

Navigating Artificial Intelligence in Neuroregulation Practice: Ethical Principles, Risk Stratification, and a Clinical Policy Framework

Leslie H. Sherlin^{1,2,3*} and Robert E. Longo⁴

¹Sherlin Consulting Group, Scottsdale, Arizona, USA

²Ottawa University, Surprise, Arizona, USA

³Grand Canyon University, Phoenix, Arizona, USA

⁴Private Practice, Retired, Wilmington, North Carolina, USA

Abstract

Artificial intelligence (AI) tools are entering neuroregulation practice through multiple simultaneous pathways, including automated qEEG analysis, protocol recommendation systems, AI-assisted documentation, and consumer-facing mental health applications that clients bring directly into the therapeutic relationship. No ethics code specific to neuroregulation has yet addressed these applications directly. This article presents a conceptual and practical ethical framework grounded in established codes, including the BCIA Code of Ethics, ISNR Code of Ethics, APA Ethical Principles, ACA Code of Ethics, and the APA (2025) Ethical Guidance for Artificial Intelligence. A risk continuum model organizing AI applications along three dimensions (opacity, clinical consequence, and distance from oversight) is proposed. Four ethical domains are addressed: informed consent and transparency, practitioner competence and scope, data privacy and vendor accountability, and three AI accuracy risks (hallucination, automation bias, and collusion). A five-element practice policy framework is presented in tabular form for direct clinical use. Existing ethics principles are sufficient to guide responsible AI integration when applied deliberately. The practitioner remains accountable for AI-assisted decisions regardless of tool design or vendor claims. Proactive policy development positions neuroregulation clinicians to shape the field's emerging standards.

Author's Note on Temporal Limitations. The landscape of AI in health care is evolving at a pace that outstrips the typical publication cycle of academic literature. The specific tools, regulatory developments, and empirical findings described in this article reflect the state of the field as of the time of writing (mid-2026). The ethical frameworks, guiding principles, and policy elements presented here are designed to be durable: they are grounded in established ethics codes that are not tool-specific and will remain applicable as the technology continues to develop. Readers should expect that particular examples, cited tools, and regulatory references may have been superseded depending on when this article is read. The authors encourage readers to treat the frameworks as stable anchors and the specific illustrations as illustrative rather than exhaustive.

Keywords: artificial intelligence; neuroregulation; neurofeedback; ethics; automation bias; clinical policy; data privacy

Citation: Sherlin, L. H., & Longo, R. E. (2026). Navigating artificial intelligence in neuroregulation practice: Ethical principles, risk stratification, and a clinical policy framework. *NeuroRegulation*, 13(3), 280–292. <https://doi.org/10.15540/nr.13.3.280>

***Address correspondence to:** Leslie H. Sherlin, PhD, Sherlin Consulting Group, 7272 E. Indian School Rd., Suite 540, Scottsdale, AZ 85251, USA. Email: leslie@drsherlin.com

Copyright: © 2026. Sherlin and Longo. This is an Open Access article distributed under the terms of the Creative Commons Attribution License (CC-BY).

Edited by:

Rex L. Cannon, PhD, Currents, Knoxville, Tennessee, USA

Reviewed by:

Rex L. Cannon, PhD, Currents, Knoxville, Tennessee, USA

Randall Lyle, PhD, Mount Mercy University, Cedar Rapids, Iowa, USA

Introduction

The Current Moment

The neuroregulation field is encountering artificial intelligence (AI) not as a distant development to anticipate but as a present reality arriving from multiple directions simultaneously. Commercial platforms are offering automated qEEG interpretation, returning clinically framed reports generated with limited or no visible reasoning. Protocol recommendation engines are suggesting training parameters on the basis of assessment data and pattern-matching algorithms trained on large datasets. AI-assisted documentation tools adopted across health care broadly are finding their way into the administrative workflows of neuroregulation practices, drafting session notes, generating treatment summaries, and composing communications submitted under the practitioner's name and credentials. And, perhaps most consequentially, clients are arriving at sessions having already consulted AI tools about their symptoms, their diagnosis, their protocols, and their progress, often without having disclosed this to their clinician.

This convergence is not a future state. It is the present one. Horvitz and West (2026) observed in a *Science* editorial that AI systems are being woven into everyday life faster than humans can understand them, while AI's capacity to process and model human behavior is simultaneously growing. That asymmetry applies directly to clinical practice: Practitioners are being asked to make responsible decisions about tools whose inner workings they often cannot inspect, while their clients are engaging with those same tools outside of any professional relationship.

The ethical challenge this creates is not that AI is inherently dangerous or that its clinical applications are without value. The challenge is that AI creates new applications for ethical obligations that have always existed, and the field has not yet systematically articulated how those obligations apply. No ethics code specific to neuroregulation currently addresses AI use directly. The American Psychological Association (APA) published guidance for psychologists in 2025 (APA, 2025). Medicine is developing its own standards. Regulation at the state and federal levels is developing more slowly, but it is developing. The practitioners who engage with these questions now will help write the standards the field adopts. Those who wait will receive standards written by others.

This article presents a conceptual and practical framework for ethical AI use in neuroregulation practice, grounded in the ethics codes that already govern the field. It does not evaluate specific AI products, and it does not prescribe which tools practitioners should adopt or avoid. It provides a framework for reasoning about AI use that will remain applicable as the tool landscape continues to evolve.

Relationship to Prior Work

This article extends a line of applied ethics analysis in neuroregulation concerned with the responsibilities practitioners carry in domains where technical complexity and professional obligation intersect. Sherlin and Longo (2025) examined ethical billing practices in neurofeedback and qEEG, mapping the obligations created by CPT code requirements, scope-of-practice considerations, and the accountability frameworks of payers and licensing boards onto the practical realities of clinical practice. The present article undertakes the same kind of analysis for AI: identifying the specific points at which established ethical principles create operative obligations and offering practical guidance for fulfilling them.

Conceptual Foundations: Relevant Ethics Codes and Guiding Principles

The Applicable Codes

Neuroregulation practitioners operate under multiple overlapping ethics codes depending on their licensure, certification, and professional affiliations. The Biofeedback Certification International Alliance (BCIA) Code of Ethics governs all BCIA-certified practitioners and addresses competence, scope of practice, client welfare, and professional responsibility. The International Society for Neurofeedback and Research (ISNR) Code of Ethics and Standards of Practice establishes conduct expectations for ISNR members, with particular attention to research integrity, clinical accuracy, and public representation of the field. For licensed mental health practitioners who provide neuroregulation services, the APA Ethical Principles of Psychologists and Code of Conduct (APA, 2017) and the American Counseling Association (ACA) Code of Ethics (ACA, 2014) establish the broader professional obligations within which neuroregulation practice is situated. The Association for Applied Psychophysiology and Biofeedback (AAPB) Code of Ethics similarly governs practitioners who hold AAPB membership. Across all of these codes, the foundational obligation is consistent: the practitioner bears responsibility for the welfare of the client, the

accuracy of their professional conduct, and the transparency of their methods.

The APA Ethics Code (2017) opens with five aspirational general principles that describe the character of ethical practice. These principles are not enforceable rules but ideals intended to guide practitioners toward the highest standards of their profession. Principle A (Beneficence and Nonmaleficence) calls on psychologists to benefit those with whom they work and to take care to do no harm. Principle B (Fidelity and Responsibility) establishes that psychologists uphold professional standards, take responsibility for their actions, and fulfill their commitments to those they serve and to the broader scientific and professional community. Principle C (Integrity) calls for accuracy, honesty, and truthfulness in all professional roles and the avoidance of deceptive methods. Principle D (Justice) calls for fairness and the exercise of reasonable judgment to protect clients from the practitioner's own biases and limitations. Principle E (Respect for People's Rights and Dignity) calls for respect for the rights of individuals to privacy, confidentiality, and self-determination.

The ACA Code of Ethics (2014) articulates six core values in its preamble: autonomy, nonmaleficence, beneficence, justice, fidelity, and veracity. Autonomy centers the client's right to make informed, self-determined choices about their own care. Nonmaleficence and beneficence together establish the dual obligation to avoid harm and to actively promote client welfare. Justice calls for fair and equitable treatment across all clients and populations. Fidelity reflects the obligation to be honest and trustworthy in the therapeutic relationship. Veracity extends this to a specific commitment to truthfulness in all professional communications.

The Convergence Between Codes and the Four-Principle Framework

The APA and ACA codes share a common ethical core. Beneficence, nonmaleficence, autonomy, justice, and fidelity appear in both, reflecting the deep convergence of psychological and counseling ethics traditions across decades of parallel development. Beyond the shared core, the APA preamble adds integrity, fidelity and responsibility as a paired principle, and respect for people's rights and dignity. The ACA preamble adds veracity. All of these principles are relevant to AI use in clinical practice, and each receives substantive attention in the sections that follow.

For the purposes of the analytical framework in this article, four principles are identified as most directly operative in AI-related clinical decisions: beneficence and nonmaleficence, autonomy, justice, and fidelity. This reduction is made for analytical clarity, not because the remaining principles are unimportant. Each maps onto one or more of these four in the AI context: Integrity is a dimension of fidelity. Veracity is a dimension of both autonomy and fidelity. Respect for people's rights and dignity is a dimension of autonomy and justice. Fidelity and responsibility as a paired APA principle encompasses both the individual and community-facing dimensions of the fidelity obligation addressed here.

The Four Principles and Their AI Applications

Beneficence and nonmaleficence ask three questions of any AI use: Does this serve the client's welfare? What are the known and reasonably foreseeable harms? How are they minimized? AI tools that claim clinical benefit without adequate evidence, or that generate errors the practitioner does not catch before they affect the client's care, create potential harm that the nonmaleficence obligation requires the practitioner to address.

Autonomy requires that clients have the information they need to make genuine, informed choices about AI involvement in their care. Meaningful autonomy is not satisfied by a signature on a consent form; it requires honest disclosure of how AI tools are used, what data they process, and what role they play in clinical decisions. The ACA's veracity principle reinforces this: Truthfulness in the therapeutic relationship includes transparency about the methods used in the client's care.

Justice requires attention to whether AI use in a given practice systematically provides advantages to some clients over others. AI tools can deepen existing disparities if access to higher-quality AI-assisted services depends on ability to pay, if the training data underlying a tool underrepresents particular populations, or if the practitioner applies AI reliance differentially across clients. APA Principle E (Respect for People's Rights and Dignity) reinforces this: Practitioners must be aware of and actively counteract the ways their tools may treat different clients unequally.

Fidelity calls on practitioners to be honest, trustworthy, and transparent in all professional communications and relationships. APA Principle B frames this alongside responsibility to the broader community of practice: The practitioner's

accountability does not stop at the consulting room door. APA Principle C (Integrity) adds the specific expectation that practitioners not use methods that are deceptive or that misrepresent the nature of their services to clients, colleagues, payers, or the public. Together, these principles establish that transparency about AI involvement is not merely good practice; it is an ethical obligation.

AI Applications in Neuroregulation Practice:

A Taxonomy

Before a framework for ethical reasoning about AI use can be applied, the range of AI applications currently entering neuroregulation practice must be understood. Five categories are identified here, each defined by the kind of clinical function AI is performing and the ethical question that function raises. This taxonomy is not exhaustive; the tool landscape is developing continuously.

Automated QEEG Analysis and Interpretation

A growing number of commercial platforms accept raw EEG data and return statistical analyses, normative comparisons, and clinically framed interpretive narratives. The output may include a suggested diagnostic impression, identified areas of concern, and protocol recommendations—all generated with minimal visible reasoning. What these platforms claim is efficiency, consistency, and access to normative databases larger than any individual practitioner could accumulate. The ethical question is foundational: How does a practitioner take professional responsibility for a recommendation whose reasoning they cannot inspect? The obligation to practice within the boundaries of competence (BCIA Code; APA Principle B) applies directly. A tool whose interpretive logic is unavailable for evaluation places the practitioner in a position of professional accountability without the information necessary to discharge it.

Protocol Recommendation Systems

Tools in this category generate training parameters from assessment data, applying algorithms trained on large case datasets to produce individualized protocol suggestions. The ethical question is clinical authority. When the system's recommendation conflicts with the practitioner's own clinical read, who holds authority, and on what basis is the disagreement resolved (e.g., a qEEG-generated protocol recommends training Fp1 and Fp2 eyes closed)? The practitioner who cannot explain their reasoning for overriding a system they cannot inspect is in a different position than the practitioner

who overrides a colleague's suggestion on the basis of clinical evidence.

Adaptive Neurofeedback Algorithms

Some neurofeedback systems adjust training parameters in real time based on client response data, shifting frequencies, thresholds, or montages without explicit practitioner direction. The ethical question is whether the practitioner remains the director of treatment. A system that adapts autonomously based on its own algorithmic logic is not simply executing the practitioner's clinical plan; it is extending or modifying it. The informed consent obligation (APA Principle E; ACA Autonomy) requires that clients understand this when it is happening. The competence obligation requires that practitioners understand the basis on which adaptations are made.

AI-Assisted Documentation

AI scribes, note-drafting tools, interpretive report generators, and AI-drafted client communications have spread rapidly across health care and are entering neuroregulation practice. Submission or sharing of a clinical document constitutes a professional attestation that its contents are accurate. That attestation does not transfer to the tool that generated the draft. Sherlin and Longo (2025), examining ethical billing practices in neurofeedback, established that the practitioner of record bears full accountability for the accuracy of submitted documents regardless of how they were produced. The same principle applies with full force to AI-generated documentation. As demonstrated in the literature on AI hallucination (discussed below), a document that reads fluently and professionally may contain fabricated citations, invented clinical characterizations, or inaccurate representations of the client's history. Additionally, in the absence of a client's complete physical and mental health history and relevant life experiences (e.g., traumatic events), AI cannot accurately generate a clinical history or summary.

Consumer-Facing AI Applications

Mental health chatbots, AI companions, and AI-powered self-help platforms are being used by clients at scale, independent of any professional relationship. McBain et al. (2026), in a nationally representative survey of adolescents and young adults published in *JAMA Pediatrics*, found that approximately one in five youth in the United States aged 12 to 21 reported using AI chatbots for mental health advice in 2025, with nearly two-thirds of those users reporting they had not disclosed this use to anyone. The client who arrives having processed

their experience through an AI tool may bring a narrative that has been shaped by that interaction in ways that affect assessment and treatment.

A Risk Continuum Framework

The ethical challenge in AI use is not binary. Risk in AI clinical applications vary continuously, and the review standards, disclosure obligations, and independence requirements that follow from the ethical principles above are proportional to where a specific use sits on that continuum. Three dimensions can be assessed for any specific use case.

The Three Dimensions of Risk

The first dimension is opacity: How much can the practitioner actually inspect, evaluate, and understand about how the tool arrived at its output? A tool that explains its reasoning in terms the practitioner can evaluate is categorically different from a system that returns a recommendation without any account of how it was reached. Opacity interacts directly with the competence obligation. The practitioner cannot exercise independent professional judgment about outputs they cannot evaluate.

The second dimension is clinical consequence: What is at stake if the output is wrong? The severity of potential errors, and whether they are recoverable once made, determine how much independent verification is required. An error in a protocol recommendation that results in adverse clinical effects (e.g., a client who consulted AI independently and received a recommendation to train delta up), or an inaccurate characterization of a client's history in a submitted insurance document, has consequences that a formatting error in an administrative message does not.

The third dimension is distance from oversight: How far is the practitioner from the decision point? A practitioner who reviews and consciously approves every AI output before it affects a client is closer to the decision than one whose practice has a system running between sessions, adjusting parameters, or sending communications automatically. Distance from oversight determines whether errors can be caught and corrected before they affect the client.

Application of the Continuum

At the low-risk end of the continuum, AI performs functions that are primarily logistical, transparent in their operation, and at a short distance from practitioner review. Drafting an appointment reminder that the practitioner reads before sending,

checking grammar and structure in a document the practitioner wrote, or organizing a literature search the practitioner will evaluate: These uses involve low opacity, low clinical consequence, and minimal distance from oversight.

At the high-risk end, AI performs functions that are clinically consequential, opaque in their reasoning, and at significant distance from practitioner oversight. Generating an interpretive qEEG narrative that is submitted to a third party without independent verification or running an adaptive algorithm that adjusts training parameters between sessions without explicit practitioner review: These uses combine high opacity, high clinical consequence, and real distance from oversight. They do not necessarily require abandonment of the tool; they require a review and verification standard proportional to the risk.

The practitioner's practical question for any new AI use is: Where does this sit? The answer determines what responsible use requires, what the review standard must be, what disclosure is owed to the client, and how much independence the practitioner must maintain in verifying outputs.

Informed Consent and Transparency

The Consent Question in the AI Era

Informed consent in clinical practice has always required that clients receive honest, comprehensible information about the methods used in their care, sufficient to allow a genuine choice about whether to proceed. AI does not change this requirement; it changes the content of what must be disclosed. The client has a right to know that an AI tool is involved in their care, what it does, and what role it plays in clinical decisions that affect them. This is an expression of the autonomy principle at the center of both the APA and ACA codes.

The current research suggests that transparency about AI involvement in healthcare affects both the accuracy of clinical encounters and the quality of the therapeutic relationship. Investigators have found that patients tend to be less forthcoming with AI-based systems than with human practitioners, particularly about sensitive health information, which creates downstream clinical consequences independent of any AI accuracy question (Nagata et al., 2026). The client who does not know they are interacting with or through an AI system cannot make an autonomous decision about how much to share.

Proportionate Disclosure

The practical question for practitioners is not whether to disclose AI involvement but how. Proportionate disclosure means providing the client with enough information to understand what is happening and to make a genuine choice, without overwhelming them with technical detail that does not serve their decision-making. For most clients and most AI uses, this requires honest, plain-language acknowledgment that an AI tool assists with certain functions in the practice, what those functions are, how it affects what the client receives, and that the practitioner reviews all AI-assisted outputs before they affect the client's care.

For AI tools that play a more direct role in clinical decision-making, including qEEG analysis or protocol recommendation systems, more specific disclosure is warranted: what the tool is, what role its output plays in clinical decisions, and that the practitioner exercises independent clinical judgment in reviewing and acting on its recommendations. Disclosure language for most documentation-assistance uses can be brief. The following meets the standard for common contexts: "This practice uses AI tools to assist with the drafting of certain documents, including session notes and clinical summaries. All such documents are reviewed and verified by your clinician before they are submitted or shared."

Disclosure Beyond the Client

The fidelity and integrity obligations in both codes extend disclosure requirements beyond the client relationship. When AI tools contribute to documents submitted to insurers, AI involvement is relevant to the accuracy attestation the practitioner makes in submitting the document. In supervision and consultation, modeling transparency about AI use is both an ethical obligation and a professional responsibility. In publications and professional presentations, the norm against undisclosed AI authorship is now well established; the practitioner who uses AI to draft a manuscript or presentation and does not disclose this is misrepresenting the nature of their work to the professional community.

This article models the transparency it recommends. The authors used Claude (Anthropic, 2026), a generative AI tool, to assist with brainstorming, literature organization, and formatting during the preparation of this manuscript. All substantive content, clinical analysis, ethical conclusions, the risk continuum framework, and the five-element policy framework are entirely the authors' own.

Competence, Scope, and the Limits of Delegated Judgment

Competence When the Reasoning is Not Available

Ethics codes across BCIA, ISNR, APA, and ACA require practitioners to practice within the boundaries of their competence. Competence has always included knowing the limits of one's instruments. AI creates a new version of this longstanding obligation. A practitioner cannot exercise independent professional judgment about an AI output whose reasoning they cannot inspect, evaluate, or explain to a client or colleague. The practitioner who cannot trace the reasoning must be able to verify the conclusion through independent means.

The accountability asymmetry the practitioner faces in this situation is structurally significant. Sherlin and Longo (2025) documented in the billing ethics context that practitioners bear full legal and professional accountability for submissions made under their credentials, while vendors are frequently shielded by contractual terms that limit liability for clinical decisions made on the basis of their tools. The AI tool will not stand with the practitioner before a licensing board, in a malpractice proceeding, or in a conversation with a client who was harmed by an error the tool generated.

Automation Bias: The Mechanism of Competence Erosion

The most immediate competence risk in AI-assisted practice is automation bias: The documented tendency to defer to algorithmic output without sufficient critical evaluation (Cummings, 2017; Skitka et al., 1999). Automation bias is not a character flaw; it is a cognitive pattern that affects highly trained experts consistently. Three mechanisms are particularly relevant in clinical contexts.

The first is cognitive surrender. When the algorithm has an answer, the practitioner may stop generating one. Cabitza et al. (2021) documented that decision support systems can reduce independent clinical reasoning even among experienced practitioners. The second is the flattery loop. AI systems calibrated for user engagement are structurally inclined to affirm rather than challenge. Cheng et al. (2026), in a preregistered study of 11 major AI models published in *Science*, found that AI systems affirmed users' actions 49% more often than human respondents, even when those actions involved harm or were deceptive. A single interaction with a sycophantic AI was sufficient to increase users'

conviction that they were right and reduce their willingness to reconsider their position. The third mechanism is efficiency pressure. Faster output at greater volume within the same time budget means less review per item.

Deskilling and the Longer-Term Concern

Beyond the session-level effects of automation bias lies a longer-term concern: the erosion of the clinical reasoning skills that effective neuroregulation practice requires. Natali et al. (2025), in a mixed-method review published in *Artificial Intelligence Review*, examined AI-induced deskilling across medical specialties and found consistent evidence of skill erosion across physical examination, differential diagnosis, clinical judgment, and physician-patient communication, with clinician over-reliance on AI identified as the primary contributing mechanism. Vaccaro et al. (2024), in a systematic review and meta-analysis published in *Nature Human Behaviour* examining 106 experimental studies, found that human-AI combinations performed on average significantly worse than the best of either humans or AI alone. The corrective is not to avoid AI use but to maintain deliberate practice alongside it.

The Scope Question

Competence in the AI era requires a clear answer to a question that becomes urgent only when it has not been answered in advance: What will this practitioner not delegate to an algorithm, under any circumstances? For neuroregulation practice, this boundary includes at minimum: the determination of clinical appropriateness for neurofeedback at intake; the integration of intake assessment findings into a clinical conceptualization; the interpretation of client responses that suggest adverse effects or contraindications; and any determination communicated to a third party under the practitioner's professional credentials as a statement of clinical fact.

The clinician should not use AI to make clinical determinations about any disorder not within the clinician's scope of practice.

A boundary that cannot be stated clearly in advance cannot be held in the moment. The practitioner who has not defined what they will not delegate before sitting down with an AI-generated protocol recommendation is making that determination under conditions that are not conducive to independent judgment: time pressure, confirmation bias, and the cognitive weight of an authoritative-sounding alternative already visible on the screen.

Data Privacy, Vendor Accountability, and HIPAA

What Happens to Client Data

Every AI tool that processes information about a client creates a data pathway that the practitioner is responsible for understanding. Cloud-based AI platforms route data to servers operated by vendors or their subcontractors. That data may be stored, analyzed, used to improve the model, shared with affiliated entities, or retained indefinitely depending on the platform's terms of service. The sensitivity of neuroregulation-specific data compounds this concern: EEG recordings, qEEG reports, clinical notes, and treatment summaries are protected health information (PHI) under Health Insurance Portability and Accountability Act (HIPAA, 1996), and their unauthorized use or disclosure creates legal liability as well as ethical violation.

Many practitioners who have adopted AI tools have not investigated these data pathways. The vendor's marketing materials are typically reassuring in tone while the actual terms of service may differ materially from the marketing characterization. Practitioners are advised to read primary vendor documents directly rather than relying on marketing summaries before routing client data through any platform.

HIPAA, Business Associate Agreements, and the Practitioner's Obligation

Under HIPAA, any entity that handles PHI on behalf of a covered entity in the course of providing a service is a business associate and must execute a Business Associate Agreement (BAA). Many AI platforms used in clinical contexts do not offer BAAs or offer them only on enterprise plans. Some platforms explicitly exclude clinical and regulated use from their standard terms of service. The practitioner who routes client PHI through a platform without a BAA is out of compliance with HIPAA regardless of the marketing language on the platform's homepage.

The accountability asymmetry documented in the billing ethics context (Sherlin & Longo, 2025) applies with equal force here. Vendors whose terms limit liability for clinical decisions made on the basis of their tools are structurally protected from the consequences of their platform's errors. The practitioner, as the covered entity responsible for PHI, is not. Due diligence before adopting any AI tool that touches client data is not optional; it is a HIPAA compliance requirement and an ethical obligation.

A Due Diligence Framework

Before adopting any AI tool that will process client information, a practitioner should be able to answer the following questions: Where does client data go when it is processed by this tool? Is the data stored after processing? If so, for how long and where is it stored? Is client data used to train or improve the model? Is a BAA available and has it been executed? What does the platform's privacy policy actually state about data use, retention, and sharing? What would happen to stored client data if the vendor were acquired, merged, or ceased operations?

AI Accuracy Risks: Hallucination, Automation Bias, and Collusion

Three distinct accuracy risks warrant systematic attention in AI-assisted clinical practice. They are not synonyms and they do not produce the same clinical effects. Each requires its own clinical response.

Hallucination

AI hallucination refers to the generation of false or invented content with the same fluency, confidence, and apparent precision as accurate content. The defining clinical problem is not that AI systems are occasionally wrong; it is that the fluency and apparent authority of the output provide no reliable signal of its accuracy. A qEEG interpretive narrative that contains a fabricated citation, an invented characterization of a normative pattern, or a misrepresentation of the client's history reads exactly the same as one that is accurate.

The evidence for the scope of this problem is substantial. Van Dis et al. (2023) identified fabricated references in published academic papers that had incorporated AI-generated content and called for priority attention to accuracy verification in AI-assisted scientific writing. Sallam (2023), in a systematic review, found that AI-generated responses to medical questions carried significant and variable rates of error across platforms and question types. For neuroregulation practice specifically, the hallucination risk is most acute in AI-generated interpretive reports and protocol narratives. The practitioner's signature attests to accuracy; the tool's fluency does not guarantee it. Content-level review, meaning active verification of clinical claims against the practitioner's own knowledge and the client's record rather than reading for comprehension, is the standard that the hallucination literature demands.

Automation Bias in the Accuracy Context

Automation bias also functions as a distinct accuracy risk. The practitioner who has internalized trust in an AI system may not apply the same critical scrutiny to AI-generated content that they would apply to a colleague's draft. Errors that would be immediately apparent in a document a human colleague submitted—a misidentified frequency band, an incorrect normative comparison, or a clinically implausible characterization—may pass without detection in an AI-generated document whose overall fluency and structure communicate authority. The sophistication of AI language generation is not evidence of AI clinical reasoning accuracy, and the two should not be conflated in the practitioner's review process (Vaccaro et al., 2024).

AI Collusion

Tahseen (2026) introduced the concept of AI collusion in a *JMIR Mental Health* analysis, defining it as the structural reinforcement of unreliable user input as ground truth. Collusion is distinct from hallucination, which is the AI generating false content independently. In collusion, the AI takes distorted or inaccurate input from the user and returns it amplified, validated, or repackaged in ways that feel authoritative. The error does not originate in the AI system; it enters through the human.

The clinical relevance is direct, and it connects to the sycophancy literature reviewed above. AI systems calibrated for user engagement are structurally inclined to affirm user input rather than challenge it (Cheng et al., 2026). A user who presents a distorted self-narrative to an AI chatbot receives back a response that takes that narrative as reliable, builds on it, and in doing so reinforces it. For clients in mental health treatment, where self-report is often distorted by the very conditions being treated, this creates a specific risk: The client may arrive at the clinical session with a self-understanding that has been structurally confirmed by repeated AI interaction, making it more resistant to the therapeutic examination that accurate clinical work requires. Asking routinely about AI use as part of intake and ongoing clinical assessment is the practical response this literature supports.

Client AI Use: What Is Coming Into the Room

Clients are using AI for mental health purposes at significant scale, independently of any clinical relationship, and bringing the results into treatment in ways that affect assessment and therapeutic work. McBain et al. (2026), in a nationally representative survey of 1,009 adolescents and

young adults aged 12 to 21, found that approximately one in five reported using AI chatbots for mental health advice, with nearly two-thirds of those users reporting they had not disclosed this use to anyone. Among users, 43% reported seeking AI mental health advice at least monthly. Nagata et al. (2026), writing in *JAMA Pediatrics* on risks and benefits of generative AI in adolescent health, noted that nondisclosure to clinicians and parents is common and that AI tools are reshaping the mental health information ecosystem largely outside of professional awareness.

Clients do not typically volunteer this information in clinical settings. Adding a standard question about AI use to intake assessment and periodically revisiting it during treatment normalizes the topic and provides clinically relevant information that is otherwise unavailable. Documentation of client AI use as a clinical variable, not merely a lifestyle fact, positions the practitioner to track its influence on the therapeutic process.

The collusion dynamic described above is particularly relevant for the neuroregulation practitioner seeing clients whose self-narratives have been shaped by AI interaction. A client who has spent weeks processing their anxiety, attention difficulties, or trauma history through AI tools calibrated for engagement and affirmation may present with a self-account that is more defended than one they would have developed without that reinforcement. The therapeutic skills required to work with this presentation are not new. What is new is the source of the shaping, and the recognition that it is happening requires the practitioner to ask about it.

A Five-Element Practice Policy Framework

The Case for a Written Policy

A written, practice-specific AI policy serves three functions that no informal practice norm can replicate. It protects clients by ensuring that AI is used consistently, with appropriate review, and with the disclosure and consent their autonomy requires; it protects the practitioner by documenting that AI use in the practice was deliberate, reviewed, and grounded in ethical principles, rather than ad hoc and unexamined; and it positions the practitioner as a participant in developing the field's emerging standards. A policy does not need to be lengthy or legalistic. It needs to be honest, specific, and reviewed regularly as the tool landscape changes.

The Five Elements

The framework is organized around five elements, each addressing a domain in which the ethical principles reviewed above create specific operative obligations. Table 1 presents the framework in condensed form for clinical reference. The following paragraphs describe each element in the depth required for practical implementation.

Transparency and consent addresses which clients are informed about AI involvement in their care, in what language, at what point in the clinical relationship, and with what opportunity to ask questions or decline. The implementation task is straightforward: an AI disclosure addendum to the intake consent form, a brief standard script for the intake conversation, and a protocol for what to do when a client asks a question about AI that the practitioner cannot answer on the spot.

Data handling addresses what client data each AI tool processes, where that data goes, whether it is stored and for how long, whether it is used to train models, and whether a BAA is in place. The implementation task is a tool audit: a systematic review of each active AI tool's data practices, documented and revisited annually or when vendor terms change.

Documentation standards address which document types AI may assist with and which are off limits, the review standard required before any AI-assisted document is submitted or shared, and how AI involvement is recorded in the clinical record. The practitioner who has defined the review standard in writing has a standard to apply; the practitioner who has not defined it is making the determination under pressure at the moment of submission.

The clinical judgment boundary addresses what the practitioner will not delegate to an algorithm under any circumstances, how the practitioner handles disagreement between the tool's output and their own clinical assessment, and what signals require independent verification before acting on AI recommendations. The implementation task is to complete one sentence before adopting any AI tool that plays a role in clinical decisions: "Regardless of what this tool recommends, I will not proceed without independent clinical analysis of the following."

Ongoing evaluation addresses the questions asked before adopting a new AI tool, the schedule and criteria for reviewing tools already in use, and the conditions that would trigger removal of a tool from

Table 1

Five-Element Practice AI Policy Framework for Neuroregulation Clinicians

Policy Element	Key Questions	Relevant Ethics Standards	Common Gap in Practice	Suggested First Step
Transparency and Consent	Which clients are informed, in what language, at what point, and with what opportunity to ask questions or decline	APA Principles A, E; ACA Autonomy, Veracity; APA AI Guidance (2025)	No consent language; no consistent disclosure practice for AI-assisted communications or reports	Develop an AI disclosure addendum for intake consent; establish timing protocol
Data Handling	What client data each tool processes, where it goes, storage duration, whether it trains models, BAA status	HIPAA Privacy and Security Rules; BCIA Code of Ethics; ISNR Code of Ethics	Unknown or unverified vendor data practices; no BAA on file for tools that process PHI	Audit each active tool; request or execute BAAs; document findings in practice records
Documentation Standards	Which document types AI may assist with; review standard before submission; how AI involvement is recorded in the clinical record	APA Principle C (Integrity); ACA Fidelity; APA AI Guidance (2025); ISNR Code	AI-drafted documents submitted without content-level review; no record notation of AI involvement	Define explicit review standard in writing; create a standard clinical record notation template
Clinical Judgment Boundary	What the practitioner will not delegate to an algorithm; how tool-clinician disagreement is handled; what signals require independent verification	BCIA Code of Ethics; APA Principles A, B; ACA Nonmaleficence, Fidelity	Following AI recommendations without independent clinical analysis, particularly for qEEG interpretation and protocol selection	Define the boundary explicitly and in writing before adopting each tool; establish an override and documentation protocol
Ongoing Evaluation	Questions asked before adopting a new tool; review schedule and criteria for tools in use; conditions that would trigger removal	APA Principle B (Fidelity and Responsibility); APA AI Guidance (2025); BCIA and ISNR competence standards	No pre-adoption evaluation process; tools adopted on vendor recommendation without independent vetting	Create a pre-adoption checklist; schedule an annual review of all active tools; designate a responsible party

Note. BAA = Business Associate Agreement. PHI = Protected Health Information. Ethics code references are illustrative; multiple codes may apply depending on the practitioner's licensure and certifications. APA = American Psychological Association (2017, 2025). ACA = American Counseling Association (2014). BCIA = Biofeedback Certification International Alliance. ISNR = International Society for Neuroregulation and Research.

practice. A platform that met due-diligence standards at adoption may change its terms of service, its ownership, its data practices, or its technical architecture. An annual review of all active tools is the minimum; any change in vendor terms or ownership should trigger an immediate review.

Implications for Training, Supervision, and Field-Wide Standards

Training and Supervision

The trainee and newly certified practitioner represent a population of heightened concern in the AI era. The deskilling evidence reviewed above is particularly consequential for individuals who are still developing independent clinical reasoning: If AI substitutes for the practice that consolidates clinical reasoning skills, those skills may never develop fully. Supervisors bear a specific obligation in this context. Competence assessment needs to include evaluation of whether supervisees are developing independent clinical reasoning or learning to rely on algorithmic output as a substitute for it. The supervisor who asks, “What is your own interpretation of this qEEG before we look at the platform's recommendation?” is doing something that has always been good supervisory practice and that is more important now than it has ever been.

BCIA and ISNR Standard-Setting Opportunities

Both BCIA and ISNR are positioned to provide field-specific AI ethics guidance before regulatory requirements arrive from outside the profession. The APA's 2025 guidance for psychologists provides a template that is directly adaptable to the neuroregulation context. Practitioners who wait for formal regulatory requirements will receive standards written by people who do not practice neuroregulation. Practitioners who engage with the organizations governing their certification and membership now will have the opportunity to help write those standards in terms that reflect clinical reality.

A Research Agenda

Several empirical questions emerge from this framework that the field is well positioned to address. The prevalence and nature of AI use among neuroregulation practitioners is a foundational descriptive question. The accuracy and reliability of specific AI-assisted qEEG interpretation tools relative to expert practitioner consensus is an evaluative question with direct clinical implications. The effects of AI-assisted protocol selection on client outcomes is an outcomes question that the field's research infrastructure is capable of addressing. The

effects of AI use on practitioner competence development over time, particularly in trainees, is a training research question with significant implications for certification standards.

Limitations

This article has several limitations that readers should consider. The AI tool landscape is changing faster than the publication cycle of academic literature. Specific tools, vendor practices, regulatory developments, and empirical findings described here reflect the state of the field at time of writing and may be superseded before or shortly after publication. The frameworks are designed to be durable and will remain applicable as the technology evolves.

The evidence base drawn on here is primarily from medicine and mental health broadly. Neuroregulation-specific evidence on AI accuracy, practitioner use patterns, and clinical outcomes with AI-assisted tools is currently limited. The frameworks proposed are extrapolations from a broader evidence base, not empirical findings from within the neuroregulation specialty. This limitation is itself a call for the field-specific research agenda described above.

This is a perspectives article, not a systematic review. The reference base reflects the authors' clinical and theoretical judgment about the most relevant current literature, not a comprehensive search strategy.

Conclusion

Three claims run through this article and are worth stating plainly at its close.

Existing ethics principles are sufficient. The principles that have guided neuroregulation practice across decades of technological change (beneficence and nonmaleficence, autonomy, justice, and fidelity) do not require revision for the AI era. They require deliberate application to it. The practitioner who asks whether this AI use serves the client's welfare, whether the client understands what is happening and has a genuine choice, whether all clients are being treated equitably, and whether they are being transparent about their methods is applying the framework that both the APA and ACA codes provide.

Risk lives on a continuum. Not all AI use in neuroregulation practice is equivalent. The practitioner who uses AI to format a reminder and

reads it before sending is not in the same ethical position as the practitioner who implements an AI-generated qEEG protocol recommendation without independent analysis. The risk continuum framework provides the structure for identifying where any specific use sits, and the review standard, disclosure obligation, and independence requirement that position demands.

The practitioner remains accountable; the tool never is. No vendor, platform, or algorithm will appear before a licensing board alongside a practitioner who is the subject of a complaint. The client who was harmed by an AI error that the practitioner did not catch will not have recourse against the tool. The professional obligation rests with the professional. The practitioner who uses AI intentionally, reviews AI output independently, discloses AI involvement honestly, and limits AI's role to functions where it can be verified is not in the same position as the practitioner who adopted tools uncritically. The field needs the former.

The window for neuroregulation practitioners to help write the standards that govern their own use of AI is open. This article is offered as a resource for that work.

Author Declaration

Neither author has relevant financial interests or relationships to disclose. The authors used Claude (Anthropic, 2026), a generative AI tool, to assist with brainstorming, literature organization, and formatting during the preparation of this manuscript. The tool assisted with structure, clarity, and the organization of source materials. All substantive content, clinical analysis, ethical conclusions, the risk continuum framework, the five-element practice policy framework, and all intellectual contributions presented in this article are entirely the authors' own. At the time of submission, the first author served as President of ISNR and Chair of the Board of BCIA; both authors served as an Advisory Board member of the International qEEG Certification Board (IQCB). These roles are disclosed in the interest of transparency, as the manuscript addresses certification standards, professional guidelines, and organizational practices relevant to these bodies.

References

American Counseling Association. (2014). ACA code of ethics. <https://www.counseling.org/docs/default-source/default-document-library/ethics/2014-aca-code-of-ethics.pdf>

American Psychological Association. (2017). Ethical principles of psychologists and code of conduct (2002, amended effective June 1, 2010, and January 1, 2017). <https://www.apa.org/ethics/code>

American Psychological Association. (2025). Ethical guidance for artificial intelligence in the professional practice of health service psychology. <https://www.apa.org/topics/artificial-intelligence-machine-learning/ethical-guidance-ai-professional-practice>

Anthropic. (2026). Claude [Large language model]. <https://claude.ai/>

Biofeedback Certification International Alliance. (n.d.). BCIA code of ethics. <https://www.bcia.org/bcia-professional-standards-ethical-principles>

Cabitza, F., Campagner, A., & Bitterman, D. (2021). The need to separate the wheat from the chaff in medical informatics. *International Journal of Medical Informatics*, 153, Article 104510. <https://doi.org/10.1016/j.ijmedinf.2021.104510>

Cheng, M., Lee, C., Khadpe, P., Yu, S., Han, D., & Jurafsky, D. (2026). Sycophantic AI decreases prosocial intentions and promotes dependence. *Science*, 391(6792), Article eaec8352. <https://doi.org/10.1126/science.aec8352>

Cummings, M. L. (2017). Automation bias in intelligent time critical decision support systems. In D. Harris (Ed.), *Decision making in aviation* (1st ed.). Routledge. <https://doi.org/10.4324/9781315095080-17>

Health Insurance Portability and Accountability Act (HIPAA) of 1996, Pub. L. No. 104-191, 110 Stat. 1936 (1996).

Horvitz, E., & West, R. (2026). A narrowing window to understand AI. *Science*, 392(6802), Article 1003. <https://doi.org/10.1126/science.aei3167>

International Society for Neuroregulation and Research. (n.d.). ISNR code of ethics and standards of practice. <https://isnr.org/interested-professionals/isnr-code-of-ethics>

McBain, R. K., Cantor, J. H., Breslau, J., Diliberti, M., Zhang, L. A., Zhang, F., Burnett, A., Kofner, A., Rader, B., Pataranutaporn, P., Stein, B. D., Mehrotra, A., & Yu, H. (2026). AI chatbot use and disclosure for mental health among US adolescents and young adults. *JAMA Pediatrics*, 180(8), 884–890. <https://doi.org/10.1001/jamapediatrics.2026.2015>

Nagata, J. M., Memon, Z., Huang, O., & Moreno, M. A. (2026). Adolescent health and generative AI: Risks and benefits. *JAMA Pediatrics*, 180(1), 7–8. <https://doi.org/10.1001/jamapediatrics.2025.4502>

Natali, C., Marconi, L., Dias Duran, L. D., & Cabitza, F. (2025). AI-induced deskilling in medicine: A mixed-method review and research agenda for healthcare and beyond. *Artificial Intelligence Review*, 58(11), Article 356. <https://doi.org/10.1007/s10462-025-11352-1>

Sallam, M. (2023). ChatGPT utility in healthcare education, research, and practice: Systematic review on the promising perspectives and valid concerns. *Healthcare*, 11(6), Article 887. <https://doi.org/10.3390/healthcare11060887>

Sherlin, L. H., & Longo, R. (2025). Clarifying the code: Historical foundations, current practices, and ethical billing in neurofeedback and qEEG. *NeuroRegulation*, 12(2), 138–153. <https://doi.org/10.15540/nr.12.2.138>

Skitka, L. J., Mosier, K., & Burdick, M. D. (1999). Does automation bias decision-making? *International Journal of Human-Computer Studies*, 51(5), 991–1006. <https://doi.org/10.1006/ijhc.1999.0252>

Tahseen, H. (2026). When AI colludes: Clinical reliability of training and preference data as a trustworthy-AI criterion. *JMIR Mental Health*, 13, Article e96894. <https://doi.org/10.2196/96894>

Vaccaro, M., Almaatouq, A., & Malone, T. (2024). When combinations of humans and AI are useful: A systematic review and meta-analysis. *Nature Human Behaviour*, 8(12), 2293–2303. <https://doi.org/10.1038/s41562-024-02024-1>

Van Dis, E. A. M., Bollen, J., Zuidema, W., van Rooij, R., & Bockting, C. L. (2023). ChatGPT: Five priorities for research. *Nature*, 614(7947), 224–226. <https://doi.org/10.1038/d41586-023-00288-7>

Received: June 21, 2026

Accepted: September 1, 2026

Published: September 28, 2026