

Use of AI Agents and Agentic Systems Standard

Introduction

Oklahoma's artificial intelligence strategy aims to establish the state as a leader in the responsible, safe, secure and proactive use of artificial intelligence. AI Agents and Agentic Systems introduce capabilities beyond traditional AI tools by enabling software to plan, reason, select tools, interact with systems and take actions in pursuit of defined objectives. These capabilities may improve government efficiency, service delivery and employee productivity, but they also create elevated risks related to delegated authority, cybersecurity, privacy, accountability, auditability and public trust. While the [Use of AI in Oklahoma State Government Standard](#) governs data and models through a point-in-time approval process, AI Agents and Agentic Systems require ongoing governance of their decisions and actions through continuous oversight. This standard establishes that governance framework.

Purpose

This standard establishes clear guidelines for the governance and responsible use of AI Agents and Agentic Systems within Oklahoma state government. It ensures that AI Agents are designed, acquired, deployed and monitored in a manner that protects state data and systems while governing the authority, data access, tool access, level of autonomy and scope of actions delegated to them to promote security, accountability and public trust. Relationship to the Use of AI in Oklahoma State Government Standard.

This standard supplements the [Use of AI in Oklahoma State Government Standard](#) by establishing additional governance requirements for the design, acquisition, approval, deployment, monitoring and use of AI Agents and Agentic Systems. While the [Use of AI in Oklahoma State Government Standard](#) applies broadly to the use of AI across Oklahoma state government, this standard addresses the unique risks and governance considerations associated with AI Agents and Agentic Systems. Oklahoma state agencies using AI Agents or Agentic Systems must comply with both standards, with this standard providing the additional requirements applicable to this specific use case.

Definitions

Unless otherwise defined in this standard, all terms shall have the meanings assigned to them in the [Use of AI in Oklahoma State Government Standard](#). The definitions in that standard are incorporated by reference and apply to this standard unless expressly modified or supplemented herein.

AI Agent - an AI-enabled technology that can autonomously or semi-autonomously carry out tasks, make decisions, interact with applications or services and adapt its behavior in pursuit of assigned objectives, with limited human guidance through perceiving information, planning tasks, selecting tools, interacting with applications or services, retrieving data, generating outputs, or executing actions to achieve an assigned objective.

Agentic System - one or more AI Agents operating within software, workflows, application programming interfaces ("APIs"), robotic process automation, enterprise applications, data

sources or other digital tools to carry out business processes, operational workflows or other objectives defined by humans.

Action Space - the specific set of tasks, tools, systems, data sources, applications, decisions, transactions, permissions or workflow steps that an AI Agent is authorized to access or perform.

Tool Use - an AI Agent's ability to invoke an external application, API, database, file repository, workflow, browser, code execution environment, communication tool or other system capability.

Automation Drift - the gradual shift of an AI Agent's behavior, usage, reliance or scope beyond its original approved purpose, controls or operating boundaries.

Agent Hijacking - an attack or failure mode in which an AI Agent is manipulated, including through prompt injection or malicious content in emails, documents, websites, files or other data sources, causing the AI Agent to take unintended or unauthorized actions.

Agent Sponsor - the accountable business owner of record for an AI Agent, as identified in the Agent Registry, who is responsible for the AI Agent's purpose, delegated authorities, operating boundaries, lifecycle, and ongoing oversight. The Agent Sponsor shall monitor the AI Agent's operation and ensure its continued compliance with applicable laws, regulations, and this standard.

Irreversible Action - an action whose effects cannot be fully reversed through a documented rollback or remediation procedure within a defined recovery period.

Untrusted Content - content originating outside state control that an AI system ingests or processes.

Agent Registry - the statewide registry of AI Agents established and maintained by the CIO pursuant to this standard.

Sponsoring Organization – the state agency, board, commission, office, shared service provider, or OMES-designated business or technology function that owns, funds, requests, operates, or is accountable for an AI Agent or Agentic System. The Sponsoring Organization is responsible for ensuring the AI Agent is registered, monitored, supported, reviewed and operated within its approved purpose, autonomy level, action space and control requirements.

Agent Sponsor – a designated state employee, official or role accountable for the business purpose, approved use, operational need, and ongoing accountability of an AI Agent or Agentic System.

Authorized Agent Action Approver - an authorized state employee, official, or approved designee with delegated authority to approve the action authority envelope for an AI Agent or Agentic System. The action authority envelope shall define the agent's approved autonomy level, permitted action types, business rules, thresholds, systems, data access, escalation conditions and actions requiring additional approval before execution. The Authorized Agent Action Approver is not required to approve every individual agent action unless required by law, policy, risk classification or the approved operating boundaries of the agent.

Standard

General requirements

All AI Agents and Agentic Systems requested, developed, procured, piloted, deployed or used by Oklahoma state agencies shall be subject to CIO authority and shall comply with this Standard and

[Use of AI in Oklahoma State Government Standard](#), including all applicable procurement, security, privacy, governance, training, incident reporting and auditing requirements.

Approval of an AI Agent or Agentic System is limited to the documented business purpose, level of autonomy, data access, tool access, action space, and operating boundaries that were reviewed and approved by the CIO. Such approval shall not be construed as blanket approval for all uses of that AI Agent or Agentic System. Any material change to the approved use or capabilities requires a new review and approval before implementation.

Agent autonomy classification

Prior to deployment, each AI Agent or Agentic System shall be classified by the CIO according to the highest level of autonomy it is permitted to exercise. The state agency shall submit a [PIR – AI Review](#) with sufficient documentation for the CIO to evaluate and classify the AI Agent and Agentic System. Classification shall be based on the AI Agent or Agentic System’s demonstrated and configured capabilities and approved action space and shall not be determined by the supplier or supplier characterization.

Level of Autonomy	Classification	Capabilities	Required Approval
0	Observe	Searches, summarizes, retrieves, classifies, or analyzes information without taking action. Provides information only. Output does not change any system.	Agency business owner
1	Recommend	Provides recommendations, drafts content, prepares analysis or proposes next steps for human review. No changes occur without human action.	Agency business owner
2	Prepare	Prepares workflow actions, records, forms, code, messages or transactions, but requires human approval before execution. Examples include but are not limited to draft emails, software logic, service tickets, reports, etc.	Agency business owner and technical owner
3	Execute	Performs approved actions within predefined rules, including but not limited to creating tickets, updating approved records, routing cases, scheduling meetings or sending approved communications with limited human interaction.	Agency business owner, technical owner, Agency executive leadership owner, and the CIO.

Level of Autonomy	Classification	Capabilities	Required Approval
4	Orchestrate	Coordinates multi-step actions across multiple systems, APIs, workflows, data repositories, or business processes with limited human interaction.	Agency business owner, technical owner, Agency executive leadership owner, and the CIO.
5	Delegate	Creates, manages, supervises, invokes, or delegates tasks to other AI Agents or multi-agent systems with limited to no human interaction.	Agency business owner, technical owner, Agency executive leadership owner, and the CIO.

AI Agents and Agentic Systems that are classified as an autonomy level of three (3) or higher under the framework established by this standard require documented threat modeling that references the OWASP Top 10 for Agentic Applications and MITRE ATLAS, informed by the NIST AI Risk Management Framework.

Accountability

Use of an AI Agent or Agentic System does not transfer responsibility from a state agency, employee, contractor, supplier or public official to the AI Agent or Agentic System. The state agency using or sponsoring the AI Agent shall retain full responsibility and accountability for all actions of the AI Agent, including outputs, decisions, communications, records, workflows, and operational impacts from or support by the AI Agent. The CIO review, approval, platform approval, security review, or vendor certification constitutes only a review of the proposed AI acquisition or use and shall not be interpreted as assuming responsibility for the AI Agent or its actions. The state agency implementing the AI Agent is solely responsible and accountable for the AI Agent.

Shared Accountability

Use of an AI Agent or Agentic System does not transfer responsibility from a state agency, employee, contractor, supplier or public official to the AI Agent or Agentic System. The agency using or sponsoring the AI Agent remains accountable for the business process, public service, decision support, communication, record, workflow, or operational impact supported by the AI Agent.

State IT review, platform approval, security review or vendor certification does not relieve the agency of responsibility for appropriate use, human oversight, compliance with applicable law and policy or management of operational impacts within the agency's business process.

Minimum Required Accountabilities & Role(s)

Before deploying or materially modifying an AI Agent or Agentic System, a state agency shall identify and document the individuals or roles responsible for the following functions, as applicable based on the AI Agent's autonomy level, data sensitivity, system access and operational risk.

- **Business owner:** an individual, team, or organization responsible for the business purpose, use case, operational fit, human oversight, adoption, value realization and accountable use of the AI Agent(s).
- **Technical owner:** an individual, team, or approver service provider responsible for technical configuration, integration, security implementation, environment, logging, monitoring, access implementation, support, incident coordination and technical change management.
- **Agency executive leadership owner:** an individual or team responsible for approving higher-autonomy or higher-risk AI Agents and accepting agency level operational accountability.
- **Data owner:** an individual, team, or organization with authority that is accountable for data sensitivity, data quality, data access, retention, privacy, and appropriate use.
- **Monitoring owner:** an individual, team, or organization responsible for reviewing logs, alerts, exceptions, usage, cost, performance and other runtime evidence.
- **Approver:** the person or body authorized to approve deployment, material changes, exceptions or continued use.
- **Intervention owner:** the person or team authorized to pause, disable, rollback, contain or escalate the AI Agent when necessary.

The first three roles above, business owner, technical owner and agency executive leadership owner shall be identified for all AI Agents and Agentic Systems, regardless of autonomy level. State agencies may assign additional responsibilities, establish more restrictive internal procedures or require additional approvals based on mission, risk, data sensitivity, system access, public impact or operational need. However, agencies shall not eliminate, delegate, bypass or replace the three required accountability roles.

Statewide AI Agent Registry

The CIO will maintain a centralized Agent Registry of all AI Agents and Agentic Systems operating across state agencies that are classified as an autonomy level of three (3) or higher using the framework established in this standard. State agencies shall provide OMES with the following information for each applicable AI Agent or Agentic System: agent name and supplier or platform; Agent Sponsor and technical owner; autonomy level; approved action space; systems and data sources accessed; authorized data classifications; deployment date; and current operational status (active, paused, deactivated). State agencies are responsible for ensuring the information provided to OMES remains accurate and up to date and shall notify OMES of any changes to the information maintained in the Agent Registry within five (5) business days of the change. The information provided by state agencies shall be maintained in the Agent Registry. The Agent Registry will be used for auditing, enterprise oversight, and incident response. No AI Agent classified as an autonomy level of three (3) or higher using the framework established in this standard shall be deployed without a corresponding registry entry.

AI Agents classified as an autonomy level of three (3) or higher using the framework established in this standard in operation before the effective date of this Standard shall be registered, placed under approved identity and access controls, and assigned an Agent Sponsor within ninety (90) days of the effective date of this standard. AI Agents that are not brought into compliance within the required timeframe shall be suspended from operation until all applicable requirements have been satisfied. No grandfathering or exemption from the requirements of this standard shall be permitted based solely on an AI Agency's prior deployment or operational status.

Required pre-deployment review

Before any AI Agent or Agentic System is deployed, the state agency shall define in writing the:

- business purpose and expected value
- assigned Agent Sponsor, technical owner, and data owner where applicable
- autonomy level
- approved action space
- systems, APIs, applications, and data sources the AI Agent may access
- actions the AI Agent may perform
- actions requiring human approval
- sensitive, confidential, regulated, or federally protected data involved
- reversibility classification of actions within the action space
- logging and audit requirements
- monitoring approach
- escalation path
- pause, disable, or deactivation method
- known risks, including cybersecurity, privacy, legal, operational, and public impact risks

For AI Agents procured from a 3rd party supplier, the supplier shall supply the state agency with an agent system card that documents the AI Agent's capabilities, limitations, testing performed, and known failure modes, which shall be incorporated into the review record, which includes but not limited to relevant documentation package retained for an AI Agent or Agentic System that evidences the intake, risk classification, review materials, approvals, approved action space, conditions of use, monitoring requirements, testing, system card, remediation actions, re-review, and audit history. This review record shall be included in the contract between the supplier and state agency.

Reversibility and recovery

Actions within an AI Agent's approved action space shall be classified as reversible, conditionally reversible, or irreversible. For any action space including irreversible actions, agencies shall implement, as applicable: pre-action confirmation or preview steps; versioning or restore points for records or systems the AI Agent may modify; and a documented rollback procedure sufficient to undo or remediate the AI Agent's actions within a defined recovery window.

AI Agents shall not be authorized to perform irreversible actions unless a documented and tested rollback or remediation capability is in place. This requirement does not apply when human approval is required for each individual action in accordance with the human accountability requirements established in this Standard.

Human accountability and oversight

Human accountability shall apply to all AI Agents and Agentic Systems. A designated Agent Sponsor shall remain accountable for the agent's approved business purpose, operating boundaries, appropriate use, and continued need.

Human approval shall be required where required by law, policy, risk classification, autonomy level or the approved action authority envelope. Human approval may apply to deployment, material changes, expanded access, higher-risk actions, exceptions or specific actions before execution. Human accountability shall remain with state agency staff regardless of the actions recommended, prepared or performed by an AI Agent. If accountable ownership becomes unclear because of role changes, vacancy, transfer or separation, the AI Agent shall be suspended or deactivated until sponsorship is formally restored.

Approval of an AI Agent's action authority envelope does not require human approval of every individual agent action. Agents may perform pre-approved actions without individual approval only when those actions remain within the approved autonomy level, action space, business rules, monitoring controls, and escalation thresholds.

Human approval shall be required before an AI Agent may:

- make or materially influence a legally binding decision
- affect eligibility for benefits, services, enforcement, licensing, employment, or public safety
- commit state funds or resources
- approve payments or financial transactions
- modify official government records
- send external communications on behalf of the State
- initiate procurement activity
- create, modify, or disable user accounts
- deploy code or configuration changes to production
- access or transmit sensitive or federally protected data beyond an approved use

Identity and access management

AI Agents shall operate only through approved identity and access management processes. AI Agents shall:

- use approved, standards-based service or runtime identities (i.e., including but not limited to OAuth/OpenID Connect or SPIFFE-compatible where technically feasible)
- be uniquely distinguishable in logs and system records
- operate under least privilege with just-in-time, task-scoped credentials rather than standing permissions where technically feasible
- receive only the permissions necessary for the approved action space
- be traceable to the Agent Sponsor and Authorized Agent Action Approver user where applicable
- have access reviewed periodically
- support timely revocation, suspension, or session termination

AI Agents shall not:

- use shared accounts
- use a state employee's credentials as a substitute for approved agent identity controls
- bypass identity governance, access review, change management, or cybersecurity controls
- retain standing access to sensitive systems unless specifically approved

AI Agents and Agentic Systems shall be reviewed on a risk-tiered basis to maintain overall health, continued business need, appropriate access, approved action space, monitoring effectiveness and compliance with this standard. AI Agents that are classified as an autonomy level two (2) or

below using the framework established by this standard may be reviewed annually by the CIO. AI Agents that are classified as an autonomy level of three (3) or above using the framework established by this standard may be reviewed by the CIO quarterly, unless a different cadence is required by the CIO based on documented risk.

AI Agents and Agentic Systems shall also be reviewed before implementation of any material change, including changes to autonomy level, action space, system integration, data access, model, vendor, business purpose, monitoring approach, or human oversight model. Additional review shall occur following significant incidents, unauthorized actions, automation drift, security events, or other conditions that materially affect risk. Periodic review shall not replace continuous monitoring where continuous monitoring is required by this standard.

Tooling, API, and system access

AI Agents shall only access tools, APIs, applications, databases, repositories, and workflows that have been explicitly approved. Access authorizations shall be documented and approved by the state agency before use. Permissions shall be limited to the minimum scope necessary to perform the intended function, configured using deny-by-default principles where technically feasible, monitored for unauthorized or unexpected behavior, and reviewed periodically to ensure continued compliance

AI Agents shall not be granted broad or unrestricted access to enterprise systems, sensitive data repositories, production environments, financial systems, identity systems, law enforcement systems, health systems or any other high-impact systems without written approval by the CIO.

Inter-agent communication shall be authenticated and encrypted, and agent-to-agent interactions shall be logged. For AI Agents or Agentic Systems that are classified as an autonomy level of five (5) using the framework established in this standard, the full delegation chain) shall be documented and approved as a condition of authorization. The delegation chain shall identify which agents may invoke other agents, the authority granted to each agent, and the permitted scope of delegated actions. Undocumented or unauthorized delegation between AI Agents is prohibited.

Untrusted content and prompt injection

AI Agents and Agentic Systems that ingest or process untrusted content shall treat such content as untrusted data not as instructions. Approval documentation for these AI Agents shall identify the prompt-injection and content-manipulation safeguards implemented. Monitoring shall include detection of injection attempts and deviation from the intended objectives. Content-derived instructions shall not expand an AI Agent's authorized action space or override established controls and boundaries.

Data protection

AI Agents shall comply with all applicable state and federal data protection requirements, including requirements governing the collection, access, use, storage, disclosure, and protection of personally identifiable information (PII), protected health information (PHI), Health Insurance Portability and Accountability (HIPAA) data, Family Educational Rights and Privacy Act (FERPA) data, federal tax information (FTI), criminal justice information (CJI), financial information,

credentials, confidential business information and other sensitive or legally protected data. Agencies shall ensure:

- data minimization
- approved data sources
- encryption in transit and at rest where applicable
- restrictions on data retention and reuse
- protection against unauthorized disclosure
- compliance with applicable records retention and litigation hold requirement
- clear handling rules for agent memory, stored context, embeddings, logs, and generated outputs

AI Agent memory shall be logically partitioned by agent and use case, protected against unauthorized modification, and subject to integrity verification controls, including cryptographic verification where technically feasible. AI Agent memory shall be subject to defined retention, review, and purge requirements. AI Agents shall not independently determine, expand, or initiate new uses of sensitive or federally protected data outside the approved purpose and authorized scope.

Auditability and logging

AI Agents and Agentic Systems shall generate logs sufficient to support oversight, auditing, investigation, reconstruction, replay of consequential actions and incident response. Logs shall capture, as applicable:

- the human user initiating the request
- agent identity
- Agent Sponsor
- timestamp
- objective or task assigned
- data sources accessed
- tools, APIs, or systems invoked
- actions attempted and completed
- records created, changed, transmitted, or deleted
- human approvals or denials; errors, exceptions, overrides, or escalations
- outputs generated
- security or policy violations

Logs generated by or associated with AI Agents shall be tamper-evident maintained in an append-only format, and protected through integrity controls, such as hash chaining, where technically feasible. Logs shall be stored in systems that the AI Agent cannot modify and shall be accessible for review and audit purposes. Log retention shall comply with applicable state record retention schedules for security requirements. For AI Agents whose action space includes benefits administration, payments, or the creation of official records, logs shall contain sufficient detail to enable an independent third party to reconstruct, replay and evaluate each consequential action.

Monitoring and runtime controls

State agencies or Sponsoring Organization shall monitor AI Agents and Agentic Systems for:

- unauthorized actions
- access outside approved scope
- prompt injection or agent hijacking
- data leakage or exfiltration
- overreliance by users
- automation drift and goal deviation
- bias or discriminatory outcomes
- hallucinations or unsupported outputs
- unexpected tool use
- performance degradation
- security events
- failure to escalate when required

Continuous monitoring is required for all AI Agents or Agentic systems that are classified as an autonomy level of three (3) or above using the framework established by this standard. Continuous monitoring shall include, at a minimum, automated anomaly detection and real-time alerting. Without continuous monitoring in place the state agency shall suspend the AI Agent's operation until monitoring is implemented or reestablished. The CIO or CISO may require independent validation, red-team testing or additional technical controls based on the AI Agent's risk, AI Agents shall be configured to fail-safe by default. If an AI Agent encounters an undefined condition, a boundary conflict, or a suspected integrity failure, the AI Agent shall stop operation and escalate to a human rather than attempting to continue or improvise.

Prohibited uses of AI Agents and Agentic Systems

AI Agents and Agentic Systems shall not:

- autonomously make or finalize legally significant decisions without authorized human approval. Independently approve benefits, enforcement actions, licenses, permits, payments, hiring decisions, disciplinary actions or procurement awards.
- circumvent human review, workflow approvals, access controls, or change management
- access systems or data outside their approved action space
- modify official records without authorization
- send external communications without required approval
- represent themselves as human users
- create or supervise other AI Agents without approval
- execute code, scripts, or system commands in production without approved controls
- perform surveillance, profiling or monitoring of individuals beyond approved legal authority
- use sensitive or federally protected data in public or unapproved AI systems

Incident reporting and deactivation

Actual or suspected misuse, unauthorized access, agent hijacking, data leakage, harmful output, unintended action, or violation of this Standard shall be reported immediately through the incident reporting procedures required by the [Use of AI in Oklahoma State Government Standard](#). Agencies

shall maintain the ability to pause, disable, revoke access or deactivate AI Agents in events such as, but not limited to:

- when ownership is unclear
- the AI Agent behaves outside approved scope
- a security event is suspected
- the AI Agent accesses unauthorized data
- the AI Agent performs or attempts unauthorized action
- continued operation presents unacceptable risk

Platforms hosting AI Agents that are classified as an autonomy level of three (3) or above using the framework established by this standard shall support automatic suspension on defined triggers and should not require manual deactivation.

Periodic review

AI Agents and Agentic Systems shall be reviewed quarterly to maintain overall health including

- continued business need
- compliance with approved purpose and action space
- appropriate autonomy level
- appropriate access privileges
- accuracy, reliability, and performance
- adequacy of logging and monitoring
- absence of automation drift
- continued compliance with the law, applicable policies, and security requirements

Procurement requirements

Any contract for the procurement of an AI Agent or Agentic System shall identify, at a minimum, the following information agent capabilities:

- autonomy level
- model provider or providers
- third-party agent components
- tool and API integrations
- data sources and data residency
- memory, retention, and logging capabilities
- security controls
- identity and access controls
- human approval mechanisms
- monitoring and audit features
- ability to pause, disable, or revoke agent access
- supplier responsibilities for incident response, audit support, and compliance

As a condition of procurement, platforms hosting state AI Agents shall require the platform to provide the following capabilities:

- observability telemetry exposed to the state
- support for deterministic action boundaries

- policy enforcement on external calls
- escalation-to-human paths
- automatic suspension capability
- standards-based agent identity
- export of logs and configurations in usable formats

Compliance

This standard shall take effect upon publication and is made pursuant to Title 62 O.S. §§ 34.11.1 and 34.12 and Title 62 O.S. § 35.8. OMES IS may amend and publish the amended standards policies and standards at any time. Compliance is expected with all published policies and standards, and any published amendments thereof. Employees found in violation of this standard may be subject to disciplinary action, up to and including termination..

Rationale

To ensure efficient, secure, and cost-effective delivery of essential public services, all state agency IT purchases and projects must receive central approval. This allows the Chief Information Officer to evaluate agency needs and capabilities, strengthen data protection and security, and streamline and consolidate systems to reduce costs for taxpayers.

References

- [Use of AI in Oklahoma State Government Standard](#), including vendor and contract requirements.
- [OMES Information Security standards](#) (incident reporting and notification timeframes).
- [NIST AI Risk Management Framework \(AI 100-1\)](#).
- [NIST AI Agent Standards Initiative \(CAISI\)](#).
- [OWASP Top 10 for Agentic Applications \(2026\)](#);
- [MITRE ATLAS](#).

Revision history

This standard is subject to periodic review to ensure relevancy.

Effective date: 07/30/2026	Review cycle: Annual
Last revised: 07/30/2026	Last reviewed: 07/30/2026
Approved by: Dan Cronin, Chief Information Officer	