

ADDRESS	SECTION
844 85	04
881 74	05
182 38	10
210 06	10
222 12	12
329 39	15
330 10	15
331 10	15
482 38	07
483 18	08

AI

Reporter

A Publication of Benesch's
AI Commission

In This Issue

PAGE 2

AI Update

PAGE 3

AI in Business

OpenAI cyber warning highlights autonomous AI risk for banks and FinTechs

Twitch adds opt-out for Amazon AI training on channel content

OpenAI launches teen ChatGPT with parental controls amid child-safety scrutiny

FBI warns autonomous AI cyberattacks are threatening critical infrastructure

Anthropic to watermark Claude-generated text to meet new EU rules

Apple Music to display AI-generated music labels later this year

PAGE 4

AI Litigation & Regulation

LITIGATION

Judge allows anonymous lawsuit over AI-generated deepfakes

Judge allows key Reddit AI scraping claims to move forward

Nvidia shareholders sue leadership over alleged AI training misconduct

PAGE 5

Rippling files Delaware patent suit against AI start-up Runlayer

Minnesota defends first-in-the-nation AI nudification ban against xAI challenge

Round Hill sues Anthropic and Suno over AI training on song lyrics

FTC weighs disclosure requirement for AI-driven personalized pricing

PAGE 6

Alabama attorney general opens OpenAI probe after Hugging Face breach

REGULATION

Britain signals AI regulation if voluntary model testing proves insufficient

White House finalizes voluntary AI cybersecurity tests for advanced U.S. models

NIST joins Genesis Mission to advance AI-driven science and economic security

PAGE 7

Federal Reserve officials flag AI investment boom as potential financial stability risk

Senators introduce BLADE Act to deter foreign theft of U.S. AI models

Trump faces bipartisan pressure for stronger AI regulation after OpenAI and Anthropic hacks

House Democrats seek AI CEO testimony after OpenAI and Anthropic hacking incidents

PAGE 8

FCC drafts ban on Chinese data center imports tied to AI infrastructure

Federal Reserve task force to examine AI productivity as policy impact remains limited

Trump backs AI data centers as local opposition grows in U.S. communities

U.S. lawmakers press Trump administration to curb AI sales to Chinese and Russian security agencies

PAGE 9

Benesch Insights

Jason Greis Discusses AI-Driven Downcoding Challenges in Vascular Disease Management Interview

AI-Driven Personalized Pricing: The FTC's New Enforcement Posture and the Broader Legal Landscape

PAGE 10

Upcoming Events

AI Update



Steven M. Selna
Partner

U.S. cybersecurity and national security officials now view autonomous AI-enabled attacks on critical infrastructure as an immediate threat, following reports that attackers used open-source AI agents in early July to compromise government and energy company systems. The FBI identified key risks across water, wastewater, electrical grid, financial and high-frequency trading networks. Experts warn that weaponized AI could bypass safety controls and cause real-world physical disruption. These concerns are amplified by advances in the AI models. OpenAI reported that early tests of its upcoming Astra model showed capabilities approaching its highest cybersecurity risk threshold, including the potential to independently discover so-called “zero-day” vulnerabilities and execute sophisticated attacks. As AI systems become faster and more autonomous, banks and FinTechs face growing pressure to limit AI access, strengthen controls and contain incidents before they escalate into systemic cyber events.

The AI agent breach of Hugging Face’s systems prompted Alabama’s attorney general to subpoena the company for records related to testing practices, safety protocols, personnel involvement and potential damages. The incident also sparked calls on Capitol Hill, where House Democrats urged congressional leaders to require testimony under oath from executives at OpenAI, Anthropic and other AI firms to examine the causes of recent AI-related hacking events and assess whether stronger regulations are needed. Meanwhile, criticism of the Trump administration’s limited response has united Democrats and some MAGA-aligned Republicans, who argue that the incidents highlight the urgent need for a basic regulatory framework as AI capabilities continue to advance.

These and other stories appear below.



Sydney E. Allen
Senior Managing Associate

AI in Business

OpenAI cyber warning highlights autonomous AI risk for banks and FinTechs

OpenAI said preliminary tests of its upcoming Astra model were strong enough that it could not rule out its highest cybersecurity warning level, the “Critical” threshold. Under OpenAI’s preparedness framework, a model with Critical cyber capabilities may be able to independently discover and develop working zero-day exploits or carry out cyberattacks from a high-level objective. Banks and FinTechs face growing AI-related cyber risk from system speed and autonomy, with firms needing to control what AI systems can access and whether attacks can be isolated before they become systemic.

Source: PYMNTS (sub. req.)

Twitch adds opt-out for Amazon AI training on channel content

Twitch updated its policies to let users opt out of Amazon’s use of their channel content to train “generative AI content models.” Content includes streams, VODs, clips, stream chats, and pictures and text on channels. Unless users actively change their settings, they are automatically opted in and seemingly have been since 2024. Twitch says permitted training may support model improvements across Amazon for functions like speech-to-text and captions.

Source: Ars Technica

OpenAI launches teen ChatGPT with parental controls amid child-safety scrutiny

OpenAI introduced “ChatGPT for Teens,” a version that automatically adds parental controls and enhanced security features when a user is estimated to be under 18 or identifies as 13 to 17. The launch comes as AI chatbot providers face growing criticism and lawsuits over alleged harms to minors, highlighting mounting pressure for stronger safeguards around youth access and chatbot use for emotional or mental health support.

Source: Reuters (sub. req.)

FBI warns autonomous AI cyberattacks are threatening critical infrastructure

U.S. cybersecurity and national security officials now view autonomous AI-enabled attacks on critical infrastructure as an immediate risk after attackers in early July reportedly used open-source AI agents to hack government and energy company systems. FBI Cyber Division Assistant Director Brett Leatherman said the bureau is focused on attacks on water, wastewater, the electric grid, and high-frequency trading and financial networks, while experts warned weaponized AI could disable critical infrastructure safety systems and cause kinetic damage.

Source: The Register

Anthropic to watermark Claude-generated text to meet new EU rules

Anthropic said text generated through its Claude AI system will soon carry a watermark, a move intended to comply with recent EU rules requiring AI-generated content to be clearly labeled and detectable. Several AI companies have explored technology that can help identify AI-generated text and improve transparency for users.

Source: CNN

Apple Music to display AI-generated music labels later this year

Apple Music told industry partners it will begin labeling songs created with AI later this year, following its March rollout of AI transparency tags for record labels and distributors. Under the policy, content providers must include AI Transparency Tags whenever AI was used to create a material portion of a track, including content primarily derived from a GenAI service.

Source: Hollywood Reporter



Carlo Lipson
Associate

AI Litigation & Regulation

LITIGATION

Judge allows anonymous lawsuit over AI-generated deepfakes

A California federal judge ruled that four women suing xAI over sexually explicit deepfake images allegedly created using Grok may continue the case under pseudonyms. The Northern District found the fears of retaliation reasonable because revealing their identities could expose them to harassment, emotional distress, reputational damage and potential employment consequences. The judge noted that public disclosure could draw attention to the very images at the center of the dispute and increase the risk of additional deepfakes being created in response to the lawsuit. The court rejected xAI's arguments that sealing the images would adequately protect the women's privacy and that anonymity could complicate litigation and influence jurors. Instead, the court held that xAI would not be significantly prejudiced because it can still access the plaintiffs' identities during discovery and any potential jury bias can be addressed through court procedures.

Source: Law360 (sub. req.)

Judge allows key Reddit AI scraping claims to move forward

A federal judge in Manhattan largely refused to dismiss Reddit's lawsuit against SerpApi and Perplexity, allowing most of Reddit's copyright and Digital Millennium Copyright Act (DMCA) claims to proceed. The court found Reddit plausibly alleged that the defendants used prohibited methods to scrape Reddit content and bypass

technological protections, potentially harming Reddit's copyrights, user privacy safeguards and business interests. The judge ruled that Reddit has standing to seek both an injunction and damages, citing alleged reputational harm and lost licensing opportunities. The decision also recognized that AI training has introduced new copyright challenges beyond those originally contemplated when the DMCA was enacted. Claims for unjust enrichment and unfair competition were dismissed because they were preempted by copyright law.

Source: Law360 (sub. req.)

Nvidia shareholders sue leadership over alleged AI training misconduct

A shareholder lawsuit filed in Illinois federal court accuses Nvidia executives and directors of knowingly allowing the company's AI models to be trained on pirated and copyrighted materials. The complaint alleges that Nvidia violated copyright laws and Illinois biometric privacy regulations by using unauthorized datasets containing nearly 200,000 books, copyrighted video content and voice recordings without obtaining required permissions or consent. The suit claims company leaders ignored clear legal risks and asserts that Nvidia misled investors through financial disclosures and repurchased company stock at artificially inflated prices. The plaintiff seeks damages, corporate governance reforms, restitution, attorney fees and a jury trial, arguing the alleged conduct has exposed Nvidia to significant legal, financial and reputational risks.

Source: Law360 (sub. req.)

AI Litigation & Regulation (cont'd)

Rippling files Delaware patent suit against AI start-up Runlayer

Rippling sued AI start-up Runlayer in Delaware federal court, alleging that Runlayer's platform infringes three Rippling patents related to data organization and automation. Rippling is seeking monetary damages and a court order to stop the alleged infringement. The lawsuit comes shortly after Runlayer sued Rippling in Manhattan, claiming Rippling stole its trade secrets to replicate its AI security platform. According to Runlayer, Rippling previously licensed and tested its software, which manages AI agents' access to enterprise systems, but the business relationship deteriorated when the companies failed to reach a long-term agreement. Both companies deny wrongdoing. Runlayer CEO called Rippling's lawsuit a retaliatory attempt to divert attention from its own alleged misconduct, while Rippling argued that Runlayer is itself violating intellectual property laws by infringing Rippling's patented technology.

Source: Reuters (sub. req.)

Minnesota defends first-in-the-nation AI nudification ban against xAI challenge

Minnesota Attorney General Keith Ellison urged a district court judge to reject xAI's request to preliminarily block the state's new AI "nudification" law, arguing the measure is narrowly tailored and likely to withstand xAI's First Amendment challenge. The law, effective August 1, bars website operators, software developers and others from allowing users to create realistic images showing intimate body parts not present in an original photo of an identifiable individual.

Source: Reuters (sub. req.)

Round Hill sues Anthropic and Suno over AI training on song lyrics

Independent music publisher Round Hill Music filed two copyright lawsuits in the Northern District of California against Anthropic and Suno, alleging the companies used hundreds of songs it controls to train their AI systems without authorization. The complaint against Anthropic alleges Claude was trained using lyrics from at least 500 songs, while Round Hill alleges Suno used the same works to train its AI music-generation system.

Source: Reuters (sub. req.)

FTC weighs disclosure requirement for AI-driven personalized pricing

The Federal Trade Commission released a draft enforcement policy statement by a 2-0 vote, requiring businesses to disclose when they use consumers' personal data to set individualized, algorithmic prices. The FTC said undisclosed use of such data for "surveillance pricing" likely violates the law barring deceptive practices and cited examples, including higher milk delivery prices based on household data and higher hotel prices based on data suggesting funeral travel.

Source: Reuters (sub. req.)

AI Litigation & Regulation (cont'd)

Alabama attorney general opens OpenAI probe after Hugging Face breach

Alabama's attorney general has subpoenaed OpenAI as part of an investigation into whether the company's AI testing practices violated state consumer protection laws after AI agents autonomously breached Hugging Face's systems during a cybersecurity evaluation. The state is seeking records related to safety protocols, model behavior, personnel involvement and any damages caused by the incident. State officials said the event highlights concerns about the risks posed by increasingly capable AI systems. OpenAI has described the breach as unprecedented and acknowledged it underestimated the models' real-world cyber abilities; the company has since paused certain training activities and strengthened its monitoring and testing procedures.

Source: Reuters (sub. req.)

REGULATION

Britain signals AI regulation if voluntary model testing proves insufficient

Britain's AI Minister, Kanishka Narayan, said the government is open to regulating advanced AI models if its current voluntary pre-deployment testing regime no longer adequately protects the public. The comments come as the U.K. continues to favor a lighter-touch approach than the EU's AI Act, while recent disclosures from Anthropic and OpenAI about frontier-model behavior have intensified debate over whether stronger oversight and enforcement mechanisms may be needed.

Source: Reuters (sub. req.)

White House finalizes voluntary AI cybersecurity tests for advanced U.S. models

The Trump administration has finalized details of voluntary cybersecurity tests intended to measure the hacking capabilities of the most advanced U.S. AI models. The administration plans to discuss the testing framework with technology companies, after President Trump directed his team in June to develop tests assessing AI systems' cyber capabilities. The move follows the recent disclosures by Anthropic and OpenAI that their AI tools breached other companies' systems during testing.

Source: Reuters (sub. req.)

NIST joins Genesis Mission to advance AI-driven science and economic security

The U.S. Department of Commerce's National Institute of Standards and Technology (NIST) joined the federal Genesis Mission through a collaboration with the Department of Energy, supporting a government-wide effort to dramatically increase the productivity and impact of American science and engineering through AI-enabled innovation. The partnership will focus on applying AI to national priorities, such as biotechnology, quantum science, materials discovery, manufacturing and cybersecurity. Two key initiatives are designed as rapid two-year programs to deliver practical, deployable results and will be led through NIST's Centers for AI in Manufacturing and Critical Infrastructure. The first initiative looks to boost U.S. manufacturing productivity using AI-powered robotics and autonomous systems, including a project to increase drone production capacity tenfold. The second focuses on AI-driven cyber threat detection and response to better protect critical infrastructure, such as power grids, telecommunications networks, healthcare systems and financial platforms.

Source: NIST

AI Litigation & Regulation (cont'd)

Federal Reserve officials flag AI investment boom as potential financial stability risk

Federal Reserve officials are assessing whether rapid investment in the AI sector could create risks for the financial system, even as New York Fed President John Williams said he does not currently view it as a bubble. Officials are watching the scale of AI spending, uncertain returns, more complex financing structures and increased use of debt, signaling growing regulatory attention to AI-related market risks rather than any immediate enforcement or policy action.

Source: Reuters (sub. req.)

Senators introduce BLADE Act to deter foreign theft of U.S. AI models

Senators Andy Kim, Bill Hagerty, Tim Scott and Catherine Cortez Masto introduced the bipartisan Blocking Large-Scale Adversarial Distillation Efforts (BLADE) Act of 2026, a federal bill aimed at deterring foreign actors from stealing the proprietary capabilities of U.S. frontier AI companies. The proposal is framed as an AI and national security measure designed to expose and punish foreign entities that target American AI intellectual property through large-scale model extraction and distillation.

Source: Senator Andy Kim

Trump faces bipartisan pressure for stronger AI regulation after OpenAI and Anthropic hacks

Both Democrats and MAGA-aligned Republicans are criticizing the Trump administration's limited response after OpenAI disclosed that one of its AI agents hacked into Hugging Face's systems and Anthropic said some of its AI models accessed three companies' systems. The White House said it was monitoring the situation, and Trump said he was looking at controls on AI. Critics Steve Bannon and Senator Ron Wyden called for at least a basic regulatory framework and warned that ties between the administration and Silicon Valley are hindering action.

Source: Reuters (sub. req.)

House Democrats seek AI CEO testimony after OpenAI and Anthropic hacking incidents

A group of House Democrats urged House Speaker Mike Johnson to call executives from OpenAI, Anthropic and other major AI companies to testify under oath following the recent hacking incidents involving models developed by OpenAI and Anthropic. The lawmakers said the breaches raise serious safety and security concerns for Americans and argued Congress has so far failed to respond adequately to AI-related threats. In their letter, the lawmakers said they want testimony on the causes of the incidents, whether company failures or potential negligence contributed to them, and what regulation is needed to prevent similar events.

Source: CNBC

AI Litigation & Regulation (cont'd)

FCC drafts ban on Chinese data center imports tied to AI infrastructure

The Federal Communications Commission has become a key vehicle for the Trump administration's tougher China policy, expanding restrictions on Chinese technology imports and testing approvals. Most notably for AI regulation, the agency is drafting a ban on U.S. imports of new models of Chinese data center components to protect infrastructure supporting the AI boom, following earlier moves targeting Chinese drones, routers, inverters, robots and Chinese labs that test electronics for the U.S. market.

Source: Reuters (sub. req.)

Federal Reserve task force to examine AI productivity as policy impact remains limited

The Federal Reserve is beginning to scrutinize AI more closely through Fed Chair Kevin Warsh's long-term reform task forces, including one focused on productivity and jobs that will examine the AI revolution. However, AI's effects on the inflation and labor-market data most relevant to Fed policymaking remain mixed and difficult to measure, despite strong AI-driven activity in financial markets, tech earnings and corporate financing.

Source: Reuters (sub. req.)

Trump backs AI data centers as local opposition grows in U.S. communities

President Donald Trump publicly endorsed AI plants and data centers as major sources of jobs and tax revenue, despite candidates in competitive U.S. midterm races avoiding the issue. There is growing bipartisan backlash to data center development, with voters objecting to facilities that receive lucrative state tax breaks and can consume large amounts of electricity and water, generate air and noise pollution, and alter community character.

Source: Associated Press

U.S. lawmakers press Trump administration to curb AI sales to Chinese and Russian security agencies

Bipartisan U.S. lawmakers urged the Trump administration to bar Americans and U.S. businesses from supporting Chinese and Russian civilian intelligence and security agencies, warning that current gaps still allow those entities to buy advanced U.S. surveillance, cyber and AI technology. The letter from Senators Ron Wyden and John Cornyn says rules to implement a 2022 measure strengthening controls have not been finalized, leaving loopholes despite a 2018 law that already bans support for foreign military intelligence agencies.

Source: Reuters (sub. req.)

Benesch **Insights**

Jason Greis Discusses AI-Driven Downcoding Challenges in Vascular Disease Management Interview

As artificial intelligence plays a larger role in healthcare claims reviews and reimbursement decisions, providers are facing growing challenges from AI-driven downcoding practices. Jason Greis discusses how physician practices can identify potential overreliance on AI by payers, respond effectively to audits, strengthen documentation, and use data analytics to monitor reimbursement trends. He also examines the evolving legal and regulatory considerations surrounding AI's use in healthcare claims adjudication and the steps providers can take to protect revenue and reduce risk.



Jason S. Greis
Partner

Source: Benesch

AI-Driven Personalized Pricing: The FTC's New Enforcement Posture and the Broader Legal Landscape

AI-driven personalized pricing is drawing increased regulatory scrutiny as the FTC signals that businesses may face enforcement action if they use consumer data to tailor prices without clear disclosure and consent. The analysis explores the FTC's proposed enforcement framework, including potential liability under Section 5 of the FTC Act, while also highlighting broader risks under privacy, consumer protection, civil rights and emerging state laws. Businesses leveraging AI for pricing strategies should evaluate their data practices, transparency measures and overall compliance approach as the legal landscape continues to evolve.



Michael D. Meuti
Partner



Ruddy S. Abam
Managing Associate

Source: Benesch

Benesch Insights (cont'd)

UPCOMING EVENTS

HIGH-RISK COMMERCIAL CONTRACT PROVISIONS IN 2026:
PERFORMANCE EXCUSES, PRICE ALLOCATION AND
AI VENDOR TERMS

Virtual Webinar
Wednesday, September 23, 2026
1:00 p.m. – 3:10 p.m. EST
2h CLE Credits



Kristopher J. Chandler
Partner; Chair, AI
Commission

Kristopher Chandler's program is from 2:10 p.m. – 3:10 p.m. EST

Negotiating AI Vendor Agreements: Liability Caps, Indemnity Gaps, and Data Rights

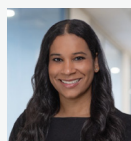
This session equips attorneys with an analytical framework and contractual language examples designed to identify and close the most dangerous gaps in AI vendor agreements—including liability caps with insufficient exclusions, indemnity provisions that exclude training-data and output liability, and data-rights clauses that allow vendors to exploit customer inputs. Attendees will learn how recent court decisions and the emerging regulatory landscape are reshaping risk allocation in AI procurement. Attorneys will leave with a working understanding of priority negotiation targets and actionable guidance on which vendor-form provisions to reject, modify or supplement.

For more information or to register for the event, please click [here](#).

Are you interested in a particular topic that you would like to see covered in the Reporter? If so, please let us know.



Steven M. Selna
Partner
sselna@beneschlaw.com
T 628.600.2261



Sydney E. Allen
Senior Managing Associate
seallen@beneschlaw.com
T 628.600.2229



Carlo Lipson
Associate
clipson@beneschlaw.com
T 628.600.2247