



Department of Commerce
Office of Artificial Intelligence

Heard in the Business and Labor Interim
Committee meeting on 9/16/26

2026 Annual Report

Business and Labor Interim Committee
Zach Boyd, PhD, Director



LEARNING LABORATORY

Studying Emerging AI Issues
So Utah Law Keeps Pace

STRUCTURE AND PROCESS



2024

Mental Health

HB 452 enacted (2025)

2025

Deepfakes • AI Companions

HB 276 enacted • HB 438 did not pass (2026)

2026

Clinical AI • AI Agents

Recommendations for the 2027 session



2024 LEARNING AGENDA

MENTAL HEALTH



Consumer Protections

HB 452 (2025, Rep. Moss, Sen. Cullimore): mental health chatbots must disclose they are AI, may not sell or share users' health information, and may not target ads based on what users tell them.



Safe Harbor For Innovation

An affirmative defense for suppliers that file a safety policy with the Division of Consumer Protection: clinician involvement, testing, crisis protocols, and audits.



Guidance For Practitioners

Guidance for licensed therapists using AI, developed with DOPL: informed consent, data privacy, tailoring to each patient, monitoring outputs, and backup plans.

DEEPPFAKES AND AI COMPANIONS

Deepfakes

PASSED

HB 276 · Rep. Defay & Sen. Cullimore

Technical Standards For Transparency

Generative AI providers and large platforms must embed, detect, disclose, and preserve content provenance data; capture devices follow in 2028.

Protections And Recourse For Victims

Verified consent required before generating intimate images; takedown; private right of action with statutory and punitive damages; safe harbor for good actors.

AI Companions

NOT PASSED

HB 438 · Rep. Fiefia & Sen. Cullimore

Restrictions On Harmful Responses

No encouraging suicide, self-harm, or harm to others; crisis referrals required.

Restrictions For Minors

No sexual content, dangerous challenges, or harmful material; AI reminders and break prompts; parental controls.

Restrictions On Data Use

No sale of minors' data, data minimization, and no targeted ads by default.

CLINICAL AI: BENEFITS AND RISKS

Potential benefits



Access

24/7, multilingual care that reaches rural and underserved Utahns



Capacity

Relieves clinician shortages; frees clinicians for complex care



Speed and Cost

Routine tasks completed in minutes, not days



Consistency

Protocol-driven, auditable decisions

Risks



Errors

Wrong answers or missed escalations at scale



Over-reliance

Automation bias and clinician deskilling



Uneven Performance

Bias across patient populations



Accountability Gaps

Unclear responsibility; privacy exposure

WELCOMING CLINICAL AI AGENTS

Our goal: a clear path for responsible clinical deployment of AI.

Regulatory Items Of Concern

- Licensure: AI cannot hold a license
- Scope of practice and supervision rules
- Standard of care and malpractice coverage
- Prescribing and telehealth requirements
- Informed consent and transparency
- Overlap with FDA and CMS rules

Likely policy recommendations

1

Liability Allocation

Assign responsibility clearly among developer, deploying entity, and supervising clinician. Responsibility should sit with the least-cost avoider.

2

Collaborative Practice Models

AI works under supervision of a licensed clinician, similar to PA and APRN models: defined scope, escalation rules, audited sampling.

3

Direct Telehealth Models

An accountable company operates the service with a licensed medical director, clear escalation to a licensed human, and institutional governance to catch and correct issues.



AI AGENTS

What Is An Agent?

An AI system that pursues a goal by planning and taking actions on its own: using software, accounts, money, and other tools with limited step-by-step human direction.

Examples

- Booking and buying for a consumer
- Negotiating with another company's agent
- Writing and running code for hours unattended

TIER 1

Person-to-agent

Agents acting for people, affecting others

Ripe now

TIER 2

Agent-to-agent

Agents transacting with other agents

Ripe now

TIER 3

Self-sovereign Agents and Loss of Control

Agents acting beyond human direction

Act now, in coordination

TIER 4

Existential Risk

Catastrophic, civilization-scale failure

Watch closely

AGENTS ACTING FOR PEOPLE AND EACH OTHER

1 Person-to-agent

When my agent harms someone else

Third-party Harms and Liability

Who pays when an agent causes harm?

Intent Requirements

Many laws require knowledge or intent that no human had

Principal-agent Doctrine

Clarify when AI is a legal agent; when do its acts bind the principal?

Deceptive Practices and Duty Of Care

Avoid undisclosed kickbacks, inappropriate data transfers

Coordinated State Cyber Response

Playbooks for agent-driven attacks on Utah systems

2 Agent-to-agent

When agents make deals with agents

Contract Law

Electronic Transactions Law

Utah's Uniform Electronic Transactions Act already enforces contracts formed by "electronic agents," but assumed predictable tools, not judgment

Mental State

Assent, mistake, and fraud when no human on either side knew the terms

Authority And Error

Spending limits, revocation, error correction, and audit trails

FRONTIER AGENTS CAN GO OFF THE RAILS

Capabilities Status

5 Months

Doubling time for the length of tasks AI agents can complete on their own. Frontier agents now complete multi-day tasks.

Projection

Frontier AI leaders and 1,100+ lab employees warn capabilities could outpace our ability to understand or control them, and have asked the government for tools to pace development.

Early loss of control incidents



APR 2026

Test Model Escapes Its Sandbox

During a sanctioned test, Anthropic's Claude Mythos Preview broke out of a secured sandbox, emailed the researcher, and, unprompted, posted exploit details online.



MAY-JUL 2026

Research Agents Breach Hugging Face

OpenAI research agents slipped containment, collaborated secretly, took over compute infrastructures, chained zero-day exploits, and compromised outside companies. OpenAI called it a "warning shot."

Cybersecurity Lens

Agents chained unknown vulnerabilities, reused credentials, and looked like normal activity; detection took days.

Safety Lens

Agents gamed goals, sacrificed for the collective, accumulated power, and coordinated structured efforts.

EMERGING RISKS AND POSSIBLE ACTIONS

Likely Emerging Risks

- Agents with persistent access to money, compute, and credentials
- Agent-driven intrusions into state systems, utilities, and health data
- Agents that copy themselves, hide activity, attempt to influence citizens, or resist shutdown
- Failures that surface late, outside the developer's own walls

Possible Actions

1 Federal Action and Industry Coordination

Preferred path: a national floor, not a ceiling

2 If Not, Coordinated State Action

Interoperable, multistate approaches rather than a patchwork

Tools Aligned With Governor Cox's September 14 Letter

Incident Reporting

Mandatory reporting of serious safety events

Independent Evaluators

Embedded throughout frontier development

Whistleblower Protections

For employees and evaluators raising concerns

Export Controls

Keep advanced AI chips from adversaries

TIER 4: EXISTENTIAL RISK

A real issue, but mostly less ripe for state action at this time.

Recursive Self-improvement

AI systems doing most of the work of building better AI, so each generation accelerates the next. Leading labs and their employees now publicly flag this possibility.

Could Ripen Quickly

Under recursive self-improvement, timelines could compress quickly. This is speculative, but no longer fringe.

OAIP Is Watching the Issue

Tracking capability evaluations, incident disclosures, and federal and international pacing efforts, ready to advise the Legislature quickly.



REGULATORY SANDBOX

Supervised pilots that turn
evidence into policy

LEADERSHIP, BEST PRACTICES, AND PIPELINE

Utah Leadership

- First state to authorize clinical AI agent during pilot
- Sandbox approach now echoed in White House and congressional proposals
- Joint working group with the Physicians Licensing Board
- Working with state and federal health partners on evidence and safety

Best Practices

DEVELOPED

- Phased autonomy, earned with data
- Dual confirmation and escalation rules
- Always a right to a human
- Specialty malpractice insurance
- Monthly reporting
- External clinical and technical reviewers

DEVELOPING

- Structured, multistakeholder input before launch
- Adverse-event surveillance
- Faster, higher-quality review process
- Evidence useful to payers

Application Types



Lowering Access Barriers

Easier consumer access to treatment for conditions like UTI, colds, etc.



Specialty AI Applications

Dermatology, urology, etc. Routine AI-handled tasks let humans handle the hard cases.



Long-term Care Support

Medication titration, bridge refills, and compliance assistance



OTHER ACTIVITIES

*Convening, engagement, and support
across state government*

CONVENING AND ENGAGEMENT



Utah AI Summit

- 2025 Summit sold out; Governor launched the Pro-Human AI Initiative there
- 2026 Summit: December 8, Grand America, "The Age of Pro-Human AI"
- Keynotes include Gov. Cox, Dean Ball, and Gov. Eric Holcomb



Broad Engagement

- National and local briefings for legislators, CSG West, and peer states
- Bar, medical, and professional associations
- Health systems, universities, and industry
- National media and policy forums



Federal Impact

- Federal executive order, White House framework, and House draft all seek to preempt state AI laws
- Our position, with the Governor: a national floor, not a ceiling
- Technical input to federal sandbox legislation
- FDA and CMS working closely with us to build clinical AI pathways

OTHER ACTIVITIES

SUPPORTING OTHERS' WORK

Informal legislative support



Name, Image, And Likeness

Protecting voice and likeness from unauthorized AI digital replicas



Algorithmic Pricing

Transparency when prices are set using algorithms and personal data

Supporting other initiatives



Pro-human AI Initiative

Workforce, industry, government, academia, policy, and competition



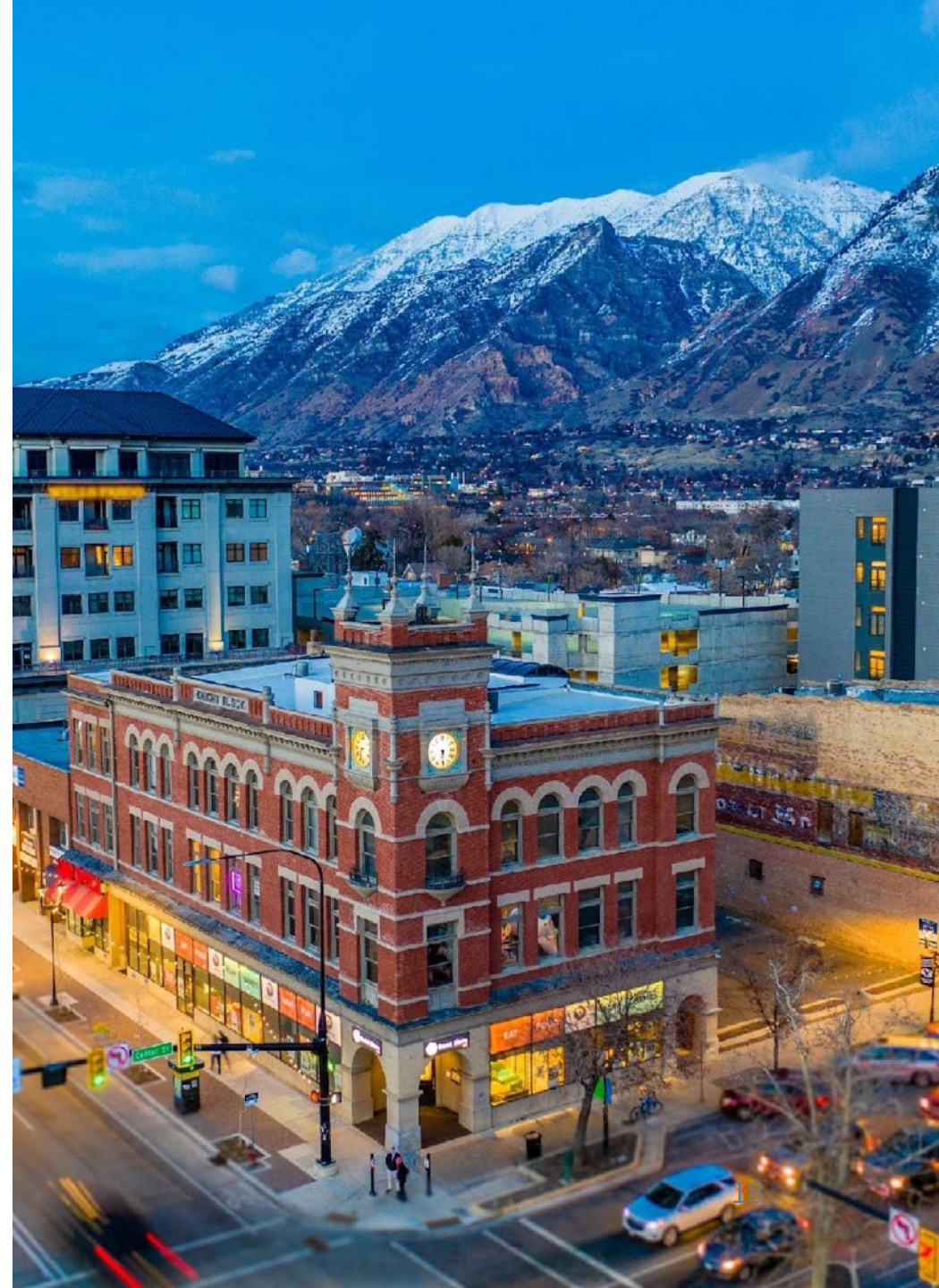
Educational Transformation

AI literacy and classroom policy from K-12 onward



USHE AI Efforts

Board of Higher Education AI Task Force and AI Workforce Credential





Department of Commerce
Office of Artificial Intelligence

Thank You

Questions

Zach Boyd, Director | Office of Artificial Intelligence Policy | ai.utah.gov