



800-767-4991 · www.theconnectgroup.com

Protecting Your Business in the Age of AI

A Practical Guide to AI Governance, Security, and Compliance

■ Why This Guide Matters

Artificial Intelligence is transforming how businesses operate—from automated customer service to intelligent data analysis. But with these powerful capabilities come significant risks: data breaches, compliance violations, and security vulnerabilities that can threaten your entire organization.

This guide provides a practical framework for implementing AI safely and responsibly, helping you harness the benefits of AI while protecting your business, your clients, and your reputation.

78%

of businesses now use AI
in some capacity

43%

have experienced an
AI-related security incident

\$4.88M

average cost of a
data breach in 2024

Prepared by The Connect Group
Your Trusted Technology Partner Since 1991

***Disclaimer:** This guide is provided for informational purposes only and does not constitute legal, regulatory, or professional advice. It is intended as a framework to help organizations implement AI safely and responsibly. Every organization's needs are unique—please consult with qualified legal and compliance professionals before adopting any policies or procedures.*



1. Understanding the AI Landscape

The rapid adoption of AI tools presents both tremendous opportunities and significant risks for businesses. Before implementing any AI governance policy, it's essential to understand what we're dealing with and why proper oversight is critical to your organization's success.

Key Definitions

- **AI System:** Any software that generates outputs—content, recommendations, or automated actions—based on prompts or data inputs. Examples include Microsoft 365 Copilot, Azure OpenAI, Amazon Bedrock, and similar platforms.
- **Client Data:** Information owned by or processed for clients under your service agreements, including personal information and confidential business data.
- **Personal Information / Sensitive PI:** Data as defined by privacy regulations like CPRA/CCPA, including names, contact information, financial data, and any information that could identify an individual.
- **Approved AI Platform:** AI tools that have been vetted and configured by your IT and Security teams, operating within enterprise security controls.
- **Consequential Decision:** Any automated decision with legal, financial, or similarly significant effects on a person or business—these require special human oversight.

■The Risk is Real

Why This Matters: When employees use unapproved AI tools or input sensitive data into consumer AI platforms, they may inadvertently expose your business to data breaches, compliance violations, and significant legal liability. A clear governance framework protects everyone.



2. Legal & Compliance Framework

Your AI governance policy doesn't exist in a vacuum—it must align with existing legal requirements and industry standards. Here are the key frameworks your organization should consider:

Privacy Regulations

- **CPRA/CCPA:** California's comprehensive privacy laws establish service-provider and contractor obligations for handling personal information.
- **GLBA:** The Gramm-Leach-Bliley Act governs financial institutions and requires safeguarding customer information.
- **HIPAA:** Healthcare organizations must ensure AI tools comply with patient data protection requirements.
- **State Data Breach Laws:** Various state laws mandate notification and response procedures when data is compromised.

Security & Assurance Standards

- **SOC 2 / ISO 27001:** Industry-standard frameworks for demonstrating security controls and practices.
- **NIST Cybersecurity Framework:** Comprehensive guidelines for managing cybersecurity risk.
- **MSPAlliance MSP/Verify or Cyber Verify:** Specialized certifications for managed service providers.

AI-Specific Governance

- **NIST AI Risk Management Framework 1.0:** The foundational framework for managing AI risks throughout the system lifecycle.
- **NIST Generative AI Profile:** Specific guidance for organizations deploying generative AI systems.
- **Microsoft Responsible AI Standard (v2):** A useful external benchmark for AI development and deployment practices.



3. Core AI Governance Principles

Every effective AI policy is built on foundational principles that guide decision-making across your organization. These principles should inform every AI-related activity:

Client-First Approach: Client data protection and interests take precedence in all AI decisions. Never compromise client security for convenience.

Least Privilege Access: AI systems should have only the minimum access necessary to perform their intended function—no more.

Security by Design: Build security controls into AI implementations from the start, not as an afterthought.

Privacy by Default: Default settings should always favor privacy protection. Opt-in to data sharing, never opt-out.

Human-in-the-Loop: Require human review and approval for any consequential decisions. AI assists; humans decide.

Documented Exceptions: Any deviation from standard policy requires written approval with clear mitigations and expiration dates.

Transparency & Auditability: Log AI access, prompts (where permitted), and maintain output provenance for material client work.

No Training on Client Data: By default, enterprise AI services should not train models on your client data.

■ Remember

These principles aren't just guidelines—they're the foundation for building trust with your clients and protecting your business from liability. When in doubt, return to these principles to guide your decision-making.



4. Approved vs. Prohibited AI Activities

Clear boundaries help employees understand what's acceptable and what's not. Here's a practical breakdown:

✓ Approved Activities (with training + logging)

- Drafting internal documents, ticket summaries, and code comments using approved enterprise AI tools (without client secrets or sensitive data).
- Summarizing publicly available information or client materials that are explicitly permitted in your service agreements.
- Automating repetitive tasks like knowledge retrieval and report templates—with mandatory human review before any external release.
- Using AI for research, brainstorming, and ideation on internal projects.
- Generating meeting notes and action items from approved recording tools.

✗ Prohibited Activities

The following activities are strictly prohibited and may result in disciplinary action:

- Entering client Personal Information (PI), Protected Health Information (PHI), Payment Card Industry data (PCI), or Criminal Justice Information (CJI) into ANY non-approved AI tool
- Inputting secrets, API keys, credentials, firmware, or proprietary code into consumer AI platforms
- Using device logs containing identifiable information in AI prompts
- Using personal or consumer AI accounts (like personal ChatGPT accounts) for any company work
- Relying on AI outputs as final decisions for legal, financial, HR, or operational matters without authorized human sign-off
- Enabling facial recognition or biometric classification without executive approval, documented testing, client consent, and legal review



5. Data Protection & Prompt Hygiene

How you interact with AI systems matters just as much as which systems you use. Good "prompt hygiene" protects sensitive information and reduces risk:

Best Practices for Safe AI Interaction

- **Use Synthetic or Redacted Data:** When testing AI capabilities or developing workflows, use fake or anonymized data instead of real client information.
- **Minimize Data Exposure:** Include only the minimum necessary information in prompts. Ask yourself: 'Does the AI really need this data to complete the task?'
- **Maintain Approved Repositories:** Store prompts and outputs related to material client work in approved systems with proper retention tags.
- **Honor Data Residency:** Respect client data residency requirements. Never transfer data across regions without explicit approval.
- **Review Before Sending:** Take a moment to review your prompt before submission—is there any sensitive information that should be removed?

Data Type	Can Use in Approved Enterprise AI?	Can Use in Consumer AI?
Public information	Yes	Yes (with caution)
Internal documents (non-sensitive)	Yes, with review	No
Client names/contact info	Approved tools only	Never
Financial data	Approved tools only, minimize	Never
Health information (PHI)	Special approval required	Never
Credentials/API keys	Never	Never
Proprietary code	Approved tools only	Never



6. Enterprise AI Platform Guidelines

Understanding the data boundaries of your AI tools is critical. Enterprise platforms offer important protections that consumer versions don't provide:

Microsoft 365 Copilot (Enterprise)

- Prompts and responses are processed within the Microsoft 365 service boundary
- Your data is NOT used to train foundation large language models
- Enforce tenant controls: SSO/MFA, Data Loss Prevention (DLP), eDiscovery
- Data stays within your compliance boundary

Azure OpenAI Service

- Prompts and outputs are NOT used to train, retrain, or improve base models by default
- Abuse monitoring may apply per enterprise terms
- Supports disabling content logging in eligible subscriptions
- Enterprise-grade security and compliance controls available

Amazon Bedrock

- Prompts and completions are NOT used to train AWS models
- Your data is NOT distributed to third parties
- Private fine-tuning available on isolated copies
- Full data isolation within your AWS account

■ Security Configuration

Critical Configuration: Disable consumer endpoints for work accounts. Only allow enterprise tenants under company SSO/MFA. For other providers (Google Cloud, etc.), apply equivalent enterprise data boundaries, prohibited-use categories, and logging requirements.



7. Risk Management & Governance

Before deploying any AI solution—even for internal pilots—conduct a structured risk assessment. This doesn't need to be complex, but it must be documented.

AI Risk Assessment Framework

For each AI implementation, evaluate and document:

- **Context:** What problem is this solving? Who will use it? What data is involved?
- **Data Inventory:** What types of data will the AI access or process? Is any of it sensitive?
- **Model Choice:** Which AI platform or model? Is it approved? What are its limitations?
- **Grounding Approach:** How will the AI access information? RAG (Retrieval Augmented Generation) is preferred for accuracy.
- **Human-in-the-Loop:** Where are human review points? Who approves outputs before action?
- **Success Metrics:** How will you measure effectiveness and accuracy?
- **Rollback Plan:** If something goes wrong, how do you revert to the previous state?

High-Risk Categories Requiring Additional Approval

■High-Risk Categories

The following AI use cases are considered high-risk and require explicit approval from both Legal and Security teams before proceeding:

- **Employment decisions:** Hiring, termination, performance evaluation
- **Biometric systems:** Facial recognition, voice identification, behavior analysis
- **Financial decisions:** Credit scoring, fraud detection, investment recommendations
- **Public safety applications:** Any AI used in law enforcement or emergency services context

Ongoing Governance

- **Model/Feature Factsheet:** Document the purpose, limitations, metrics, and support contacts for any external-facing AI.
- **Quarterly AI Oversight Review:** Review exceptions, incidents, supplier attestations, and monitor for model drift or bias.
- **Continuous Improvement:** Update your policy as technology evolves and new risks emerge.



8. Incident Response & Monitoring

Even with the best policies, incidents can occur. Having a clear response plan ensures you can act quickly and effectively:

Logging & Monitoring Requirements

- Log AI access, tool identifiers, and output hashes for material deliverables where legally permitted.
- Use privacy-preserving logging when prompts contain personal data.
- Maintain audit trails that can demonstrate compliance during reviews.
- Set up alerts for unusual activity patterns or potential policy violations.

Incident Classification

Treat the following as security incidents requiring immediate response:

- Client data entered into non-approved AI tools (data misrouting)
- Credentials, API keys, or secrets exposed through AI prompts
- Unauthorized AI outputs shared externally
- Discovery of AI-generated content in production without required human review
- Evidence of model hallucination in client-facing materials

Response Timeline

- **Notification within 72 hours:** Where contractually required, notify affected parties within 72 hours of confirming a material incident.
- **Document everything:** Maintain detailed records of the incident, response actions, and outcomes.
- **Root cause analysis:** Determine what went wrong and implement controls to prevent recurrence.
- **Policy updates:** If the incident reveals a gap in your policy, update it promptly.

Pre-Deployment Testing

- Run red-team tests before external AI launches to identify vulnerabilities.
- Conduct bias checks to ensure AI outputs don't discriminate unfairly.



9. Vendor & Third-Party Controls

Your AI policy must extend to vendors and subprocessors. Their practices directly impact your security and compliance posture:

Subprocessor Management

- Maintain an approved list of AI subprocessors and model providers.
- Publish trust pages and assurance documentation to clients upon request.
- Review and update the approved vendor list at least annually.
- Assess new AI vendors against your security and compliance requirements before onboarding.

Flow-Down Requirements

All AI vendors and subprocessors must contractually agree to:

- No training on client data by default
- Encryption of data in transit and at rest
- Access controls aligned with least-privilege principles
- Breach notification within specified timeframes
- Data deletion upon contract exit
- Compliance with applicable privacy regulations

Certifications & Attestations

Require annual attestations or certifications from AI vendors. Look for:

Certification	What It Covers
SOC 2 Type II	Security, availability, processing integrity, confidentiality, privacy controls
ISO 27001	Information security management system
SOC 2 + AI	Extended controls for AI-specific risks
NIST AI RMF	AI risk management practices



10. Implementation Checklist

Ready to put this guide into action? Use this checklist to ensure your organization is prepared:

Phase 1: Foundation (Week 1-2)

- Identify and inventory all AI tools currently in use across the organization
- Classify AI tools as approved, pending review, or prohibited
- Document current data flows involving AI systems
- Identify key stakeholders (IT, Security, Legal, HR, Operations)
- Establish an AI governance committee or owner

Phase 2: Policy Development (Week 3-4)

- Customize this framework for your organization's specific needs
- Define approved AI platforms and configuration requirements
- Establish the exception request and approval process
- Create incident response procedures specific to AI
- Have Legal review the final policy document

Phase 3: Training & Rollout (Week 5-6)

- Develop training materials for all employees
- Conduct mandatory training sessions
- Collect signed acknowledgments from all staff
- Implement technical controls (SSO, DLP, endpoint blocks)
- Establish monitoring and logging infrastructure

Phase 4: Ongoing Operations

- Schedule quarterly AI governance reviews
- Set up annual training refreshers
- Establish vendor attestation collection schedule
- Create metrics dashboard for AI usage and incidents
- Plan for policy updates as technology evolves



800-767-4991 · www.theconnectgroup.com

Ready to Secure Your AI Future?

Implementing a comprehensive AI governance policy can feel overwhelming, but you don't have to do it alone. The Connect Group has been helping Montana businesses navigate technology challenges since 1991—and we're ready to help you with AI.

■How We Can Help:

- **AI Readiness Assessment:** We'll evaluate your current AI usage and identify risks
- **Technical Implementation:** Configure enterprise AI tools with proper security controls
- **Training Programs:** Get your team up to speed on responsible AI use
- **Ongoing Managed Services:** Continuous monitoring and support for your AI infrastructure

Contact The Connect Group Today

Billings (Main Office)

1306 Central Avenue
Billings, MT 59102
Phone: (406) 248-8900

Bozeman Office

251 Evergreen Drive
Bozeman, MT 59715
Phone: (406) 922-5000

Helena Office

Phone: (406) 324-7160

Visit us online at www.theconnectgroup.com
Schedule a free consultation to discuss your AI governance needs