

HYPE
SHORTCUTS
RISK
BIAS
OVERPROMISES

Thinking With AI

TRUST
TRANSPARENCY
ACCOUNTABILITY
FAIRNESS
HUMAN IMPACT

A Storybook for Better Decisions,
Responsible Leadership,
and Enterprise AI



Different roles. Shared responsibility. **One AI future.**

Contents & Chapter Guide

Select any chapter to jump directly to its page.



1 The AI Landscape

Understand rapid AI change and lead with clarity rather than hype.



2 The Decision Maker's Playbook

Use a practical checklist to make responsible, evidence-based AI decisions.



3 Start with Better Questions

Frame opportunities, success criteria, and organizational change before choosing technology.



4 Trust but Verify

Test claims, data, safeguards, human oversight, and rollback readiness before proceeding.



5 Choose Winnable Use Cases

Prioritize use cases that are operationally, organizationally, economically, and humanly viable.



6 Measure What Matters

Assess ROI, employee impact, and future readiness instead of focusing only on cost savings.



7 Align the Organization

Connect strategy, structure, capability, and compliance to enable responsible delivery.



8 Evaluate Continuously

Monitor drift, bias, outcomes, vendors, and regulatory signals throughout the AI lifecycle.



9 Educate Continuously

Build role-based learning and ongoing practice across the organization.



10 Pilot Safely

Run small, controlled pilots with clear objectives, measures, guardrails, oversight, and rollback plans.



11 Keep Humans in the Loop

Match human review and decision rights to the impact and risk of each AI use case.



12 Strengthen the Data Foundation

Improve data quality, governance, privacy, lineage, and observability for reliable AI.



13 Set Guardrails and Policy

Define allowed, review-required, and prohibited uses with clear accountability.



14 Decide: Build, Buy, or Partner

Compare control, speed, cost, security, integration, and vendor risk before selecting an approach.



15 Prepare for Failure

Establish incident response, escalation, containment, recovery, and learning mechanisms.



16 Scale Responsibly

Expand only after value, controls, ownership, monitoring, training, and platform readiness are proven.



17 Examples in Action

Compare realistic use cases to see how context, readiness, and risk change the decision.



18 Ask Before You Approve

Apply a repeatable pre-approval checklist to ensure ownership, evidence, guardrails, and monitoring.



19 The Journey Continues

Reinforce the leadership principles that sustain responsible AI adoption over time.



Click a chapter card to navigate. Use the PDF bookmarks to move between sections.

The AI Landscape

AI is evolving quickly; leaders must make choices with clarity, not hype.



The landscape never sits still.

New models. New tools.
New players. New rules.
Every week brings something new.



More information is not the same as better decisions.

Leaders must cut through the noise, focus on what matters, and take responsibility for the impact.



Responsible AI starts with clarity.

Understand the change.
Weigh the trade-offs.
Act with purpose.
Earn trust.

Lead with Clarity



1. See the Change

Stay informed about models, tools, markets, and regulations. Understand the pace and direction.



2. Separate Signal from Noise

Filter hype. Look for patterns, evidence, and real business impact.



3. Know What Matters

Focus on your strategy, people, data, risks, and values. Context makes the difference.



4. Decide Deliberately

Make choices with transparency, accountability, and a long-term view of impact.



Key Takeaway

- ✓ Rapid change increases the need for judgment.
- ✓ The right response is not fear or frenzy—it is disciplined decision-making.



The Decision Maker's Playbook

A simple checklist for responsible AI decisions that create real-world impact.

You don't need perfect information—just a clear process, trusted evidence, and a commitment to do what's right.

Lead with purpose. Decide with confidence.

BEFORE YOU SAY YES

- ☒  Clarify the Problem
- ☒  Name the Owner
- ☒  Test Assumptions
- ☒  Verify Evidence
- ☒  Define Guardrails
- ☒  Plan Monitoring
- ☒  Train the People
- ☒  Review and Adapt

THE GOAL

Make the right choice with the right knowledge—then keep learning.



KEY TAKEAWAYS



Responsible AI is a continuous practice.



Trust is earned through verification.



Good decisions improve with education and evaluation.



Different roles. Shared responsibility. **Better AI decisions.**

Thank you for being a leader who chooses what's right—for today and the future.

Start with Better Questions

The right questions shape the right path.

1

What could we be doing with AI?



Explore broadly. Look across the business to find opportunities to create value with AI.

2

What can we be successful at doing with AI?



Focus where we can win. Match opportunities to our strengths, data, and capabilities.

3

What organizational changes do we need to make?



Align people, processes, and policies to enable responsible and sustainable AI impact.



Good questions help avoid bad decisions.



risk.



They expose assumptions.



They open the door to better outcomes.

Better questions lead to better AI decisions.



Key Takeaways



Better questions create better options.



Questions should be tied to evidence.



Questions should reveal what must change.

Trust but Verify

Great AI starts with trust.
Great AI lasts with verification.

We don't rely on hope—we
validate, measure, and
continuously check
our AI decisions.



OUR VERIFICATION PROCESS



WHAT TO VERIFY



VERIFICATION QUESTIONS

- What evidence supports the claim?
- What would disprove it?
- How will we monitor failure?
- Who can stop it?



KEY TAKEAWAYS

- Trust begins with evidence.
- Every AI use case needs a way to test, monitor, and stop.

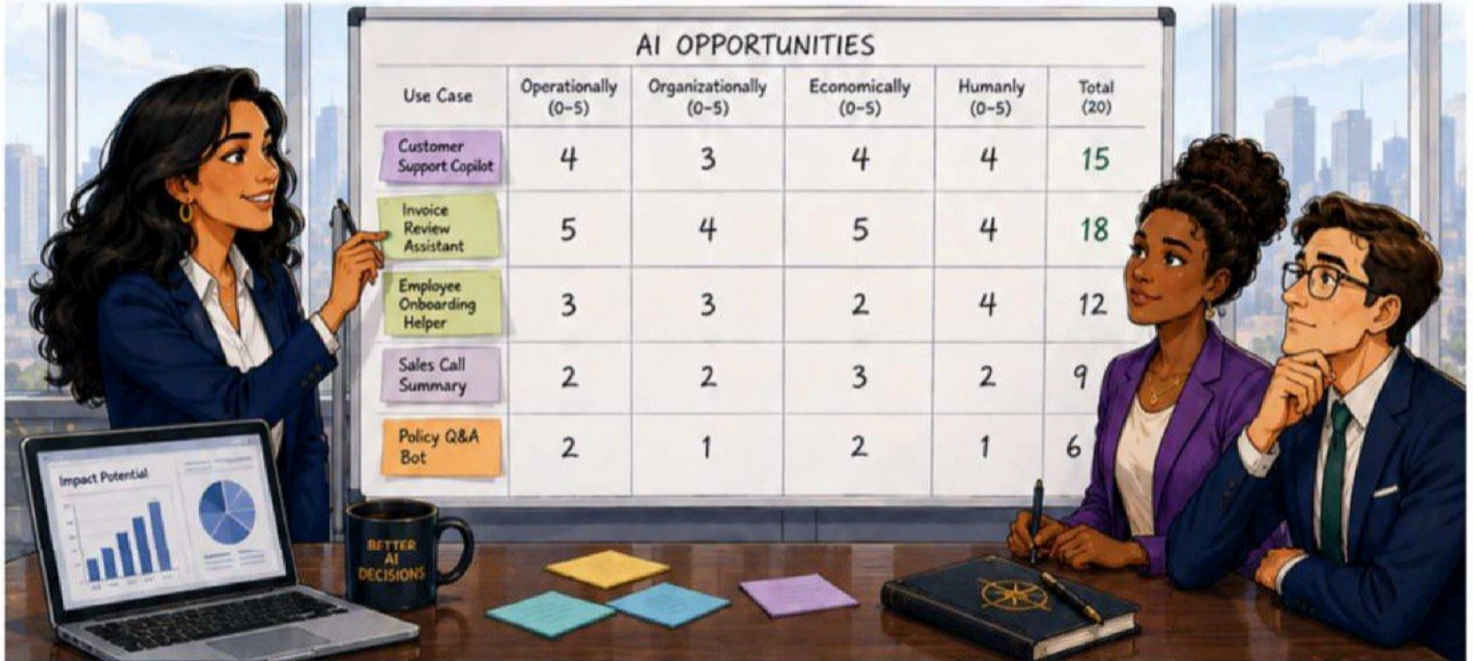
Choose Winnable Use Cases

Pick use cases you can operate, govern, and sustain.



Start small. Prove value. Scale with confidence.

Focus on opportunities that meet the right conditions today—so your team can deliver impact quickly and build trust for the future.



What Must Be True



Operationally
monitor, detect,
rollback.



Organizationally
clear owner, trained
users, workable
process.



Economically
sustainable cost
and real value.



Humanly
reduced friction,
not hidden labor.

Scorecard: Is It Winnable?

Scoring Guide (0-5)

0 = Not at all
1 = Very limited
2 = Limited
3 = Moderate
4 = Strong
5 = Excellent

Winnable Bet ✓

- ✓ Strong across most dimensions (total ≥ 15)
- ✓ Clear owner and path to delivery
- ✓ Measurable value within 90 days
- ✓ Low to moderate risk
- ✓ Builds trust and capability

Not Ready Yet ✗

- ✗ Low scores in any dimension
- ✗ Unclear ownership or process
- ✗ High cost, uncertain value
- ✗ High risk or compliance gaps
- ✗ Requires more data, tools, or change



Key Takeaway

- ✓ Not every idea is ready to scale.
- ✓ Choose the use cases you can operate, govern, and sustain.

Measure What Matters

Great decisions start with the right measures.

Responsible AI success is more than numbers. It's value for the business, for people, and for the future.

We measure what matters—across today, people, and tomorrow.



RoX Return on Everything that Matters

ROI

ROI —
Return on Investment

Does it create measurable financial value?

FINANCIAL IMPACT



\$ VALUE CREATED



Balanced Decision



A strong initiative should create value now, help people, and reduce the cost of future change.

ROE

ROE —
Return on Employee

Does it help people work better, with less friction?

EMPLOYEE EXPERIENCE

- ☒ Time saved
- ☒ Fewer handoffs
- ☒ Less rework
- ☒ More focus



ROF

ROF —
Return on Future

Does it improve adaptability, governability, and readiness?

FUTURE READINESS



ADAPTABILITY



GOVERNABILITY



READINESS



Key Takeaways



Do not measure only cost savings.



Hidden labor destroys value.



Future readiness is a real return.

Align the Organization

Strong outcomes start with strong alignment.

These four pillars work together to make Responsible AI real.

Alignment turns intent into impact—at every level.



Strategy

Align AI to business goals.



Structure

Create clear ownership and a workable operating model.



Capability

Build skills, tools, and support.



Legal & Compliance

Embed guardrails, accountability, and auditability.

Why Alignment Matters



Misalignment creates duplication.



Weak ownership creates risk.



Good governance enables scale.



Responsible AI is an organizational design challenge—not just a technology project.

Evaluate Continuously

Stay aware. Detect early.
Act fast. Improve always.

Responsible AI doesn't end at launch. Continuous evaluation helps you catch issues early, adapt with confidence, and keep people at the center.

Monitoring today protects people tomorrow.



VENDOR UPDATES

- 3 Updates Available
- ☒ Model Version 2.4.1 Security patch **New**
 - ☒ Data Source Refresh Schema change **Review**
 - ☒ Hosting Region Update Policy change **Review**

INCIDENTS

Open: **2** | In Review: **1** | Resolved (7d): **5**

Recent Incident
Hallucinated output in customer email draft

High

REGULATORY CHANGE

EU AI Act guidance updated
High-risk system documentation requirements expanded.
Effective: July 1

Monitor across the system, not just the model.

What to Watch

- Model drift
- Bias and safety
- Usage and outcomes
- Vendor changes
- Regulatory signals

Continuous monitoring replaces one-time approval.

The Monitoring Cycle



1. OBSERVE

Continuously collect signals from models, systems, and users.



2. MEASURE

Track metrics and compare against baselines and thresholds.



3. REVIEW

Investigate findings with context and stakeholder perspectives.



4. RESPOND

Mitigate issues, communicate clearly, and document actions.



5. IMPROVE

Update models, processes, and policies to reduce future risk.



Key Takeaway

In a rapidly changing AI landscape, yesterday's answer may not be safe tomorrow.

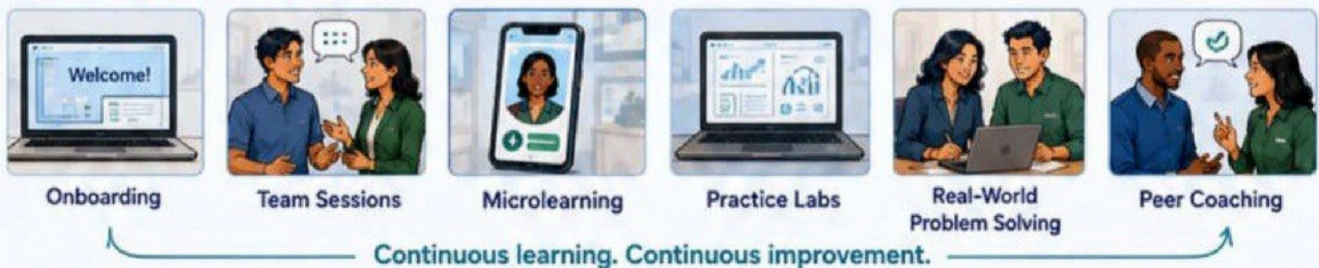
Educate Continuously

Learning never stops. Responsible AI grows with us.

Amara leads a learning session for the whole organization. Together, we build the knowledge and confidence to use AI responsibly—today and every day.



Learning happens everywhere.



Key Takeaway

- Responsible AI requires organization-wide learning.
- Education is not a one-time event—it is a capability.

Pilot Safely

Start small. Learn fast. Control risk.

Pilots help us validate value, refine our approach, and strengthen controls before we scale.

A good pilot proves value and tests controls.

Responsible AI Pilot Charter

Pilot Example

Customer Support Summarizer

	Problem to solve	Agents spend too much time reading long customer threads.
	Users	Customer support agents
	Success metrics	Summary accuracy, time saved, agent satisfaction
	Guardrails	No personal data exposure. No policy or financial advice.
	Human review	Agents review and approve summaries before use.
	Rollback plan	Disable feature and revert if issues are detected.

What a good pilot includes

- Clear objective
- Named owner
- Limited scope
- Measurable outcomes
- Human oversight
- Monitoring plan

A smart pilot path.

1

Define



- Clarify the problem and goal.
- Identify users and success metrics.

2

Prepare



- Set guardrails and data rules.
- Build and test the pilot.

3

Launch



- Start with a small group.
- Provide training and guidance.

4

Review



- Monitor results and user feedback.
- Check accuracy and guardrails.

5

Decide



- Adjust, expand, or stop.
- Document lessons and next steps.

EXAMPLE



Example:

A support summarizer pilot is tested with 25 agents, weekly reviews, accuracy checks, and a rollback option.

- 25 agents
- Weekly reviews
- Accuracy checks
- Rollback option



Key Takeaways



Pilot to learn—
not to overpromise.



Small, controlled
pilots build trust.

Keep Humans in the Loop

AI can assist. People remain accountable.



Example Use Case	MEETING SUMMARIZER (LOW RISK)	INVOICE REVIEW ASSISTANT (MEDIUM RISK)	POLICY ELIGIBILITY RECOMMENDATION (HIGH RISK)
	AI drafts a summary of a team meeting.	AI flags unusual items and recommends next steps.	AI recommends eligibility based on customer data.
AI Role	Drafts and summarizes information.	Analyzes and recommends based on patterns.	Evaluates and recommends outcomes.
Risk Level	Low	Medium	High
Required Oversight	AI ASSISTS No human review required for publication.	HUMAN REVIEW A person reviews before action is taken.	HUMAN APPROVAL / DECISION A person approves or makes the final decision.

Example

A drafted customer reply can be suggested by AI, but a human agent must approve before sending.



AI Suggestion (Draft)

Thank you for reaching out. We're sorry for the delay. We're looking into this and will follow up shortly.



Human Agent Review

Looks good. I'll make a small edit and approve.



Approved & Sent



Human oversight is a design choice, not an afterthought.



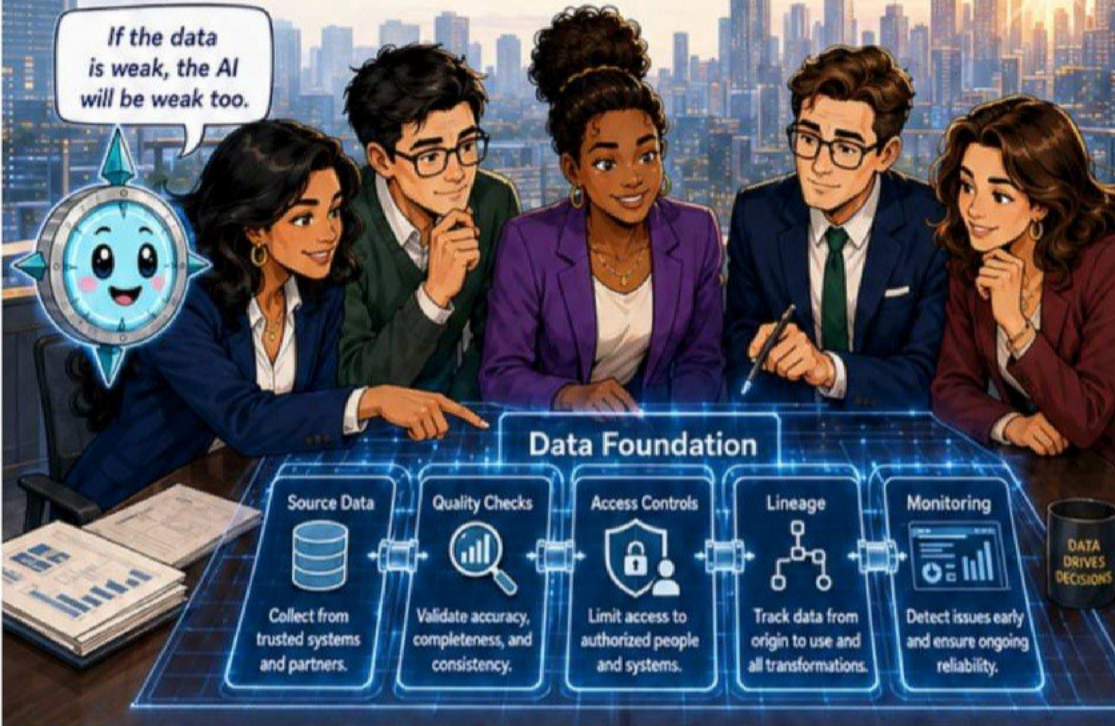
AI supports judgment—it does not replace accountability.



Strengthen the Data Foundation

Responsible AI depends on responsible data.

If the data is weak, the AI will be weak too.



What strong data practice looks like

-  Accurate and complete
-  Representative
-  Permissioned access
-  Protected sensitive data
-  Clear lineage
-  Retention rules

Data Foundation

Source Data



Collect from trusted systems and partners.

Quality Checks



Validate accuracy, completeness, and consistency.

Access Controls



Limit access to authorized people and systems.

Lineage



Track data from origin to use and all transformations.

Monitoring



Detect issues early and ensure ongoing reliability.

DATA DRIVES DECISIONS

Build a Strong Data Foundation

1 Data quality



High-quality data is accurate, complete, consistent, and timely. Poor quality leads to poor AI outcomes.

2 Data governance



Clear policies, roles, and standards ensure data is used responsibly and aligned with business goals.

3 Data privacy



Protect personal and sensitive data with privacy-by-design, minimization, encryption, and de-identification.

4 Data observability



Continuously monitor data health, quality drift, anomalies, and usage to maintain trust and performance.

Example:

An invoice review assistant performs well only when invoice samples are current, representative, and properly labeled.



Collect current invoice samples



Ensure representative across vendors and regions



Label accurately (e.g., line items, taxes, totals)



Pass quality checks



Monitor and refresh regularly



Better performance.
More reliable results.



Key Takeaway



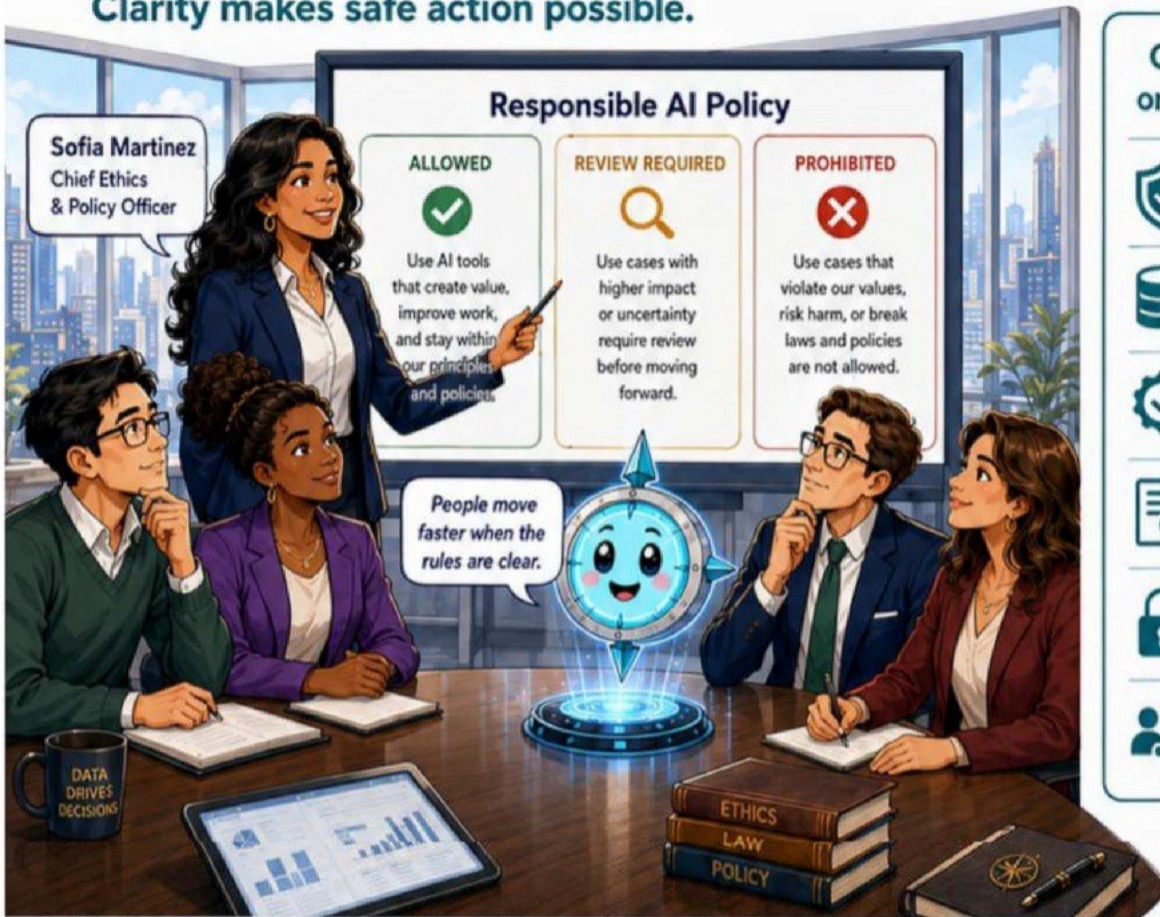
Better data creates better decisions.



Trust but verify the data before trusting the model.

Set Guardrails and Policy

Clarity makes safe action possible.



Guardrails every organization needs



Acceptable use



Data handling



Model approval



Logging and traceability



IP and confidentiality



Escalation path

Our Policy Ladder: Match Risk to the Right Level of Oversight



LEVEL 1 ALLOWED

Low risk, high value.
Use with confidence
within our principles.

Examples

- Brainstorming ideas
- Summarizing internal documents
- Drafting internal communications



LEVEL 2 REVIEW REQUIRED

Moderate to high impact
or uncertainty.
Human review required
before use.

Examples

- Customer-facing content
- Recommendations or decisions that affect people
- External sharing or publishing



LEVEL 3 PROHIBITED

High risk or against
our values or laws.
Do not use.

Examples

- Unlawful discrimination or bias
- Unsafe automation of decisions
- Confidential data leakage



Example

An internal note summarizer may be allowed, while an external policy chatbot may require legal review.



Key Takeaways



Good guardrails
reduce confusion.



Policy turns principles
into daily practice.

Decide: Build, Buy, or Partner

Choose the path you can govern and sustain.

There's no one-size-fits-all answer.

Build vs Buy vs Partner

	BUILD	BUY	PARTNER
Control	●●●●●	●●●●●	●●●●●
Speed to Value	●●●●●	●●●●●	●●●●●
Total Cost	●●●●●	●●●●●	●●●●●
Customization	●●●●●	●●●●●	●●●●●
Risk	●●●●●	●●●●●	●●●●●

Match the choice to your goals, risks, and capabilities.

What to compare

-  Speed to value
-  Data fit
-  Security
-  Integration effort
-  Cost of ownership
-  Vendor risk



BUILD

We develop the solution in-house.

Pros

- ✓ Full control and customization
- ✓ Deep alignment with needs
- ✓ Protects unique data assets

Cons

- ✗ Slower to value
- ✗ Higher upfront investment
- ✗ Requires specialized talent

Relative Score (example)



BUY

We purchase a ready-made solution.

Pros

- ✓ Fastest time to value
- ✓ Lower upfront cost
- ✓ Proven capabilities

Cons

- ✗ Less customization
- ✗ Data may not fit perfectly
- ✗ Ongoing vendor dependency

Relative Score (example)



PARTNER

We co-develop or work with a partner.

Pros

- ✓ Shared expertise and speed
- ✓ More flexibility than buy
- ✓ Access to partner IP

Cons

- ✗ Coordination complexity
- ✗ Shared control
- ✗ Potential IP or roadblock risk

Relative Score (example)



Example



A customer support copilot may be bought from a vendor, but it still needs local controls, testing, integration, and training.



Key Takeaway



Do not outsource accountability.



Even when you buy AI, you still own the outcome.

Prepare for Failure

The safest teams rehearse what could go wrong.

When something goes wrong



Detect

Spot issues early with monitoring and signals.



Contain

Stop or limit impact quickly.



Notify

Inform the right people internally and externally.



Investigate

Find root causes with data and evidence.



Fix

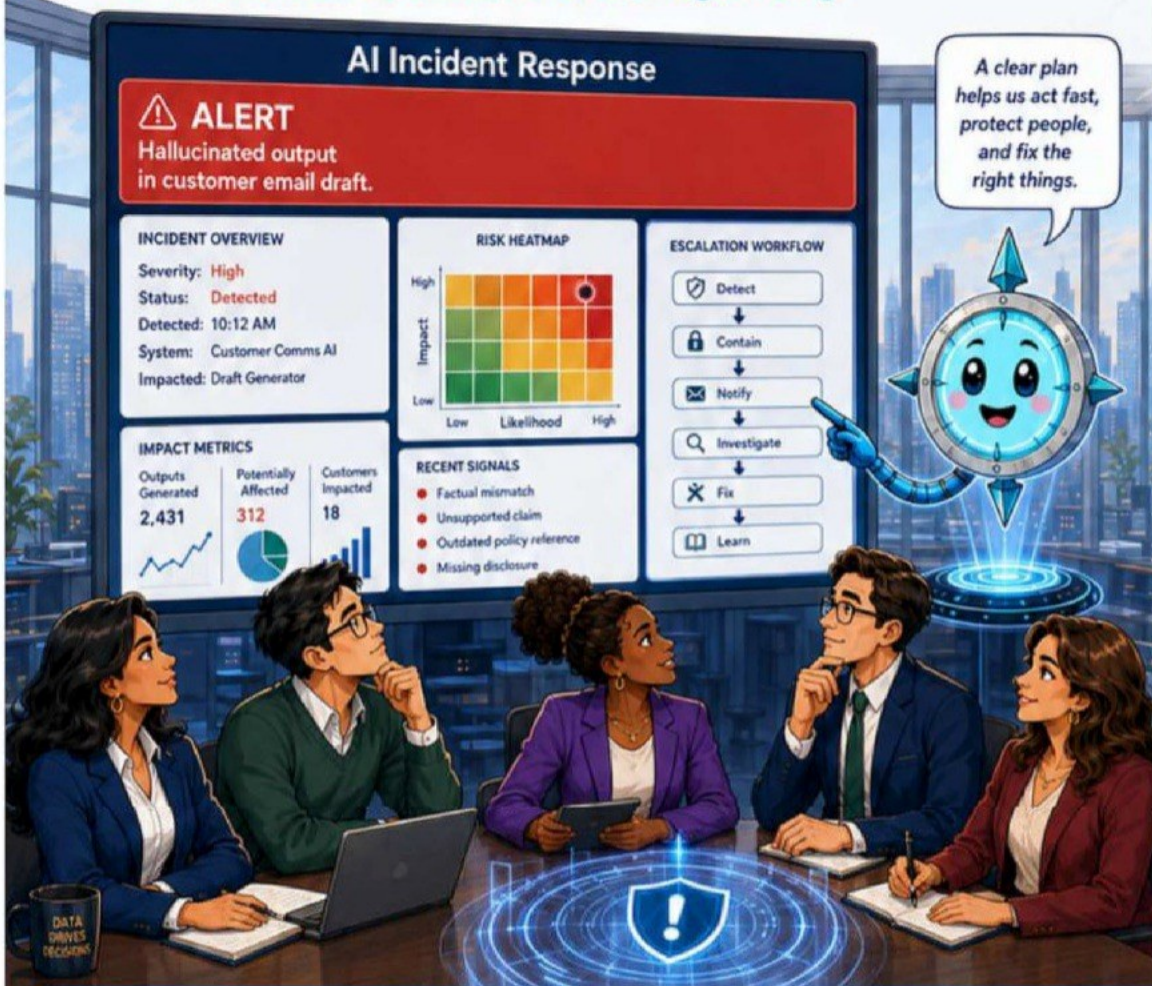
Address the issue and validate the solution.



Learn

Capture lessons and improve controls.

A clear plan helps us act fast, protect people, and fix the right things.



Our 6-Step Incident Response Chain

1



Detect

Automated and human signals flag unusual or harmful behavior.

2



Contain

Pause, throttle, or disable the affected capability or output.

3



Notify

Alert responders, leaders, and customers based on severity and policy.

4



Investigate

Analyze logs, data, prompts, and model behavior to find causes.

5



Fix

Implement mitigations, patch, or update models and controls.

6

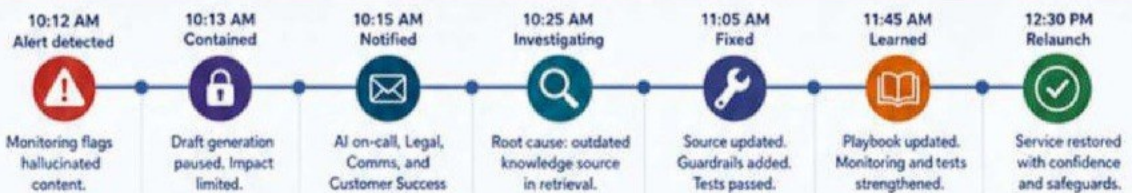


Learn

Document findings, update playbooks, and strengthen prevention.

Example:

A customer-facing draft is paused, affected outputs are reviewed, the root cause is investigated, and controls are updated before relaunch.



Key Takeaways



Every AI use case needs a stop button.



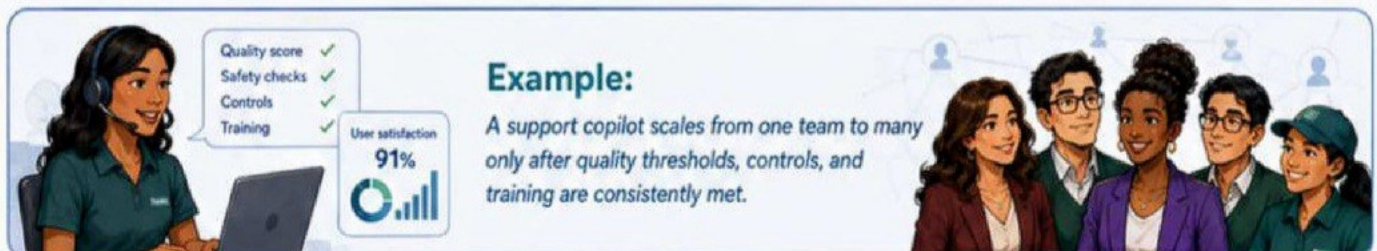
Rollback is part of responsible design.

Scale Responsibly

Scale only when value, trust, and readiness are real.



The Responsible Scale Maturity Ladder



Scale is earned.

A pilot becomes a platform only when the organization is ready.

Examples in Action

Concrete examples make better choices easier.

EXAMPLE 1



Customer Support Summarizer



PROCEED WITH PILOT

- Clear problem and scope
- High-quality transcripts
- Low risk to customers
- Easy to monitor outcomes
- Strong efficiency gains

EXAMPLE 2



Employee Onboarding Helper



PROCEED WITH TRAINING AND CONTROLS

- Helpful, but needs oversight
- Some data gaps to close
- Moderate risk if inaccurate
- Human review required
- Long-term value is high

EXAMPLE 3



Policy Q&A Bot



NOT READY YET

- Policies are complex and changing
- High risk of incorrect or harmful advice
- Hard to verify answers
- Limited ability to monitor
- Low tolerance for errors

Why these decisions differ



Problem clarity



Data readiness



Human oversight



Risk level



Ability to monitor



Value proof



Responsible decisions are specific, not generic.

Comparing the Examples

Criteria	Customer Support Summarizer	Employee Onboarding Helper	Policy Q&A Bot
Problem clarity	High	Medium	Low
Data readiness	High	Medium	Low
Human oversight	Light touch	Required	Heavy (not yet feasible)
Risk level	Low	Moderate	High
Ability to monitor	Easy	Moderate	Difficult
Value proof	Strong	High (over time)	Unproven
Overall decision	READY NOW (PILOT)	NEEDS MORE WORK (TRAINING & CONTROLS)	PAUSE (NOT READY YET)



Key Insight

The same idea can have different answers depending on context, readiness, and risk.



Key Takeaway



Examples help teams see the pattern.



Not every idea deserves the same answer.



Ask Before You Approve

A simple checklist can improve major decisions.

Before You Say Yes

- ☒ 1 Clarify the problem
- ☒ 2 Name the owner
- ☒ 3 Test assumptions
- ☒ 4 Verify evidence
- ☒ 5 Define guardrails
- ☒ 6 Plan monitoring
- ☒ 7 Train the people
- ☒ 8 Review and adapt

The goal



Make the right choice with the right knowledge—then keep learning.



RESPONSIBLE AI

1

Clarify the problem



Be clear on the real problem we're solving.

2

Name the owner



Assign a person or team who is accountable.

3

Test assumptions



Challenge what we think is true. Look for bias.

4

Verify evidence



Use reliable data and credible sources.

5

Define guardrails



Set ethical boundaries and risk limits.

6

Plan monitoring



Decide what to track, how, and how often.

7

Train the people



Make sure users understand how to use AI well.

8

Review and adapt



Evaluate results. Learn and improve.



Example:

If there is no owner, no monitoring plan, or no human oversight, the answer is not yes yet.



No Owner



No Monitoring Plan



No Human Oversight

=



Not Yes Yet



Responsible AI is a continuous practice.



Good decisions improve with education and evaluation.

Lead with Purpose

Different roles. Shared responsibility. Better AI decisions.

Our Responsible AI Journey



What leaders should remember



Clarity
over hype



Evidence
over assumption



Oversight
over blind
automation



Learning
over one-time
training



Monitoring
over one-time
approval



Shared
responsibility
across the
organization

The Responsible AI Promise



We ask
better questions.



Great questions
create better
AI outcomes.



We verify
before we trust.



Evidence and oversight
build confidence
in every decision.



We educate and
evaluate
continuously.



Learning and measurement
drive improvement
that never stops.



We keep humans
accountable.



Human judgment
ensures ethics, fairness,
and responsibility.



We scale only
what we can
govern.



Sustainable scale
creates lasting value
for people and business.



Make the right choice with the right knowledge.

Then keep learning.

Responsible AI Leadership Checklist

A practical conversation guide for senior leaders before scaling AI.

Strategic fit

- ☐ What business outcome are we improving?
- ☐ What evidence proves demand is real?
- ☐ What would falsify this idea?
- ☐ Which RoX lens matters most here?

Operational readiness

- ☐ Can we monitor quality, failures, and drift?
- ☐ Who owns the AI outcome and the human decision?
- ☐ What is the rollback or kill-switch plan?
- ☐ Can we support the tool at real usage volumes?



Human impact

- ☐ Does AI reduce cognitive load or add hidden work?
- ☐ Have employees been trained for new responsibilities?
- ☐ Does the workflow preserve human judgment?
- ☐ Do teams understand the boundaries of use?

Trust and compliance

- ☐ Have privacy, security, IP, and bias risks been reviewed?
- ☐ Are third-party dependencies governed?
- ☐ Can decisions be audited and explained?
- ☐ Are controls embedded early, not bolted on later?

Executive Takeaways

Key lessons for leading **responsible AI** that delivers lasting impact.



1. Start with questions, not hype.

Curiosity and clarity drive better AI decisions.



2. Measure value across ROI, ROE, and ROF.

Define and track impact that matters.



3. Align strategy, structure, capability, and compliance.

Integrated alignment turns plans into performance.



4. Design governance early.

Trust is built in the architecture, not added later.



5. Use AI to augment people.

Empower your teams. Elevate human potential.



6. Scale only what you can operate and trust.

Sustainable impact grows with confidence.



Lead with purpose.
Build with responsibility.
Shape a better future with AI.



Thank you for being part of the journey toward **trusted AI** that creates value for people and the world.

About OACA

Open Alliance for Cloud Adoption

Mission

This graphic novel is not meant to be read once from cover to cover and then shelved. Instead, think of it as a companion you can return to throughout your organization's AI journey. The graphic novel format provides an accessible, engaging, and memorable way to communicate complex AI concepts across different roles, functions, and levels of technical expertise, helping create a shared understanding of responsible AI adoption across the organization.

As new opportunities, challenges, and decisions arise, return to the chapters that are most relevant, reflect on the ideas presented, and consider what may have changed in your perspective. The goal is not to provide every answer, but to help you develop the context, confidence, and critical thinking needed to make informed decisions that shape AI adoption.

As you revisit these ideas, this storybook will help you:

- Translate responsible AI principles into practical decision steps;
- Align business, legal, technology, and operational stakeholders;
- Encourage measurement, evaluation, education, and human accountability; and
- Support sustainable scaling with evidence, guardrails, and governance.

Disclaimer

In alignment with Responsible AI principles, this graphic novel was created with the assistance of AI using 'OACA organization' published white papers on AI adoption as its primary reference material. While the format and visuals were AI-generated, the core concepts, recommendations, and themes are based on the referenced source documents and have been adapted for educational and illustrative purposes.

Website and Reference

OACA Website: <https://oaca-project.org/>

Documentation: <https://oaca-project.org/library/>

Legal statement

This guidance is intended for educational and informational purposes only. It does not constitute legal, regulatory, privacy, security, compliance, risk, or professional advice. Each organization remains responsible for assessing applicability, meeting its own obligations, and obtaining appropriate expert review before implementation or reliance on any recommendation.

Legal Notice

Thinking With AI – A Storybook for Better Decisions, Responsible Leadership, and Enterprise AI

© 2026 Open Alliance for Cloud Adoption – A Linux Foundation Initiative, Inc. ALL RIGHTS RESERVED.

This “Thinking With AI – A Storybook for Better Decisions, Responsible Leadership, and Enterprise AI” is proprietary to the Open Alliance for Cloud Adoption – A Linux Foundation Initiative (the “Alliance”) and/or its successors and assigns.

This OACA document is licensed under the Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License (CC BY-NC-ND 4.0). To view a copy of the license, visit <https://creativecommons.org/licenses/by-nc-nd/4.0/>

If any derivatives of this document are published, the following statement must be identified: “This document is based on the Thinking With AI – A Storybook for Better Decisions, Responsible Leadership, and Enterprise AI document created by the Open Alliance for Cloud Adoption – A Linux Foundation Initiative (OACA) but may contain changes to the original OACA document which have not been reviewed or approved by the OACA.”

LEGAL DISCLAIMER

THIS DOCUMENT AND THE INFORMATION CONTAINED HEREIN IS PROVIDED ON AN “AS IS” BASIS. TO THE MAXIMUM EXTENT PERMITTED BY APPLICABLE LAW, THE ALLIANCE (ALONG WITH THE CONTRIBUTORS TO THIS DOCUMENT) HEREBY DISCLAIM ALL REPRESENTATIONS, WARRANTIES AND/OR COVENANTS, EITHER EXPRESS OR IMPLIED, STATUTORY OR AT COMMON LAW, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE, TITLE, VALIDITY, AND/OR NONINFRINGEMENT. THE INFORMATION CONTAINED IN THIS DOCUMENT IS FOR INFORMATIONAL PURPOSES ONLY AND THE ALLIANCE MAKES NO REPRESENTATIONS, WARRANTIES AND/OR COVENANTS AS TO THE RESULTS THAT MAY BE OBTAINED FROM THE USE OF, OR RELIANCE ON, ANY INFORMATION SET FORTH IN THIS DOCUMENT, OR AS TO THE ACCURACY OR RELIABILITY OF SUCH INFORMATION. EXCEPT AS OTHERWISE EXPRESSLY SET FORTH HEREIN, NOTHING CONTAINED IN THIS DOCUMENT SHALL BE DEEMED AS GRANTING YOU ANY KIND OF LICENSE IN THE DOCUMENT, OR ANY OF ITS CONTENTS, EITHER EXPRESSLY OR IMPLIEDLY, OR TO ANY INTELLECTUAL PROPERTY OWNED OR CONTROLLED BY THE ALLIANCE, INCLUDING, WITHOUT LIMITATION, ANY TRADEMARKS OF THE ALLIANCE.

TRADEMARKS

OPEN ALLIANCE FOR CLOUD ADOPTIONSM, OACASM, and the OPEN ALLIANCE FOR CLOUD ADOPTION logo[®] are trade names, trademarks, and/or service marks (collectively “Marks”) owned by Open Alliance for Cloud Adoption – A Linux Foundation Initiative, Inc. and all rights are reserved therein. Unauthorized use is strictly prohibited. This document does not grant any user of this document any rights to use any of the OACA’s Marks. All other service marks, trademarks, and trade names reference herein are those of their respective owners.

PRODUCTION NOTE

In alignment with responsible AI principles, the format and illustrations in this document were generated with AI assistance, using OACA’s published white papers on AI adoption as the source material. The core concepts, recommendations, and themes are drawn from those referenced source documents and have been adapted for educational and illustrative purposes.

INTENDED USE

This document is intended for educational and informational purposes only. It does not constitute legal, regulatory, privacy, security, compliance, risk, or professional advice. Each organization remains responsible for assessing applicability, meeting its own obligations, and obtaining appropriate expert review before implementation or reliance on any recommendation.