaustion               | Reduce interpretive load; offer a small number of clear next steps rather than fuller explanation              | Low to moderate |
| Testimonial | Narrative, identification-based sources carry authority; lived accounts outweigh aggregate evidence          | Narrative persuasion; perceived warmth of non-clinical sources (Howcroft et al., 2025); prevalence of health misinformation on social platforms (Suarez-Lledo & Alvarez-Galvez, 2021)                                               | Challenging the source can be experienced as dismissing experience itself                | Engage the narrative before the evidence; separate the experience described from the claim generalized from it | Moderate        |
| Guarded     | Material is deliberately withheld, or disclosure is partial                                                  | Privacy concern; anticipated judgement; prior dismissal (Samulowitz et al., 2018; Sommerhalder et al., 2009); documented non-disclosure (Levy et al., 2018); lower reported stress with a machine interlocutor (Mayer et al., 2026) | The discrepancy is not in the room and cannot be addressed                               | Ask directly and without evaluation; make non-judgement explicit rather than implied                           | Moderate        |

**Note:** Evidence status refers to the strength of published support for the orientation as a distinct consultation pattern, not to the strength of the underlying mechanism. *Low* = practice-derived and inferred indirectly from adjacent literature (Anchored, from work on processing fluency and the weighting of algorithmic advice; Aggregative, from uncertainty management theory). *Low to moderate* = consistent with uncertainty management theory and with qualitative accounts of information seeking in the AI era (Depleted). *Moderate* = a directly relevant empirical literature exists, although not at the level of the consultation itself (Testimonial, from research on the prevalence of health misinformation; Guarded, from documented rates of patient non-disclosure and of withholding to protect clinician standing). No orientation has been empirically validated.

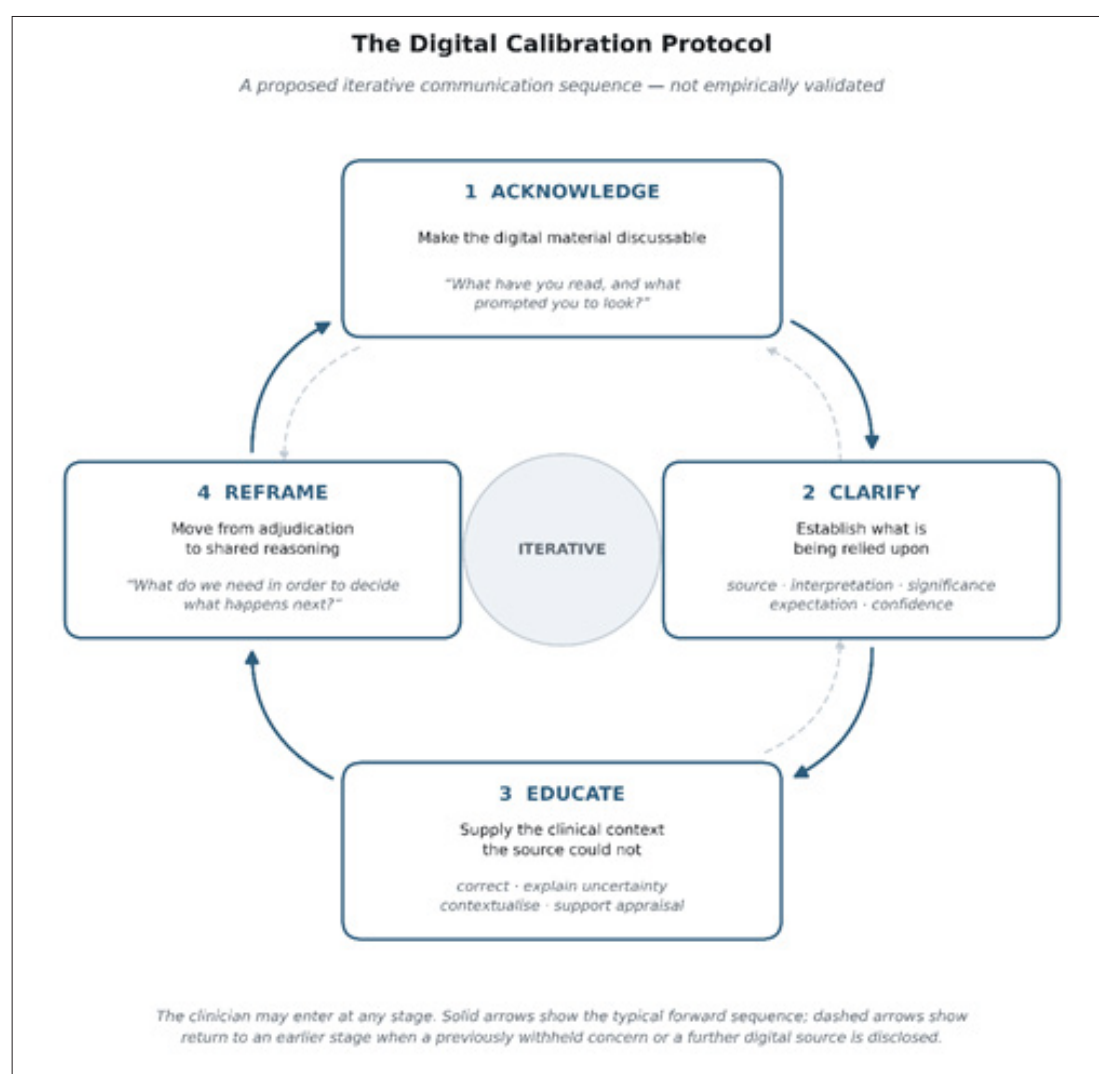
## Development of the Digital Calibration Protocol

Taken together, these five themes point to a progression. If online information is normative (Theme 4.1), it needs to be invited, not merely tolerated. If fluency is deceptive (Theme 4.2) and conversational resources remain unverified (Theme 4.3), both the origin and its warrant need to be identified. If mismatch is frequently obscured (Theme 4.4), clinicians need to uncover it, rather than wait for it. And if calibration is conversational labor (Theme 4.5), the interaction needs to conclude with shared understanding instead of judgement.

We suggest this sequence: Digital Calibration Protocol (DCP): Acknowledge — Clarify — Educate — Reframe. The steps are recursive, not linear (Figure 2). One might begin in medias res, return to clarify as new information arises, and complete multiple cycles within a single visit. This framework is a hypothesis generated from synthesis. It has not been validated,

and we offer no evidence that it enhances trust, precision, concordance, or any other metric.

DCP exists in a particular relationship to two prior communication models, which warrants clarification instead of obfuscation. The Therapeutic Dialogue Framework articulates the overall structure of an encounter as relational rather than transactional (Aghanya, 2025a); Fear-Informed Care establishes the necessary emotional context for disclosure and understanding (Aghanya, 2025b). Neither, however, was designed with this clinical problem in mind, because neither incorporates the scenario where patients arrive pre-furnished with data about their body made elsewhere. As such, DCP should be considered a specialized offshoot built upon those prior models: it assumes the relationships those papers establish, then layers on the epistemic work necessitated by digitally sourced inputs.



**Figure 2:** The Digital Calibration Protocol, shown as an iterative rather than linear sequence. Entry is possible at any stage; return to earlier stages occurs whenever new digital material or a previously withheld concern emerges. Original figure created by the authors.

**Stage 1 – Acknowledge. Goal:** bring online material into the conversational realm. Why? Patients only withhold online material if they expect it to be rejected (Levy et al., 2018; Sommerhalder et al., 2009), and clinician openness is one variable (among many) that influences whether they do (Paige et al., 2023); calling out fear prior to information exchange is its analogue at the start of any visit (Aghanya, 2025b). Action: provide a clear, nonjudgmental invitation – ask about what patients have read and what led them to look, before making any judgment about its value. Limits: acknowledgement is not acceptance, and patients must be able to read that. Empathy that patients interpret as agreement only delays the inevitable dismantling ahead.

**Stage 2 – Clarify. Goal:** understand what information is actually being relied upon. Why? Patients' reliance varies by search "format" (such as diagnosis websites versus medication labels), and their fact-checking behavior may vary by format too (Wardle et al., 2025; Yun & Bickmore, 2025). Five components merit individual attention: the source (which website, app, or tool); the interpretation (what they believed it meant, often at odds with what it said); the emotional import (what would happen if it were true), which often matters more than its factual content and may itself be the driver of continued searching (Aghanya, 2021b, 2025b); patients' expectation (what they think will happen now as a result of this visit); and the level of confidence they placed in the information, including whether anyone fact-checked it. Limits: curiosity, not cross-examination. Asked bluntly, these questions will feel like you're testing them.

**Stage 3 – Educate. Goal:** provide the clinical context the internet did not. Education, despite what its name suggests, carries over from when this framework was first drafted, describes how something fits into clinical thinking, not teaching patients per se. Three distinct processes happen under this umbrella term and should be unpicked: correction, where online information is factually inaccurate; explanation, where you cannot yet say for certain; and contextualization, where a statement is both true and untrue of the patient in front of you. Why? Both symptom checkers and ChatGPT fail patients most frequently on the very skill that differentiates clinicians from algorithms (Gong et al., 2025; Wallace et al., 2022). Limits: brevity is essential. Patients bombarded with clinical information have had their callipers reset to overconfidence-about-the-internet. They'll understand nothing you say thereafter (Aghanya, 2025b; Mayer et al., 2026).

**Stage 4 – Reframe. Goal:** transition patients from triage to team membership. Stop asking "whose information is right?" and start asking "what do we both need to find out to decide the best way forward?" This includes history, examination, investigations, watchful waiting, and, crucially, an explicit statement of what we don't yet know. Why? Shared decision-making offers the mechanics for how to do this well (Elwyn et al., 2012), and uncertainty management tells us why a plan made under acknowledged uncertainty is acceptable to patients, but ambiguity is not (Brashers, 2001). Limits: reframing shouldn't

become a conduit for surreptitious professional privilege. If what your patient found was correct, congratulations – that is the reframing.

## Discussion

It's important to note that what follows isn't so much a description of something previously unknown as it is a point of differentiation. Digital distrust is not health misinformation (content property) nor is it medical mistrust (longstanding or well-founded lack of trust in institutions and systems). Digital distrust is different from cyberchondria because cyberchondria implies anxiety amplification, and from automation bias because that literature is clinician-facing rather than patient-facing.

What the DCP brings to patient-clinician conversations that's novel is narrow but, we believe, non-trivial. Shared decision-making frameworks concern choice between options, and motivational interviewing concerns getting someone from ambivalence about taking action to making a decision. Neither concerns the conversation where someone comes to you with a different mental model of what's happening to them than the one you arrive at based on your tools, your training, and your conversation with them. Stages 2 and 3 are for that conversation.

One potential pitfall deserves another. Validating misinformation presented by a patient is harmful, but noxious, because it lends clinical credibility to information that doesn't deserve it and can lead to patients receiving care too late. Dismissing what someone has learned online, though, isn't the harmless inverse. It triggers the reason they hid their web-searching in the first place (Levy et al., 2018; Sommerhalder et al., 2009), and erases the data point the clinician needs most. Literature around dismissal exists in chronic pain; we recommend it to those who want to learn where that road goes (Samulowitz et al., 2018). The DCP is meant to help clinicians avoid both behaviors.

Finally, how far beyond General Pediatrics you take this framework will depend on context. In primary care and urgent care settings where time with patients is finite, most of the empathic calories are in Stage 2; in Emergency Medicine, Stage 1 may be skipped in the service of triaging. Women's health clinicians will be familiar with taking histories where patients' actual symptoms have been previously dismissed (Samulowitz et al., 2018), and may need to work through why being believed by clinicians is new before any digital history is revealed. Pediatric clinicians will have children's health literacy to consider, but often the one who's doing the Googling is parent, not the patient. Chronic disease patients may come to appointments with extensive topic-relevant knowledge that will again make straight application of Stage 2 inappropriate; clinicians will need to know what a calibration conversation sounds like when a patient is already right. Mental health clinicians will know that the diagnosis chapter their patient learns about from Me.You.AI will precede what they tell you about their mental state.

Application of the DCP to non-white-majority settings should proceed with more caution than the literature will comfortably allow. Health online requires literacy skills as well as internet access, and there are measurable differences in both across cultural settings. What we can say is true: there are substantial barriers to internet-mediated health information access across African contexts (Chereka et al., 2025), and non-white-majority English speakers and people with low health literacy use conversational AI for health more frequently than their peers do (Ayre et al., 2025). Past that, however, the research literature is thin. Why it's thin is at least partially structural: academic research output on health questions across African settings is impoverished by the localized development of health research and policy infrastructure (Kalu, 2024), and the epidemiologic datasets the rest of this literature depends on — vital registration systems that record births, deaths, and sometimes other health events at the population level — are themselves incomplete. Death registration completeness in Nigeria was estimated at 9.9% in 2017 (Makinde et al., 2020). Contexts where African-descended people live globally can similarly be pluralistic medical information environments where Buddhist, Traditional, and biomedical sources co-exist with varying degrees of trust Attention.

If you replaced every instance of “AI” in the previous paragraph with “Twitter,” it would remain relevant. We flag this not to fill in the blanks with speculation but to highlight it as a gap. That said, gaps exist in all populations.

Medical ethicists will have notes. Language models produce convincing but inaccurate health information; doctors have not meaningfully implemented safeguards against AI in health known to mitigate health disinformation spread (Menz et al., 2024), and language models have their largest performance deficits in the languages where mistakes are most consequential (Gong et al., 2025). With regard to the conversational AI patients will use, consumer-grade privacy profiles and data governance still need to be hashed out between doctors and tech companies. Algorithmic discrimination in health is not theoretical (Obermeyer et al., 2019). The WHO has published both broad (2021) and AI-specific (2024) ethics and governance guidance on AI for health. The supplementary material on source credibility for identifying helpful vs harmful health content online is a start on the medical intelligence side (Kington et al., 2021), as is the literature on health misinformation spread on social media platforms (Suarez- Lledo & Alvarez- Galvez, 2021). But you knew that. None of this is fear-mongering, either. These models democratize access to vocabulary and medical knowledge that had been the purview of prior experts, give voice to patients who were struggling to describe their symptoms, and make some questions askable that were previously unthinkable. But they do not absolve clinicians of responsibility for clinical decisions.

### Clinical and Policy Implications

Three implications can be drawn for practice and policy. Digital-material elicitation should be a component of communication training taught as standard history-taking competence rather than as advanced interviewing skill, both

because clinicians' own capacity to appraise predicts whether they have conversations like the ones proposed (Paige et al., 2023) and because communication skill is likely a driver of care quality rather than contributory to it already (Aghanya, 2023). Public-facing efforts to bolster appraisal must account for the verification asymmetry I describe here (Yun & Bickmore, 2025), because it makes poor sense to tell people to check sources when they're interacting with an interface that presents none to check; patient-facing communication guidance is perhaps a more actionable complement to source-appraisal education alone (Aghanya, 2021a). And developers could reasonably be expected, at minimum, to make uncertainty and provenance evident at the point of output, for which purpose existing governance frameworks provide the lexicon (World Health Organization, 2021, 2024).

The only implication I draw for researchers: Nothing I propose in this paper should be added to practice guidelines until it has been evaluated.

### Limitations and Future Research

This article isn't an empirical study; we recruited no participants and tested no intervention. The orientations listed in Table 1 are hypothetical groupings of differing evidential weight; they're deliberately ordinal and unvalidated, and may not be comprehensive, mutually exclusive, or consistent across contexts. We haven't tested the DCP in a clinical setting, so we make no claims regarding its impact. Several caveats about the evidence base itself delimit our findings. The asymmetry underpinning Theme 4.3 is based on one cross-sectional study with a small conversational- AI subgroup. Measures used to compare empathy across Theme 4.2 are mostly text-based and unvalidated. Experiments examining stress and recall used mock rather than actual consultations. Our synthesis was targeted rather than systematic and so carries selection biases; it's also skewed toward high-income, English-language contexts. The first author's clinical impressions informed analysis at multiple points and, while mitigated as outlined in Section 3.4, cannot be discounted. And of course the technology will have advanced by the time this article is published – the specific performance metrics we discuss will become dated, if they aren't already; the structural case we make regarding fluency, format, and verification should hold longer, but still needs testing.

Which leads us to future research. The orientations need validation with qualitative, consultation-based research. The asymmetry identified in Theme 4.3 merits experimental testing to see if format is really the causal factor. The DCP needs feasibility studies before we can even begin to consider its effectiveness. And the cultural gap we identified in Section 6 needs to be investigated with primary research in linguistically pluralistic healthcare settings, not inferred from extant samples.

### Conclusion

Digitally-mediated information has changed the epistemic context in which consultations are initiated. Patients come into consultations having already heard something. They will often come in having been told fluently. They will often have been

told by a system that doesn't know what it doesn't know.

The clinical response shouldn't be either wholesale acceptance or reflex rejection. Both approaches mischaracterize the problem, which isn't digital information itself but an appropriately calibrated relationship to it. Separating digital engagement (which is overwhelmingly positive and should be encouraged) from digital dependency (which is contextual and fixable) is far more productive.

Patient-centered communication is one way to achieve that calibration. The Digital Calibration Protocol represents an attempt at a provisional, falsifiable method for doing just that. Does it work? That, too, is an empirical question. One we haven't yet answered.

### Declarations

Ethics approval and consent to participate: Not Applicable. The study has no human participants, contains no personal data and involves no new data collection of any kind. It is a synthesis of published evidence. Formal ethics review was not required nor obtained.

**Consent for Publication:** Not Applicable. No individual patient data/images or other identifiable clinical information is included in this work. No described case corresponds to a real patient.

**Availability of Data and Materials:** Not Applicable. No datasets were generated or analyzed for this study. All sources underpinning the review are cited in the manuscript and publicly available.

**Competing Interests:** The authors declare that they have no competing financial interests. They do have two non-financial interests which they feel warrant disclosure. NT Authority has previously published books and articles exploring clinical communication, some of which are cited here as references (Aghanya, 2021a, 2021b, 2021c, 2023, 2025a, 2025b). These works support individual claims made within the article, but are not otherwise referenced or alluded to indiscriminately. OAK is an Executive member of TAFD's, a publisher of three of the above works. TAFD has neither input into nor funded this study, and neither author is remunerated based on article citations or inclusion of these references.

**Funding:** No funding was secured for this research from any agency, commercial or otherwise.

**Authors' Contributions:** NTA conceived of the study and its clinical framework, providing practice-based information used to develop the orientations. OAK contributed the theoretical lens, ethical framing and analysis regarding epistemic authority. Both authors contributed to construct development, literature appraisal and manuscript revision. All authors approved the final manuscript and take responsibility for the integrity of the work. There are no other people who meet ICMJE criteria for authorship that we have not included.

**Acknowledgements:** None.

**Use of AI-assisted Technologies:** AI language tool used for literature identification and language editing. The authors evaluated all AI recommendations and verified the information against original sources. All facts, statistics, quotes, and figures were re-derived and cross-checked with original sources. Any and all assertions, conceptual claims, and clinical interpretations are those made by the authors.

### Disclaimer

The Digital Calibration Protocol is an emerging hypothesis for discussion. It has not been tested. It is not an assessment. It is not a clinical algorithm. It is not a triage tool. It is not a replacement for clinical judgment. The Consultation types presented are hypothetical constructs that are not validated groupings of patients, diagnoses, or psychological types. They should not be documented as patient traits. Healthcare providers should continue to adhere to all relevant clinical guidelines, standards of practice, scopes-of-practice, and local regulations as they see fit. Nothing in this document should be taken as a clinical recommendation for any specific patient.

### References

1. Ayre, J., Cvejic, E., & McCaffery, K. J. (2025). Use of ChatGPT to obtain health information in Australia, 2024: Insights from a nationally representative survey. *Medical Journal of Australia*, 222(4), 210–212. DOI: <https://doi.org/10.5694/mja2.52598>
2. Shahsavari, Y., & Choudhury, A. (2023). User intentions to use ChatGPT for self-diagnosis and health-related purposes: Cross-sectional survey study. *JMIR Human Factors*, 10,. DOI: <https://doi.org/10.2196/47564>
3. Luo, A., Qin, L., Yuan, Y., Yang, Z., Liu, F., Huang, P., & Xie, W. (2022). The effect of online health information seeking on physician-patient relationships: Systematic review. *Journal of Medical Internet Research*, 24(2). DOI: <https://doi.org/10.2196/23354>
4. Tan, S. S.-L., & Goonawardene, N. (2017). Internet health information seeking and the patient-physician relationship: A systematic review. *Journal of Medical Internet Research*, 19(1). DOI: <https://doi.org/10.2196/jmir.5729>
5. Sommerhalder, K., Abraham, A., Caiata Zufferey, M., Barth, J., & Abel, T. (2009). Internet information and medical consultations: Experiences from patients' and physicians' perspectives. *Patient Education and Counseling*, 77(2), 266–271. DOI: <https://doi.org/10.1016/j.pec.2009.03.028>
6. Wallace, W., Chan, C., Chidambaram, S., Hanna, L., Iqbal, F. M., Acharya, A., Normahani, P., Ashrafian, H., Markar, S. R., Sounderajah, V., & Darzi, A. (2022). The diagnostic and triage accuracy of digital and online symptom checker tools: A systematic review. *npj Digital Medicine*, 5. DOI: <https://doi.org/10.1038/s41746-022-00667-w>
7. Gong, E. J., Bang, C. S., Lee, J. J., & Baik, G. H. (2025). Knowledge-practice performance gap in clinical large language models: Systematic review of 39 benchmarks. *Journal of Medical Internet Research*, 27. DOI: <https://doi.org/10.2196/84120>

8. Yun, H. S., & Bickmore, T. (2025). Online health information-seeking in the era of large language models: Cross-sectional web-based survey study. *Journal of Medical Internet Research*, 27. DOI: <https://doi.org/10.2196/68560>
9. Kaneda, Y. (2026). When patients consult artificial intelligence before clinicians: Restoring clinical prioritization in mental health care. *The British Journal of Psychiatry*.1-2. DOI: <https://doi.org/10.1192/bjp.2026.10741>
10. Semigran, H. L., Linder, J. A., Gidengil, C., & Mehrotra, A. (2015). Evaluation of symptom checkers for self-diagnosis and triage: Audit study. *BMJ*, 351. DOI: <https://doi.org/10.1136/bmj.h3480>
11. Howcroft, A., Bennett-Weston, A., Khan, A., Griffiths, J., Gay, S., & Howick, J. (2025). AI chatbots versus human healthcare professionals: A systematic review and meta-analysis of empathy in patient care. *British Medical Bulletin*, 156(1). DOI: <https://doi.org/10.1093/bmb/ldaf017>
12. Ayers, J. W., Poliak, A., Dredze, M., Leas, E. C., Zhu, Z., Kelley, J. B., Faix, D. J., Goodman, A. M., Longhurst, C. A., Hogarth, M., & Smith, D. M. (2023). Comparing physician and artificial intelligence chatbot responses to patient questions posted to a public social media forum. *JAMA Internal Medicine*, 183(6), 589–596. DOI: <https://doi.org/10.1001/jamainternmed.2023.1838>
13. Alter, A. L., & Oppenheimer, D. M. (2009). Uniting the tribes of fluency to form a metacognitive nation. *Personality and Social Psychology Review*, 13(3), 219–235. DOI: <https://doi.org/10.1177/1088868309341564>
14. Mayer, C. J., Swysen, T. L., Kurz, T. R., Stephan, T. I., Geisel, D., Doll, E. S., Mahal, N., Lerch, S. P., Sailer, S., Schaaf, C. P., Merz, C. J., Wüstenberg, T., Ehrenthal, J. C., Walter, S., Mahal, J., & Ditzen, B. (2026). Perceptions of simulated artificial intelligence in medical consultations: Associations with stress, memory, and perceived credibility. *npj Digital Medicine*, 9. DOI: <https://doi.org/10.1038/s41746-026-03022-5>
15. Menz, B. D., Kuderer, N. M., Bacchi, S., Modi, N. D., Chin-Yee, B., Hu, T., Rickard, C., Haseloff, M., Vitry, A., McKinnon, R. A., Kichenadasse, G., Rowland, A., Sorich, M. J., & Hopkins, A. M. (2024). Current safeguards, risk mitigation, and transparency measures of large language models against the generation of health disinformation: Repeated cross-sectional analysis. *BMJ*, 384. DOI: <https://doi.org/10.1136/bmj-2023-078538>
16. Wardle, C., Urbani, S., & Wang, E. (2025). Evolving health information-seeking behavior in the context of Google AI Overviews, ChatGPT, and Alexa: Interview study using the think-aloud protocol. *Journal of Medical Internet Research*, 27. DOI: <https://doi.org/10.2196/79961>
17. Williamson, L. D., & Bigman, C. A. (2018). A systematic review of medical mistrust measures. *Patient Education and Counseling*, 101(10), 1786–1794. DOI: <https://doi.org/10.1016/j.pec.2018.05.007>
18. Fonagy, P., & Allison, E. (2014). The role of mentalizing and epistemic trust in the therapeutic relationship. *Psychotherapy*, 51(3), 372–380. DOI: <https://doi.org/10.1037/a0036505>
19. Sperber, D., Clément, F., Heintz, C., Mascaro, O., Mercier, H., Origgi, G., & Wilson, D. (2010). Epistemic vigilance. *Mind & Language*, 25(4), 359–393. DOI: <https://doi.org/10.1111/j.1468-0017.2010.01394.x>
20. Hatherley, J. J. (2020). Limits of trust in medical AI. *Journal of Medical Ethics*, 46(7), 478–481. <https://doi.org/10.1136/medethics-2019-105935>
21. Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. DOI: <https://doi.org/10.1037/xge0000033>
22. Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103. DOI: <https://doi.org/10.1016/j.obhdp.2018.12.005>
23. Longoni, C., Bonezzi, A., & Morewedge, C. K. (2019). Resistance to medical artificial intelligence. *Journal of Consumer Research*, 46(4), 629–650. DOI: <https://doi.org/10.1093/jcr/ucz013>
24. Goddard, K., Roudsari, A., & Wyatt, J. C. (2012). Automation bias: A systematic review of frequency, effect mediators, and mitigators. *Journal of the American Medical Informatics Association*, 19(1), 121–127. DOI: <https://doi.org/10.1136/amiajnl-2011-000089>
25. White, R. W., & Horvitz, E. (2009). Cyberchondria: Studies of the escalation of medical concerns in web search. *ACM Transactions on Information Systems*, 27(4). DOI: <https://doi.org/10.1145/1629096.1629101>
26. McMullan, R. D., Berle, D., Arnáez, S., & Starcevic, V. (2019). The relationships between health anxiety, online health information seeking, and cyberchondria: Systematic review and meta-analysis. *Journal of Affective Disorders*, 245, 270–278. DOI: <https://doi.org/10.1016/j.jad.2018.11.037>
27. Norman, C. D., & Skinner, H. A. (2006). eHEALS: The eHealth Literacy Scale. *Journal of Medical Internet Research*, 8(4). DOI: <https://doi.org/10.2196/jmir.8.4.e27>
28. Brashers, D. E. (2001). Communication and uncertainty management. *Journal of Communication*, 51(3), 477–497. DOI: <https://doi.org/10.1111/j.1460-2466.2001.tb02892.x>
29. Chereka, A. A., Shibabaw, A. A., Butta, F. W., Tadesse, M. N., Abebe, M. T., Atanie, F. A., & Kitil, G. W. (2025). Explore barriers to using the internet for health information access in African countries: A systematic review. *PLOS Digital Health*, 4(1). DOI: <https://doi.org/10.1371/journal.pdig.0000719>
30. Jaakkola, E. (2020). Designing conceptual articles: Four approaches. *AMS Review*, 10(1), 18–26. DOI: <https://doi.org/10.1007/s13162-020-00161-0>
31. Thorne, S. (2016). Interpretive description: Qualitative research for applied practice (2<sup>nd</sup> ed.). *Routledge*. DOI: <https://doi.org/10.4324/9781315545196>
32. Braun, V., & Clarke, V. (2021). One size fits all? What counts as quality practice in (reflexive) thematic analysis? *Qualitative Research in Psychology*, 18(3), 328–352. DOI: <https://doi.org/10.1080/14780887.2020.1769238>

33. Aghanya, N. T. (2021a). Effective communication: Guide book for patients. TAFD's Publishers.
34. Aghanya, N. T. (2021b). Principles for overcoming communication anxiety and improving trust. TAFD's Publishers.
35. Aghanya, N. T. (2021c). Think, communicate, heal: Discover the healing power of clear communication. TAFD's Publishers.
36. Aghanya, N. T. (2023). A brief commentary on improving the quality of healthcare delivery practices using the soft skills of communication. *Integrative Journal of Nursing and Medicine*, 4(5), 1–2. DOI: <https://doi.org/10.31038/IJNM.2023442>
37. Aghanya, N. T. (2025a). Communicate to heal: Reclaiming the clinician's voice in the pursuit of health and longevity. *Journal of Medical Clinical Case Reports*, 7(4), 1-10. DOI: <https://doi.org/10.47485/2767-5416.1129>
38. Aghanya, N. T. (2025b). Fear-informed care: A new clinical paradigm for trust restoration in patient communication. *Japan Journal of Research*, 6(6).
39. Kalu, O. A. (2024). Policy formulation as a catalyst for anti-aging research: Opportunities and challenges in Africa. *Advances in Clinical and Medical Research*, 5(4), 1–10. DOI: [https://doi.org/10.52793/ACMR.2024.5\(4\)-90](https://doi.org/10.52793/ACMR.2024.5(4)-90)
40. Tracy, S. J. (2010). Qualitative quality: Eight "big-tent" criteria for excellent qualitative research. *Qualitative Inquiry*, 16(10), 837–851. DOI: <https://doi.org/10.1177/1077800410383121>
41. Levy, A. G., Scherer, A. M., Zikmund-Fisher, B. J., Larkin, K., Barnes, G. D., & Fagerlin, A. (2018). Prevalence of and factors associated with patient nondisclosure of medically relevant information to clinicians. *JAMA Network Open*, 1(7). DOI: <https://doi.org/10.1001/jamanetworkopen.2018.5293>
42. Samulowitz, A., Gremy, I., Eriksson, E., & Hensing, G. (2018). "Brave men" and "emotional women": A theory-guided literature review on gender bias in health care and gendered norms towards patients with chronic pain. *Pain Research and Management*, 2018. DOI: <https://doi.org/10.1155/2018/6358624>
43. Paige, S. R., Bylund, C. L., Wilczewski, H., Ong, T., Barrera, J. F., Welch, B. M., & Bunnell, B. E. (2023). Communicating about online health information with patients: Exploring determinants among telemental health providers. *PEC Innovation*, 2, DOI: <https://doi.org/10.1016/j.pecinn.2023.100176>
44. Epstein, R. M., & Street, R. L., Jr. (2007). Patient-centered communication in cancer care: Promoting healing and reducing suffering (NIH Publication No. 07-6225). National Cancer Institute.
45. Street, R. L., Jr., Makoul, G., Arora, N. K., & Epstein, R. M. (2009). How does communication heal? Pathways linking clinician-patient communication to health outcomes. *Patient Education and Counseling*, 74(3), 295–301. DOI: <https://doi.org/10.1016/j.pec.2008.11.015>
46. Elwyn, G., Frosch, D., Thomson, R., Joseph-Williams, N., Lloyd, A., Kinnersley, P., Cording, E., Tomson, D., Dodd, C., Rollnick, S., Edwards, A., & Barry, M. (2012). Shared decision making: A model for clinical practice. *Journal of General Internal Medicine*, 27(10), 1361–1367. DOI: <https://doi.org/10.1007/s11606-012-2077-6>
47. Charon, R. (2001). Narrative medicine: A model for empathy, reflection, profession, and trust. *JAMA*, 286(15), 1897–1902. DOI: <https://doi.org/10.1001/jama.286.15.1897>
48. Suarez-Lledo, V., & Alvarez-Galvez, J. (2021). Prevalence of health misinformation on social media: Systematic review. *Journal of Medical Internet Research*, 23(1). DOI: <https://doi.org/10.2196/17187>
49. Makinde, O. A., Odimegwu, C. O., Udoh, M. O., Adedini, S. A., Akinyemi, J. O., Atobatele, A., Fadeyibi, O., Abdulaziz-Sule, F., Babalola, S., & Orobato, N. (2020). Death registration in Nigeria: A systematic literature review of its performance and challenges. *Global Health Action*, 13(1). DOI: <https://doi.org/10.1080/16549716.2020.1811476>
50. Obermeyer, Z., Powers, B., Vogeli, C., & Mullainathan, S. (2019). Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464), 447–453. DOI: <https://doi.org/10.1126/science.aax2342>
51. Kington, R. S., Arnesen, S., Chou, W.-Y. S., Curry, S. J., Lazer, D., & Villarruel, A. M. (2021). Identifying credible sources of health information in social media: Principles and attributes. *NAM Perspectives*. DOI: <https://doi.org/10.31478/202107a>
52. World Health Organization. (2021). Ethics and governance of artificial intelligence for health: WHO guidance.
53. World Health Organization. (2024). Ethics and governance of artificial intelligence for health: Guidance on large multi-modal models.

**Copyright:** ©2026. Nonye Tochi Aghanya. This is an open-access article distributed under the terms of the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.