

PUBLICATION

OpenAI Updates Usage Policies: Key Considerations and Next Steps for Organizations Deploying AI

Authors: Alexandra P. Moylan, Alisa L. Chestler

November 18, 2025

OpenAI recently announced significant updates to its Usage Policies, explicitly clarifying that its tools – including ChatGPT, API-based services, and integrated products – are not substitutes for professional medical, legal, or other regulated advice. The changes apply to all OpenAI products and services, including enterprise and business accounts.

The revisions to the Usages Policies further serve as a reminder of the critical importance of utilizing subject matter expertise and human oversight over artificial intelligence (AI) solutions and their output. OpenAI's tools will continue to operate in the same manner, such that users can still query the large language model (LLM) for legal, medical, financial, and other advice. The revisions essentially attempt to limit liability for the output by expressly notifying users that the output is not a substitute for professional legal, medical, or financial advice, and by highlighting the potential that generative AI (GAI) output may be incorrect.

This alert provides organizations deploying OpenAI or other AI technology with key operational and compliance implications for consideration in addition to practical next steps for regulated entities.

Key Policy Changes

OpenAI's Usage Policies now expressly prohibit the use of its tools for the "provision of tailored advice that requires a license, such as legal or medical advice, without appropriate involvement by a licensed professional." Additional prohibitions include use of the products for "suicide, self-harm, or disordered eating promotion or facilitation," and "sexual violence or non-consensual intimate content."

OpenAI's updated policy on professional advice mirrors [similar policy revisions by Anthropic](#) announced in August of this year, both emphasizing the need for human oversight and clear policy safeguards when deploying AI in sensitive or regulated areas such as legal, health care, financial, and insurance, among others.

Alignment with the Current Legal and Regulatory Landscape

AI providers' updates to their usage policies align with the evolving U.S. legal landscape. California and other states have enacted new laws to ensure that GAI is only deployed in health care when overseen by licensed professionals and to prohibit the use of AI for professional advice unless this oversight exists. California's Assembly Bill 3030 (effective January 1, 2025) requires that any AI-generated patient communication containing clinical information include a clear disclaimer and instructions on how to contact a human provider – unless the message has been read and approved by a licensed professional. Similarly, Assembly Bill 489 (effective January 1, 2026) prohibits AI product developers and deployers from using any terminology or representations that falsely suggest licensure or certification in medicine, such as terms implying the system is operated by a doctor, nurse, or other credentialed health professional.

We previously covered the [pioneering legislation enacted by Illinois](#) with its passage of the Wellness and Oversight for Psychological Resources Act (effective 2025). The law prohibits the provision of therapy or psychotherapy services through AI unless the services are delivered by a licensed professional. The law restricts AI chatbots from performing therapeutic communication, making independent treatment

recommendations, or generating therapy plans without direct oversight from a clinician. Other states, including Nevada and Utah, have followed Illinois' lead in prohibiting or strictly regulating the provision of therapy through AI chatbots.

These laws reflect a regulatory consensus that aligns closely with OpenAI's and Anthropic's revised use policies. Indeed, recent high-profile lawsuits against AI developers involving chatbots marketed or used for companionship or therapeutic support continue to spur legislation. Emerging case law may redefine the boundaries of liability and acceptable use for AI solutions in high-risk or sensitive domains.

Organizations that deploy GAI in health care – whether as innovators, health systems, or digital health vendors – should maintain robust human involvement and transparent disclosures, ensuring that any AI-enabled clinical interaction is either overseen or ultimately reviewed by a licensed provider. This approach is designed to reduce patient risk, prevent consumer/patient confusion and harm when it comes to professional advice, and close regulatory gaps that could otherwise expose organizations to enforcement or liability for the unauthorized practice of medicine.

Implications for Organizations

For organizations deploying OpenAI or any other AI technology, these updates carry several operational and compliance implications:

- **Governance and Risk Management:** Companies should review their AI governance frameworks to ensure that AI tools are not positioned or marketed as providing regulated professional guidance, including medical, legal, or financial advice. Policies should make clear that human expertise and oversight by a licensed professional are required for all professional recommendations.
- **Acceptable Use Policies:** Enterprises integrating AI models through APIs or custom applications should update internal acceptable use policies to ensure proper subject matter and professional oversight over AI-generated output that constitutes professional advice or "high-risk use cases."
- **Training and Education:** Provide continuous employee training regarding permissible and prohibited uses of GAI tools in professional contexts.
- **Disclaimers and Client Communication:** Entities embedding AI technology into consumer- or client-facing interfaces (such as digital health chatbots or legal information tools) should update disclaimers and user terms to align with applicable use policies and maintain consistency with federal and state consumer protection laws.
- **Regulatory Compliance Alignment:** For health care users, these updates underscore HIPAA and FDA compliance boundaries, as outputs engaging in clinical interpretation or decision support may implicate regulatory oversight. Legal organizations must similarly ensure that reliance on AI-generated output does not constitute the unauthorized practice of law.

Related Reminders

The input of sensitive personal information, particularly protected health information (PHI) and confidential or privileged information, into any public/open AI tool should be strictly prohibited in all circumstances. AI models learn and generate output based on the data they process, creating a significant risk of unauthorized disclosure or a reportable breach if PHI or other personal data is exposed without appropriate business associate or data processing agreements.

Furthermore, the leakage of sensitive confidential trade secrets or corporate data remains an ongoing risk for all companies. There should be no expectation of privacy when using public/open tools, and exposing confidential data is a significantly underappreciated organizational risk.

All AI use cases should be carefully reviewed and examined, with particular attention to whether the tool is an internal enterprise deployment or a publicly accessible platform. For example, in health care, the ubiquitous use of ambient AI for medical record documentation requires robust training for providers to critically review the output.

In the legal space, numerous instances have shown AI tools "hallucinating" and fabricating answers, citing non-existent cases, statutes, or other legal precedent, often with confident but false assurances of accuracy. Indeed, there are numerous cases in the legal context regarding the misuse of GAI in preparation of legal pleadings, some of which have been tracked [here](#).

Recommended Next Steps

- Review current and planned integrations to confirm compliance with updated policy terms. The exercise should include OpenAI, Anthropic, and other tools whose terms may also have been amended, albeit with less publicity.
- Review and update internal policies and procedures for employees and vendors for using any AI tools in assisting with professional determinations.
- Update internal AI acceptable use and risk classification frameworks to reflect OpenAI's clarified restrictions.
- Reassess end-user terms, disclaimers, and documentation to ensure they align with OpenAI's revised representations about advice limitations.

OpenAI's clarification sends a clear signal that it seeks to delineate system-use boundaries between informational and professional advice. Organizations, especially those that operate in regulated domains, will need to mirror this clarity in their governance policies to mitigate compliance and liability exposure.

For more information or assistance on this topic, please contact [Alexandra P. Moylan, CIPP/US, AIGP](#), [Alisa L. Chestler, CIPP/US, QTE](#), or another member of Baker Donelson's [AI Team](#).