

The AI Risk-Tiering Model

Matching Governance to Business Consequences

By Kimberly Wiethoff, MBA, PMP, PMI-ACP

Not every AI initiative creates equal exposure. Effective governance should be proportional to the consequences of the system's decisions — directing rigorous oversight where it matters most, while keeping the path clear for low-risk innovation.

#AIGovernance #AIRiskManagement #ResponsibleAI #AITransformation
#AgenticAI #AIPortfolioManagement #AIOperatingModel #EnterpriseAI
#DigitalTransformation #ProjectManagement



The Governance Gap: Two Failing Extremes

Over-Governed

Applying the same extensive review process to every AI use case — regardless of consequence. This slows low-risk experimentation, frustrates teams, and pushes employees toward unapproved tools.

Under-Governed

Allowing each team to determine its own controls. Consequential applications — those that authorize payments, evaluate employees, or communicate directly with customers — operate without sufficient oversight.

 Neither approach scales. Both create avoidable organizational risk.

The solution is governance that is **proportional to consequences**. Higher-risk applications require stronger evidence, tighter decision boundaries, more monitoring, and clearer human authority. Lower-risk applications should meet minimum standards without being forced through controls designed for systems that can cause material harm.



Function Doesn't Determine Risk

It is tempting to classify AI applications by technology category — chatbot, forecasting model, recommendation engine, document generator, autonomous agent. But the category does not determine the business risk.

Two Chatbots

One answers general questions using public information. Another provides customers with account-specific financial guidance. Same technology label — vastly different exposure.

Two Predictive-Maintenance Models

One recommends an inspection a technician can decline. The other automatically changes equipment settings in a safety-sensitive environment. Similar algorithms — completely different consequences.

Risk depends on **what the AI is allowed to do**, what information it uses, who is affected, and how easily a harmful outcome can be stopped or reversed.

The Seven Dimensions of AI Risk

A practical model evaluates several dimensions rather than relying on a single score. Each factor contributes to the overall tier — and any one dimension can escalate classification if its severity is high enough.



Autonomy

Can the AI act independently, or does every output require human review before any action is taken?



Financial Exposure

What is the maximum dollar impact of a single action, and how many actions could occur before detection?



Safety & Human Impact

Could errors affect physical safety, employment, health, or access to essential services?



Data Sensitivity

Does the system access confidential, regulated, or personally identifiable information?



Reversibility

Can the action be stopped or corrected before permanent financial, legal, or reputational harm occurs?



Reach & Scale

How many people, transactions, or systems could be affected — and how quickly?

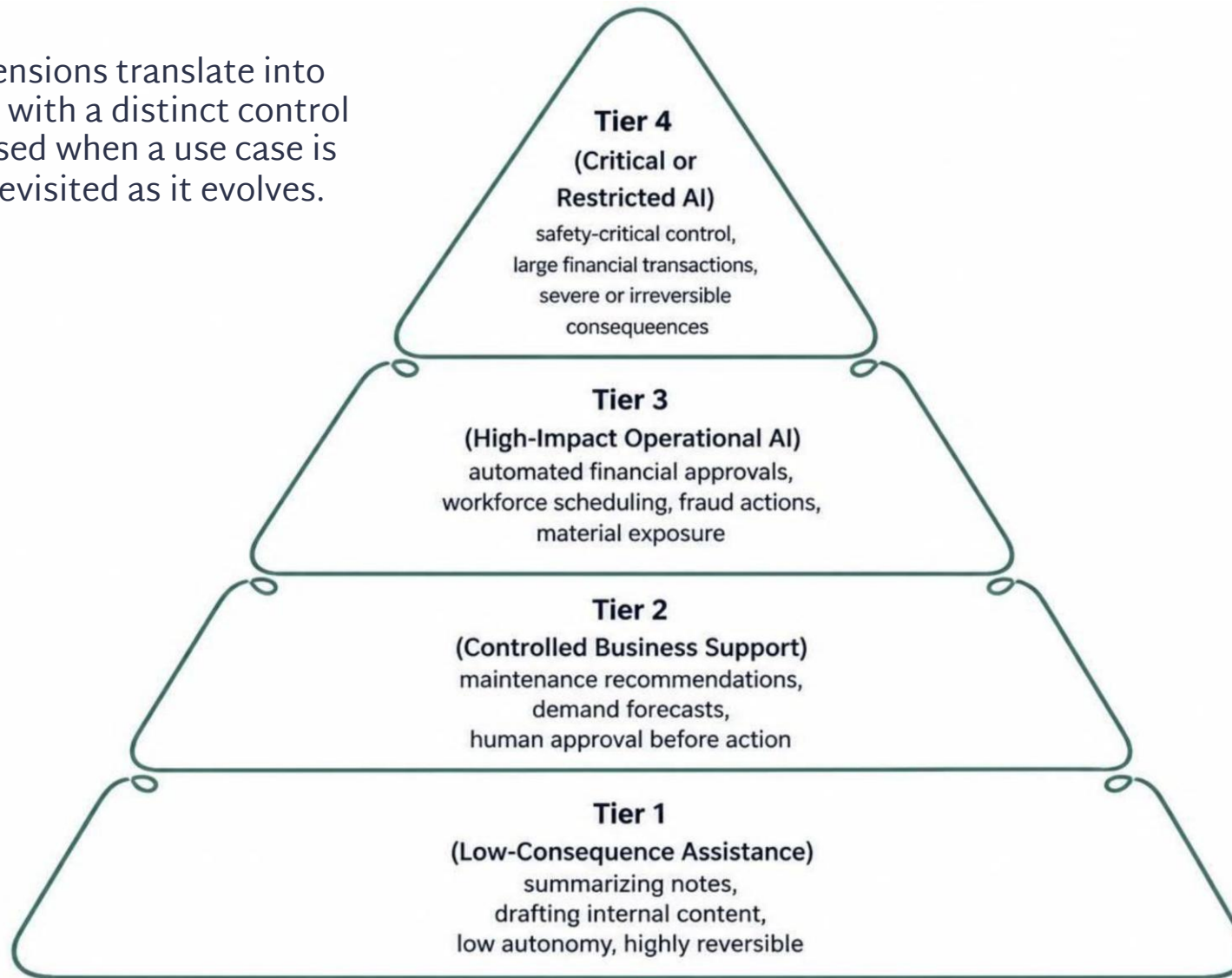


Detectability

How quickly will an error become visible, and is monitoring sufficient to reconstruct what happened?

A Four-Tier AI Risk Framework

Seven risk dimensions translate into four tiers, each with a distinct control package, assessed when a use case is proposed and revisited as it evolves.



Tier 1: Low-Consequence Assistance

Typical Characteristics

- Low autonomy — purely generative or informational
- Limited or non-sensitive data
- Low financial and safety exposure
- Small number of affected users
- Easy human review of all output
- Highly reversible — drafts can always be edited

Example Use Cases

- Summarizing non-sensitive meeting notes
- Drafting internal communications
- Organizing public information
- Brainstorming and ideation support
- Reformatting documents
- Low-impact administrative work

Appropriate Controls

Approved tools, basic acceptable-use requirements, data-handling rules, employee training, and standard access logging. The process must be **lightweight enough** to encourage use of approved capabilities over shadow AI.

Tier 2: Controlled Business Support

Typical Characteristics

- Advisory or limited-action role
- Moderate operational or financial impact
- Human approval required before consequential action
- Reversible decisions with manageable reach
- Defined business ownership

Example Use Cases

- Service-request categorization and routing
- Maintenance recommendations and anomaly detection
- Demand forecasting and scheduling recommendations
- Draft customer responses and contract analysis

Appropriate Controls

Documented use case and owner, data and privacy assessment, performance testing, human-review requirements, override and escalation paths, user training, and periodic performance evaluation.

- ❏ The primary governance challenge at Tier 2: ensuring that human review is real, informed, and adequately staffed — not a rubber stamp.

Tier 3: High-Impact Operational AI

Typical Characteristics

- Greater autonomy — executes within defined limits
- Material financial or operational exposure
- Sensitive data and significant customer/employee impact
- Broader reach, more difficult reversal
- Strong dependence on continuous monitoring

Example Use Cases

- Automated financial approvals within limits
- Workforce scheduling affecting employee conditions
- Customer eligibility recommendations
- Fraud or security actions
- AI agents that update business systems

Key Control Requirements

Formal risk and impact assessment, independent validation, documented decision rights, defined financial and operational limits, predeployment failure testing, continuous monitoring, automated stop conditions, tested rollback capability, and detailed audit trails.

- ❏ A named executive or senior business owner must be accountable for accepting the residual risk.

Tier 4: Critical or Restricted AI

These capabilities can produce severe or irreversible consequences. For some Tier 4 uses, the correct decision may be **not to deploy** the capability, or to restrict it to a strictly advisory role.

Example Use Cases

Safety-critical control systems, high-impact employment decisions, decisions affecting essential services, large or unrestricted financial transactions, autonomous actions with legal consequences, systems affecting many people before intervention is possible.

Required Controls

Executive and specialized risk approval; legal, compliance, security, privacy, and ethics review; independent testing and assurance; mandatory human authorization; real-time monitoring; kill switches and automatic suspension; red-team and adversarial testing; formal audit evidence; board or risk-committee visibility where appropriate.

Key Characteristics

High autonomy or authority, severe human/financial/legal/safety impact, low reversibility, sensitive or regulated data, large-scale reach, difficult-to-detect failure, and limited organizational tolerance for error.

Don't Let Averages Hide a Critical Risk

A scoring formula creates consistency — but a total score must never be allowed to dilute severe exposure. An AI application may score low on financial exposure, data sensitivity, and scale yet still control safety-sensitive equipment. Averaging all factors could place it in a moderate tier despite the potential for serious harm.

Safety Override

Any critical safety impact requires Tier 4 review — regardless of scores in other dimensions.

Data Minimum

Use of highly sensitive data establishes minimum privacy and security requirements regardless of total score.

Reversibility Escalation

Low reversibility can raise the final tier assignment, even when other factors score moderately.

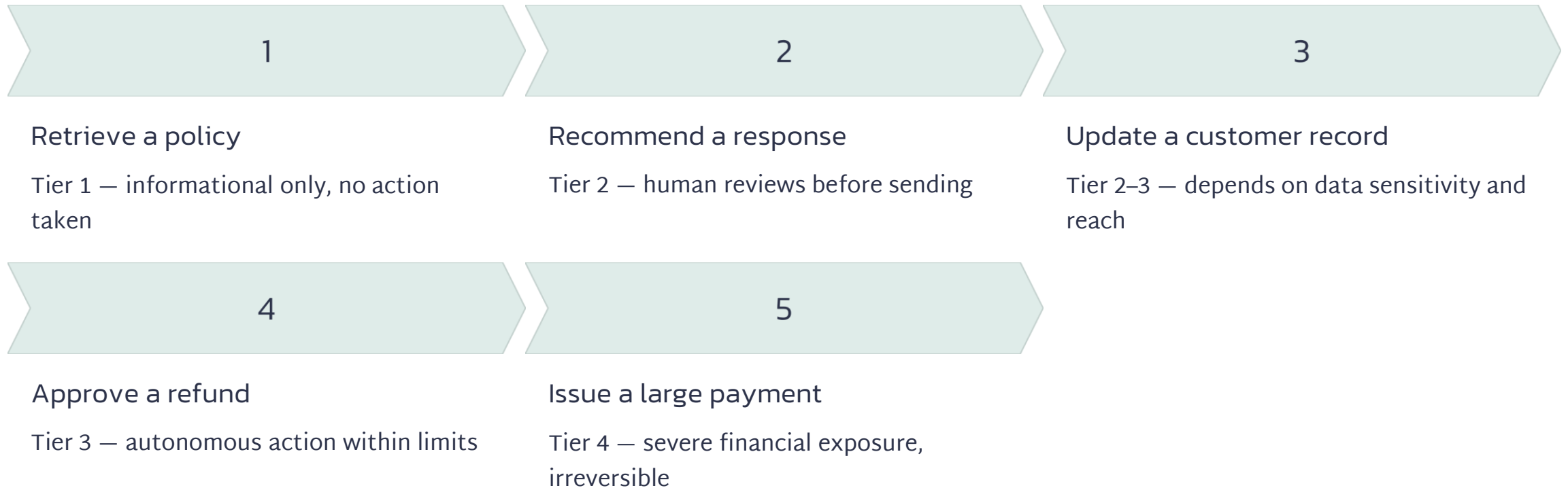
Financial Threshold

Autonomous execution above an approved financial threshold requires senior approval and documented risk acceptance.

- ❏ The model should support judgment, not replace it. Overrides must record who made the decision, why, and which additional controls are required.

Classify the Decision, Not Just the Application

A single AI system may perform tasks at multiple risk levels. Classifying the entire application at one tier produces controls that are either too weak for the highest-risk action — or unnecessarily restrictive for every low-risk task. Classify the **decisions and actions the AI is authorized to perform**.



This approach also supports **progressive autonomy**: a capability can begin as an advisor, build an evidence record, and later receive limited authority within tested boundaries.

Worked Example: The AI Refund Agent

Scenario A — Tier 2

A customer-service assistant retrieves return policy and drafts a response. A representative checks every answer before sending. The AI influences a customer interaction, but a person retains the decision and can correct the draft. Controls: documented ownership, human-review requirement, override logging, periodic performance evaluation.

Scenario B — Tier 3

The same assistant connects to the payment system and may issue refunds up to \$100 without review. Requires assessment of transaction volume, aggregate daily exposure, mistaken or fraudulent refunds, detection time, and fund recovery ability. Controls: per-transaction and aggregate limits, exception alerts, audit records, and automatic stop when patterns depart from expectations.



If the refund ceiling rises substantially or the agent can modify records across multiple brands, the team must reassess before enabling that authority. Controls belong to the *action and its consequences* — not to the label "customer-service chatbot."

Connect Each Tier to a Predefined Control Package

Risk classification has limited value if every initiative still negotiates its controls from scratch. Each tier should map to a standard package so teams know what evidence is required before development is complete.

Control Area	Tier 1	Tier 2	Tier 3	Tier 4
Approvals	Basic AUP	Business owner	Formal + executive	Board/risk committee
Testing	Standard	Performance testing	Failure + scenario	Red-team + adversarial
Human Oversight	User review	Approval required	Defined thresholds	Mandatory authorization
Monitoring	Standard logging	Override tracking	Continuous + alerts	Real-time + kill switch
Recertification	Annual inventory	Periodic review	Frequent governance	Formal + frequent

Standard packages reduce inconsistency among reviewers and allow governance teams to concentrate attention on higher-risk capabilities rather than relitigating controls at every intake meeting.

Reassess Risk When the System Changes

A risk tier is not permanent. A Tier 2 assistant can silently become a Tier 3 operational system without anyone formally recognizing the change. Risk-tier reassessment must be part of **AI change control**.

1

Authority Expands

The AI gains new decision rights, human review is removed or reduced, or transaction limits increase.

2

Reach Grows

The capability expands to new locations, customers, business units, or connected systems.

3

Data Changes

New sensitive data is introduced, or a model, prompt, vendor, or data source is substituted.

4

Performance Declines

Accuracy degrades, a new failure mode is identified, or new regulatory requirements apply.

Review dates and change triggers should be recorded alongside the original classification — making tiering a **continuing management decision**, not a one-time intake form.

What Leaders Should See at the Portfolio Level

Executives do not need every technical detail, but they must understand the **distribution of risk** across the AI portfolio. An executive dashboard should give clear visibility into where exposure is concentrated and whether governance capacity is aligned with organizational ambition.

→	→
Number of initiatives and investment by tier	High-tier initiatives approaching deployment
→	→
Control gaps and risks awaiting acceptance	Incidents, near misses, and overdue reassessments
→	→
Concentration in shared vendors or models	Initiatives whose autonomy or reach has increased

Four Questions Every Governance Model Must Answer

1. What can the AI do?

Define permitted actions and decision boundaries explicitly.

2. What could happen if it is wrong?

Assess consequence, not just probability.

3. How quickly could the organization detect and stop the harm?

Detectability and speed of response determine real exposure.

4. What evidence and controls are required before granting that authority?

Know the bar before development begins.

When governance is matched to business consequences, organizations can move quickly where risk is low — and deliberately where consequences are high.