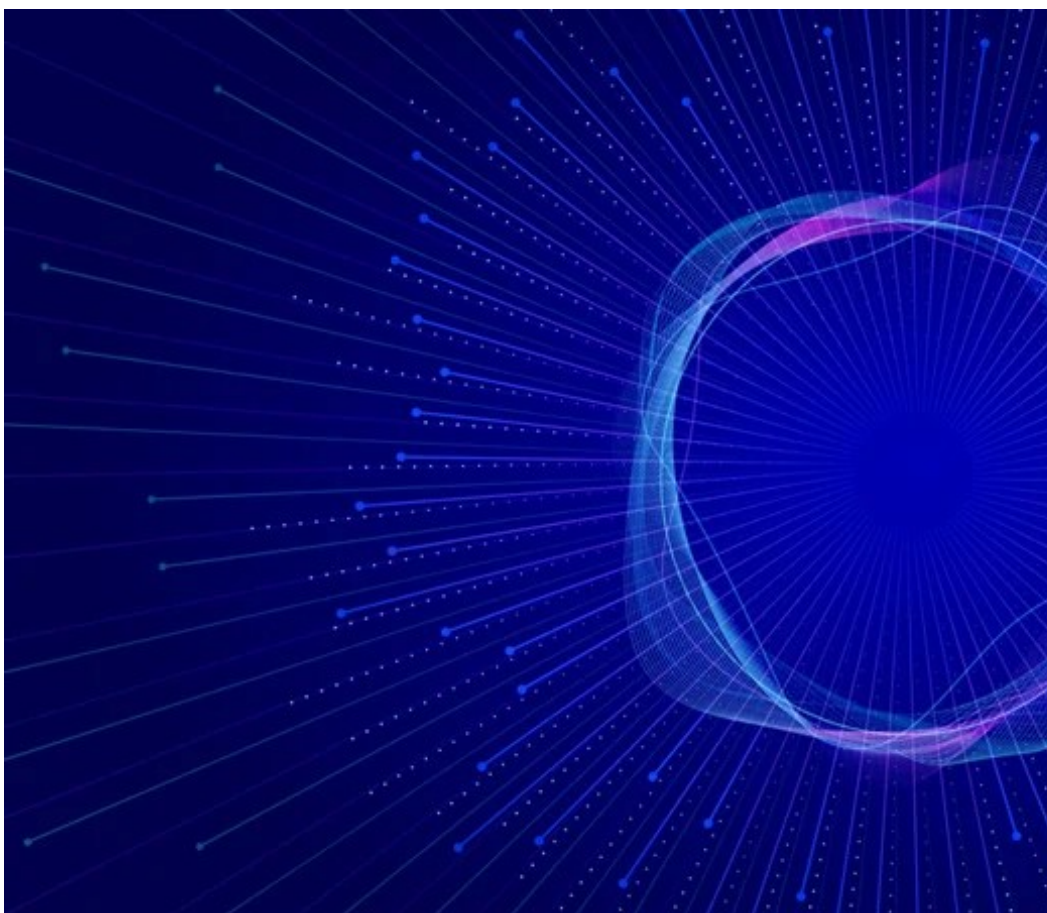


Artificial Intelligence Launching Agentic AI in an Uncertain U.S. Regulatory Landscape

February 11, 2025 – *The TMCA*

by Sarah Robertson, Jamie Nafziger, and Bianca Tillman



Related People

- Sarah Robertson

Related Capabilities

- Trademark, Copyright + Advertising
- Intellectual Property Litigation

Are you ready to begin adding AI Agents to your human teams? You will soon be getting requests to do so. While business teams will be wowed by what AI Agents and their subagents can do, the artificial intelligence regulatory environment is increasingly uncertain, and we advise caution.

Within its first few days, the new administration revoked the 2023 Executive Order on Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence (the “Old AI

Executive Order”) – which had been intended to mitigate risks associated with AI technologies by imposing safety guidelines such as a watermarking system for AI-generated content. Next, the president issued a new executive order Removing Barriers to American Leadership in Artificial Intelligence (the “New AI Executive Order”), providing White House staff 180 days to develop a plan “to sustain and enhance America’s global AI dominance in order to promote human flourishing, economic competitiveness, and national security.” Under the New AI Executive Order, federal agencies are required to immediately review and suspend, revise, or rescind any prior governance and safety measures undertaken to comply with the Old AI Executive Order that are not consistent with the New AI Executive Order.

Meanwhile, on January 21, the president announced a new \$500 billion private sector deal, “The Stargate Project,” which is a new joint venture that relies on SoftBank, OpenAI, Oracle, and MGX as initial equity funders in a bid to rapidly expand U.S. AI infrastructure by building massive new data centers, first in Texas then in other potential sites across the country. In the same week, leading big tech companies including Oracle and OpenAI each announced the next phase of the AI revolution with the release of new AI Agents designed to autonomously handle specific tasks and ultimately “join the workforce and materially change the output of companies.”

What Are AI Agents?

While traditional AI chatbots using large language models (“LLMs”) are designed to respond to users within the AI system, the latest generation of advanced AI Agents (aka AI Super-Agents or Agentic AI) are now empowered to interact with computer, network, and internet environments to automate tasks. Open AI states that its new AI Agent, Operator, “can be asked to handle a wide variety of repetitive browser tasks such as filling out forms, ordering groceries, and even creating memes.” Accordingly, Operator “can ‘see’ (through screenshots) and ‘interact’ (using all the actions a mouse and keyboard allow) with a browser, enabling it to take action on the web without requiring custom API integrations.” By combining the power of LLMs with new technologies such as the Oracle Cloud Infrastructure and retrieval-augmented generation, Oracle’s AI Agents can interact with enterprise data and can apparently be used by companies to recruit qualified job candidates, perform complex customer data analytics, optimize call centers, and expedite legal, financial, and academic research tasks.

AI Agents can be trained or instructed to follow the values of a user or business. Whether that can be consistently implemented remains to be seen. In testing, some AI systems have been shown to engage in “scheming” where they change results or take unexpected actions toward an ultimate goal set by the user. Although the usefulness of these technologies is undeniable, companies looking to deploy AI Agents should be aware of the legal risks and pitfalls that may accompany the use of these revolutionary new tools.

Legal Challenges and Mitigation Strategies

- **AI Agency Liability.** In July 2024, a California district court allowed a case against

HR and finance platform Workday to proceed, stating that an employer's use of Workday's AI-powered HR-screening algorithm may create direct liability for *both* the employer and Workday under the theory of agency liability. In this hiring discrimination case, the plaintiff alleged that his prospective employer delegated "traditional hiring functions, including rejecting applicants, to the algorithmic decision-making tools provided by Workday." The court found that by designing an AI technology to make decisions that would normally be made by a live employee, Workday should be treated as an agent of the employer for purposes of liability so long as the employer actually relied on the AI technology in its hiring process. While this case is still making its way through the courts, AI Agent vendors and deployers could subsequently be exposed to both civil and criminal liability based on the actions of the AI Agent, a theory that had previously only been applied by the courts to the actions of live humans.

- **AI Product Liability.** Given the increased responsibility and autonomy being granted to AI Agents, it is possible that developers and deployers of AI Agents could also be required to contend with product liability claims. Manufacturers and sellers can be held responsible for injuries caused by defective or unreasonably dangerous products. If an AI Agent makes a poor and costly decision, a plaintiff may claim that the AI Agent was defectively designed and/or that the developer failed to adequately warn the plaintiff of the AI Agent's limitations. This theory of liability is currently being tested in cases against AI-chatbot platform Character.AI. In the fall of 2024, the mother of a deceased minor in Florida brought a product liability case against the developers of Character.AI – an AI-chatbot role-playing platform that allows users to create and converse with AI-powered characters – alleging that her son died by suicide after becoming harmfully dependent on his relationship with his Character.AI companion "Dany." In December, parents in Texas brought additional product liability claims, stating that, without warning, Character.AI chatbots exposed their children to hypersexualized content, self-harm, and violent behaviors, with one chatbot allegedly encouraging a child to kill his parents when they tried to limit his screen time.
- **AI Contractual Considerations.** Businesses seeking to onboard AI Agents should implement clear contractual provisions to allocate and manage risk. When reviewing warranties, limitations of liability, and indemnification clauses, companies should want to know if and to what extent the AI Agent vendor will indemnify the company for an AI Agent's decision making, especially if those decisions are illegal or cause harm to end-users or others. Businesses should also consider who is responsible for training employees to use AI Agents safely and whether that training relates to the AI provider's contractual liability. In addition to requiring compliance with emerging state and international AI regulations such as the Colorado AI Act and the EU AI Act, companies should also carefully contemplate intellectual property ("IP") and data ownership, rules for algorithmic training inputs, and customization of standard liability shifting terms. Whereas typical software-as-a-service contracts usually address ownership of software and/or underlying data, AI Agent vendor contracts should also address IP-ownership of AI-generated content such as images, text, and even new software. Regarding algorithmic training inputs, companies must decide whether the AI Agent's algorithm can train on company data, and if so, how such company data must be protected and stored. In addition, businesses should be wary of AI Agents violating third party terms of service which may prohibit access and use by bots. In non-negotiated use of Agentic AI, such as when a company deploys an off-the-shelf AI Agent from one of the major providers, companies should similarly review the provisions outlined above and weigh the ultimate risk of deploying this novel technology.

- **AI Privacy and Cybersecurity Considerations.** Using an AI Agent may result in the processing of enormous amounts of personal information as defined by state, federal, and international data privacy laws. For example, AI Agents designed to read and prepare automated responses to emails may digest any personal information contained in a user's inbox, while an AI Agent designed to make investment decisions will process sensitive financial information about the user. By collecting such large amounts of personal information, companies may grow their cybersecurity attack target, as bad actors are typically attracted to companies known to collect large amounts of detailed and/or sensitive personal information. Furthermore, by deploying an AI Agent, a company in scope for comprehensive state privacy laws, such as the California Consumer Privacy Act, may be required to offer consumers the opportunity to opt-out of the AI-Agent's automated decision making, while also disclosing to its consumers how their personal information is collected and used by the AI Agent, including whether their personal information subsequently trains the AI Agent's algorithm. AI Agents that are not adequately directed and supervised could also perpetrate scams, develop vulnerable software code, or cause cybersecurity incidents, so human supervision and real-time monitoring will be essential to reduce legal risk, especially with initial uses of AI Agents.

With the launch of AI Agents, concerns about AI scheming have become immediate. Adding AI Agents to your teams may have serious, unintended consequences and should involve significant testing and implementation of controls prior to and during use. Businesses will soon be training AI Agents in their corporate values and adding AI Agents to their teams. The way we work will never be the same.