

# HOW TO MANAGE RISKS IN AI PROJECTS

A compact guide to agile learning, structured risk decisions,  
and responsible AI adoption

---

**Agile iterations manage learning.  
Stage gates manage exposure.**

# Why AI risk management needs a different approach

AI is not just another IT component. Its behavior depends on a larger system.

**Conventional IT risk management remains necessary - but it is no longer sufficient.**

## Code

rules and integration

## Data

quality, bias, lineage

## Model

probabilistic behavior

## Provider

versions, pricing, terms

## Users

trust, misuse, adoption

## Context

process, impact, regulation

**Risk changes even when the project team does not change the code.**

# The risk profile does not shift. It expands.

AI projects keep the classic IT risks - and add new domains of uncertainty.

## Conventional IT

Technical feasibility

Scope & dependencies

Security & integration

Budget & schedule

## + AI-specific uncertainty

Data quality & rights

Model behavior & drift

Provider changes

Human oversight

Ethics, law & trust

**AI risk is socio-technical: it emerges from technology, data, people, processes, providers, and consequences.**

# The six risk domains of AI projects

A shared model helps teams avoid fragmented reviews and blind spots.

1

## Value & Adoption

Does the solution solve a relevant problem?

4

## Model & Algorithm

Does the model behave reliably enough?

2

## Business & Viability

Will the business case survive scaling?

5

## Technology & Operations

Can it be integrated, monitored, and supported?

3

## Data

Can the data be trusted and used?

6

## Compliance, Ethics & Trust

Is the use lawful, fair, transparent, and accepted?

**The domains are connected. A problem in one domain can create risk across the system.**

# Agile iterations manage learning

AI teams need short cycles that turn uncertainty into evidence.



## Inside every phase:



**Do not define iterations only by functionality. Define them by what the team needs to learn.**

# Stage gates manage exposure

A gate is not a status meeting. It is a forward-looking risk and investment decision.

**Does the available evidence justify the investment, permissions, and exposure of the next phase?**

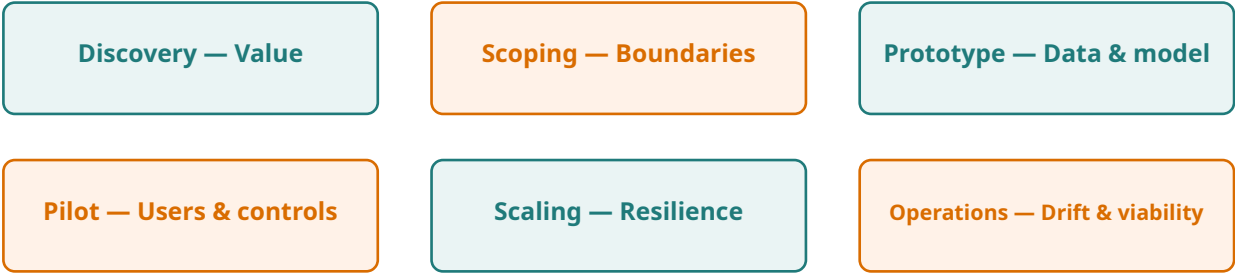
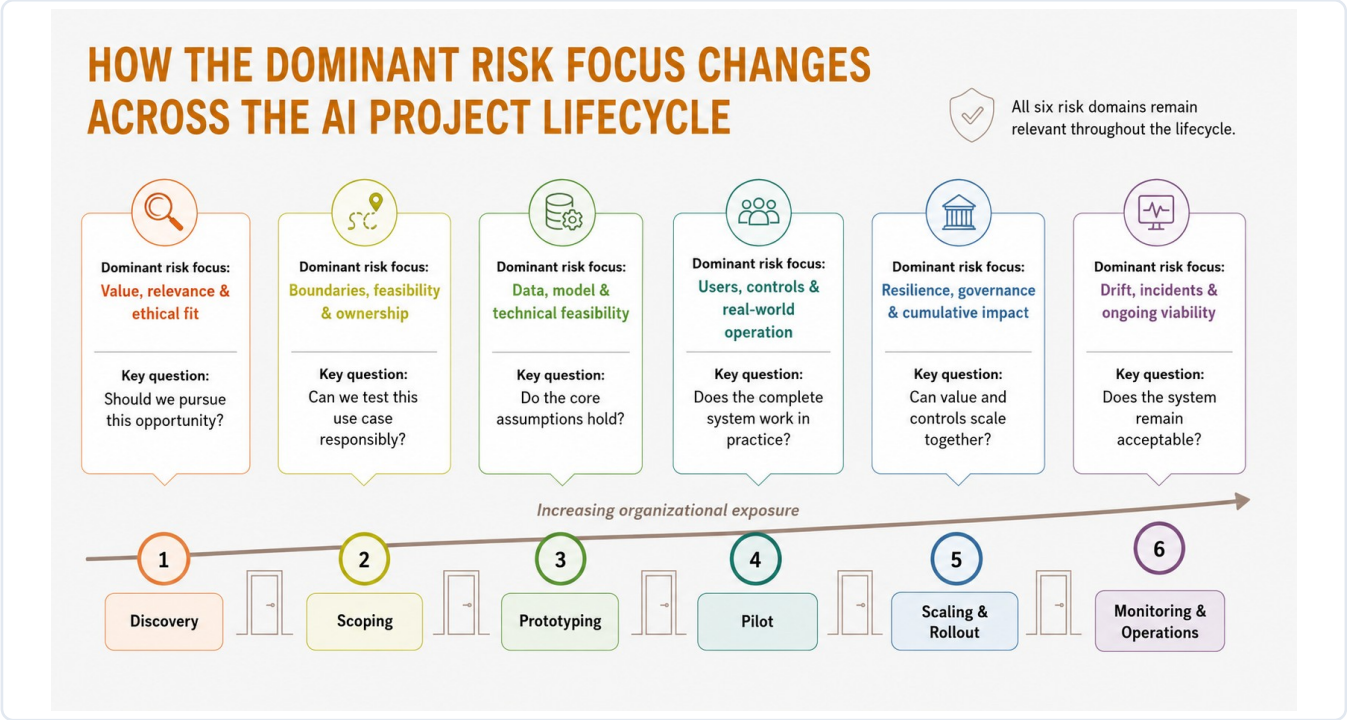
## Five possible decisions

- Go** proceed within defined conditions
- Rework** learn more before increasing exposure
- Restrict** continue only within explicit boundaries
- Pause** wait for missing conditions
- Kill** stop because value no longer justifies risk

**The strength of the gate should increase with organizational exposure.**

# The AI lifecycle changes what risk means

All six domains remain relevant. What changes is the dominant risk focus and the evidence available.



Risk management must mature at the same speed as the AI solution and the organization’s exposure.

# Evidence at the gates

Gate reviews should focus on claims, evidence, residual risks, and next exposure.

**Claim**

What the initiative currently believes

**Evidence**

What was tested and observed

**Controls**

Which safeguards exist and whether they work

**Residual risk**

What remains after treatment

**Decision**

What is approved, restricted, paused, or stopped

**A completed test is an activity. A test result is evidence.**

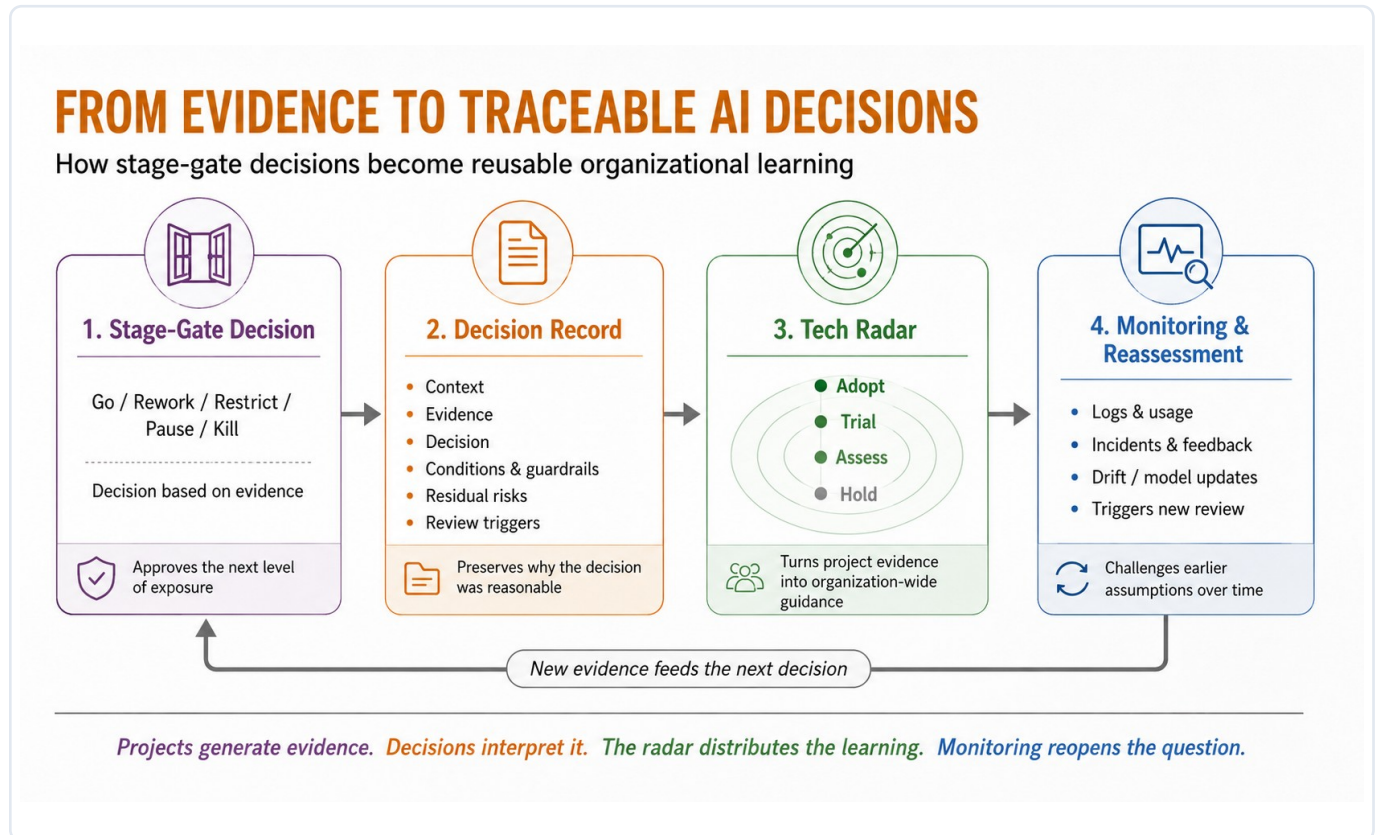
# Governance and stakeholders

Cross-functional participation is necessary. Explicit accountability is critical.



# Decision Records and Tech Radar

Traceability turns local project evidence into reusable organizational learning.



**Projects generate evidence. Decisions interpret it. The Radar distributes the learning.**

# Practical guardrails for AI projects

Good guardrails make the safe path easier - not all experimentation harder.

## PRACTICAL GUARDRAILS FOR AI PROJECTS

How teams can innovate safely within clear ethical, technical, and operational boundaries

### 1. Use Approved Tools

Only use approved AI tools, models, and platforms.  
Approval depends on the use case, data, and controls — not just the tool itself.  
Request an assessment before introducing new tools.

### 2. Protect Data

Classify data before using AI.  
Do not enter unapproved personal, customer, confidential, or IP-sensitive data.  
Use anonymized, masked, or synthetic data wherever possible.

### 3. Define Permitted Uses & Human Oversight

Make intended use, restricted use, and prohibited use explicit.  
Do not silently expand the use case.  
Human oversight must be meaningful: trained people need time, authority, and the ability to intervene.

### 4. Ensure Accountability & Traceability

A named person remains accountable for consequential AI-supported actions.  
Keep enough traceability to reconstruct what happened.  
Document exceptions, approvals, and review triggers.

### 5. Monitor, Report & Improve

Make unexpected behavior easy to report.  
Track incidents, drift, model updates, usage changes, and costs.  
Feed important findings back into risk reviews and governance decisions.

*Guardrails define the boundaries. Teams create value within them.*

|                        |                         |
|------------------------|-------------------------|
| Use case boundaries    | Data & privacy controls |
| Human oversight        | Model/output controls   |
| Operational safeguards |                         |

David Theil | DigitalisierungsCoach

Guardrails

# Common AI risk-management anti-patterns

Most failures happen when governance disconnects from learning, ownership, and action.

## One-time risk assessment

the register becomes outdated

## Framework as checklist

compliance replaces judgment

## Documents as evidence

policies are not proof

## Model-only evaluation

the work system is ignored

## Monitoring without intervention

dashboards do not decide

## Rewarding confirmation

teams prove the project right

**A gate without a clear decision is a presentation, not a control mechanism.**

# Three principles to keep

Use them as a quick test for your AI governance model.

1

## Start small enough to fail safely

Early experiments should create evidence without exposing the organization to full impact.

2

## Scale evidence, not confidence

Strong sponsorship and impressive demos are not substitutes for results under realistic conditions.

3

## Keep decisions reviewable and reversible

Every approval should define scope, conditions, owners, and review triggers.

**Innovation × Trust × Compliance = Sustainable AI Success**

**DOWNLOAD THE COMPLETE HANDBOOK**

# **How to Manage Risks in AI Projects**

A Practical Handbook for Combining Agile Learning with  
Structured Risk Decisions

[digitalisierungscoach.com/ai-risk-management-handbook](https://digitalisierungscoach.com/ai-risk-management-handbook)

**Free PDF download**

Follow David Theil on LinkedIn for practical AI strategy, risk management, and organizational change content.

---

**Move fast - but responsibly.**