



ENTERPRISE
information services

AI at the State of Oregon

PMUG July 2026





Agenda

- ▶ Responsible AI Usage Policy overview
- ▶ AI for Project Managers
- ▶ (time permitting): brief demo of Notebook tool in M365 Copilot





Goals of the Responsible AI Usage Policy

- ▶ **Primary Goal:** Governance that instills confidence and trust in the state's use of AI
- ▶ The statewide anchor for safe, transparent, and effective use of AI by state agencies
- ▶ Provide clear pathways and obligations for agencies
- ▶ Provide clear obligations and support from EIS: tools, templates, training, and a new, responsive governance forum





Responsible AI Policy – 11 sections

1. Principles of Responsible AI
2. Enterprise AI Advisory Committee
3. General Recommended Usage
4. AI Risk Management Framework
5. Agency AI Adoption Plans

6. Statewide AI Registry
7. Security
8. Unauthorized Usage
9. Workforce Development and Training
10. Transparency and Public Reporting
11. Sandbox for Custom Development and AI Model Installation





1. Principles of Responsible AI

- ▶ Human Oversight and Review
- ▶ Equity and Representation
- ▶ Privacy and Confidentiality
- ▶ Risk Management
- ▶ Safety
- ▶ User Experience, Disclosure and Feedback
- ▶ Workforce Preparedness and Understanding





2. Enterprise AI Advisory Committee

- ▶ **Goal:** Provide a forum to discuss policies, procedures and guidance supporting the responsible use of AI.
- ▶ **Key Document:** *Enterprise AI Advisory Committee Charter*
- ▶ **Details:**
 - Chaired by State CIO
 - EIS and Agency representatives
 - Both IT and business leaders – CIO Council co-chairs, DAS Director or designee, Governor’s Director of DEI, and seven senior agency leaders





3. General Recommended Usage (Copilot Use)

- ▶ **Goal:** Provide foundational guidance to agencies and employees supporting the use of Copilot
- ▶ **Key Document:** *Copilot Usage Policy* and associated procedure
- ▶ **Details:**
 - Key direction: a) Appropriate data use, and b) Appropriate human review
 - Recommended Use Guidance matrix provided as attachment
 - Agencies may supplement with their own specific guidance and use cases



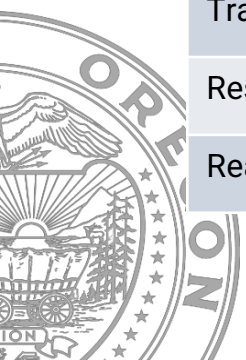


Copilot Recommended Use Guidance Matrix

DRAFT

Copilot Use	Audience / Scope		
	Limited / Personal Use	Moderate / Agency Use	Broad / Public Use
Proofread/rewrite (<~25% AI generated)	✓	✓+	⚠
Generate content from detailed human notes or prompts (<~50% AI generated)	✓	✓+	⚠
Generate content from limited human notes or vague prompts (>~75% AI generated)	⚠	⚠	⛔
Summarization of documents or meetings	✓	✓+	⛔
Transcription of voicemail	✓	✓+	⚠
Translation of documents	✓	✓+	⚠
Researching	✓	✓+	⚠
Real-time translation	✓	⚠	⛔

- ✓ **Routine Use:** Generally safe, routine uses with limited distribution. Human review is still needed.
- ✓+ **Appropriate Use, with Documented Human Review:** Generally safe, routine uses with broader distribution that should have documented human review.
- ⚠ **Use Caution:** Higher-risk scenarios that require significant human review due to large or unknown audiences or complex outputs.
- ⛔ **Discouraged:** Uses that pose significant risk and should not be attempted.





4. AI Risk Management Framework

- ▶ **Goal:** Help agencies understand the risk involved with potential AI uses, and provide recommendations for mitigations and safeguards
- ▶ **Key Documents:**
 - *AI Risk Management Framework*
 - *AI Use Form*
 - **Standards for: User Feedback, Disclosure, Human review**
- ▶ **Details:**
 - AI Use Form serves as initial intake and helps determine what constitutes higher-risk use cases





Higher-risk AI uses

Field on AI Use Form	Condition
Audience	Public/Resident audience
Level 3 or Regulated Data	L3 (Restricted) or presence of HIPAA, FERPA, CJIS, etc.
Rights Impacting	Yes, if selected (see list below)
Safety Impacting	Yes, if selected (see list below)
AI processing categories	Yes, if selected (see list below)

Rights impacting flags:

- Predictive analysis of crime
- Evaluating a person in an employment context
- Evaluating a person in a learning, training or educational context
- Evaluating a person's eligibility for a service or good
- The administration of justice, law enforcement or immigration

Safety impacting flags:

- Acting as a safety component in a system or products
- Informing safety decisions (i.e. water treatment ratios, load-bearing specifications for bridges)

AI processing category flags:

- Biometric identification / verification, including facial recognition
- Scraping images from databases
- Surveillance / tracking
- Sentiment / emotion recognition





Standards for Disclosure, Feedback, and Human Review

- ▶ **Disclosure:** Must disclose to users when they are interacting with AI-generated content.
- ▶ **Feedback:** Must provide an opportunity for users to provide feedback that is reviewed by a human. Critical for catching errors, continuous improvement, and measuring success.
- ▶ **Human review:** “A human is always responsible and accountable”:
 - At minimum, a written policy for human review of AI outputs, including what must be reviewed and who is approved to conduct the review.
 - Human review standards for higher-risk use cases still need to be developed.





5. Agency AI Adoption Plans

- ▶ **Goal:** Support agencies in establishing good AI governance and integrating with business and IT strategic plans and data governance
- ▶ **Key Document:** *AI Agency Adoption Plan Template*
- ▶ **Details:**
 - Recommendations for action-oriented AI governance, including AI use case sourcing and evaluation
 - Also includes data management, AI tool use and training, and defining success metrics.





6. Statewide AI Registry

- ▶ **Goal:** Provides complete inventory of AI tools and uses across the enterprise
- ▶ **Key Document:** Application Portfolio Management system is in development by the EIS Strategy and Design team. More info forthcoming.
- ▶ **Details:**
 - A single source of truth for application inventory and AI details, enabling monitoring, risk management, and reuse through EIS's Application Portfolio Management tool, developed in partnership with pilot agencies including DHS/OHA, PERS and DAS IT.
 - EIS will use information from the AI Registry for the annual report on AI usage.





7. Through 11.

- ▶ **Security and Unauthorized Usage**
- ▶ **Workforce development and training:** provide enterprise training on AI literacy and risk awareness
- ▶ **Public disclosure and transparency:** produce an annual report on AI use in state government
- ▶ **Sandbox for AI development:** provide requirements for safe AI custom development





ENTERPRISE
information services

AI in Action: Habits, Guardrails, and Tips





The AI Mindset

- ▶ It's moving fast and the tools change week to week. Nobody's an expert yet; we're all learning this together.
- ▶ Give yourself grace, you learn by using it, not by waiting until you feel ready. Fumbling is part of the process.
- ▶ AI generates, so verify outputs. Calculator-like power, but it creates rather than computes. A single prompt can give different outputs, and it can be confidently wrong.
- ▶ Find your lane, leverage AI's compute and match it to the right tasks (drafting, summarizing, exploring), then let it carry the heavy lifting at a scale you couldn't by hand.





Why AI Projects Succeed or Fail

- ▶ Start small, simple and local-Avoid over-engineering; prefer RAG or fit-for-purpose solutions.
- ▶ Start with a real business problem tied to measurable outcomes.
- ▶ Validate data readiness early: quality, access, bias, and legal constraints.
- ▶ Treat launch as the starting point; governance and monitoring matter.

From the Trenches

Starting with too ambitious of an idea along with underestimating the work around data readiness are the two top time sucks for AI projects.





Seven Habits of Successful AI Projects

1. Business Understanding: Define measurable success criteria.
2. Data Understanding: Inventory and validate data quality early.
3. Data Preparation: Standardize, label, and build repeatable pipelines.
4. Model Deployment: Test in the use case, but many frontier models are great generalist
5. Evaluation: Test accuracy, safety, bias, and baselines.
6. User Testing: SME are great first round testers, followed by standard UAT.
7. Operationalization: Monitor drift; assign product ownership.

From the Trenches

AI projects are fundamentally IT projects, so if you follow Software Development Lifecycle Cycle (SDLC) and Application Lifecycle Management (ALM) you will cover many of the best practices and recommendations!





Operationalization: Running AI as a Product

- ▶ Monitor accuracy, drift, bias, and pipeline health.
- ▶ Use versioning, rollback, and CI/CD practices for safe updates.
- ▶ Change management: train users, build feedback loops.

From the Trenches

No-one has seen the full lifecycle of AI in a business environment, so full costs and resource needed can be both difficult to nail down and predict.



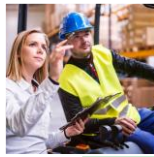


Organizational Approach to Using AI



Leaders

- Set direction
- Pick the right problems
- Ensure support
- Celebrate success



Supervisors

- Design workflows
- Set expectations
- Create “Taste” models
- Share success



Staff

- Explore use cases
- Try Operationalizing
- Share with Managers and Team Members

From the Trenches

The power and speed of AI mean the whole team will need to be thinking and focusing on how to best fit it into the organization.





Resources

- ▶ **Guide to Redesigning the Project Manager Job in the Age of AI.** Gartner Analysts, Version 20260105.
- ▶ **Guide to Leading and Managing AI Projects (PMICPMAI)** Project Management Institute (PMI).

