

AI-Readiness for L&D

A practical system for turning AI access into safe, measurable performance.

AI adoption is not AI readiness. Readiness means people can choose, use, verify, and improve AI-supported work within clear boundaries, with evidence that performance changed.

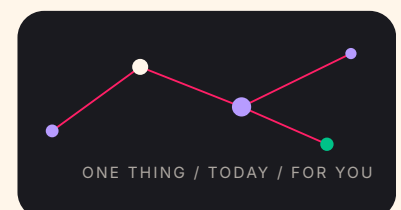
FOR L&D, HR, AND WORKFORCE TRANSFORMATION LEADERS

Playbook

10 pages

17-min read

8 source documents



READINESS MODEL

Adoption is not readiness

Tool access shows availability. Readiness requires aligned work, roles, safeguards, capability, learning support, and measurement.

88%

of surveyed organizations used AI in at least one function

Stanford reports 2025 adoption. Survey evidence is directional.

AI1: Stanford Institute for Human-Centered Artificial Intelligence, 2026

79%

reported regular generative AI use

Organizational use does not establish workforce proficiency.

AI1: Stanford Institute for Human-Centered Artificial Intelligence, 2026

26%

of AI users said leadership was aligned

Microsoft Work Trend Index respondent finding.

AI2: Microsoft WorkLab, 2026

67% vs 32%

organizational versus individual feature importance

Reported model contribution to AI impact, not a causal split.

AI2: Microsoft WorkLab, 2026

READINESS DIMENSION	QUESTION YOUR EVIDENCE MUST ANSWER
Business strategy	Which outcome and workflow justify AI use?
Workflow and roles	Which tasks change, and who remains accountable?
Governance and risk	Which uses are allowed, reviewed, logged, or prohibited?
Workforce capability	Can each role direct, verify, protect, and escalate?
Learning enablement	Can people practise with realistic work and feedback?
Measurement	Can you separate access, behavior, quality, risk, and value?

Omie readiness model. Score each dimension from 0 to 3 using local evidence.

OMIE SELF-ASSESSMENT

six dimension scores × 0 to 3 = 0 to 18

Use 0 to 6 as exposed, 7 to 11 as exploring, 12 to 15 as operating, and 16 to 18 as scaling. These are Omie working ranges, not external benchmarks.

DECISION Do not fund broad training until the organization can name the work, guardrails, capability, and evidence it needs.

TASK MAP

Map work before you train people

Occupational exposure is not the same as job loss. Break roles into tasks, decisions, inputs, outputs, and failure modes before choosing learning.

1 in 4

workers are in an occupation with some generative AI exposure

ILO global estimate. Transformation is more likely than full redundancy.

AI3: International Labour Organization, 2025

TASK FIELD	WHAT TO RECORD	WHY IT MATTERS
Trigger	Event that starts the task	Keeps the map grounded in workflow
Inputs	Data, context, policy, and source quality	Reveals privacy and provenance risk
Decision	Judgment the person must make	Protects human accountability
AI contribution	Draft, retrieve, classify, summarize, or recommend	Defines the capability to practise
Output	Artifact or action and its audience	Makes quality observable
Failure cost	Customer, legal, safety, fairness, or rework impact	Sets review intensity

Value and risk screen

Task, role, frequency, and current cycle time

Current quality measure and recurring failure

Proposed AI contribution and expected user decision

Data sensitivity, affected people, and error consequence

Required review, escalation, and audit evidence

Pilot, redesign, defer, or prohibit decision and owner

Omie working template. Prioritize high-value, bounded tasks with reviewable outputs.

DECISION Train the changed judgment inside a task. A generic prompt course cannot resolve unclear accountability or unsafe workflow design.

ROLE CAPABILITY

Define what competent AI use looks like by role

A shared capability language makes expectations teachable while role-level evidence keeps the standard relevant to the work.

CAPABILITY	OBSERVABLE ROLE BEHAVIOR
Understand	Explains the use case, limits, and when AI is inappropriate
Direct	Provides the goal, context, constraints, examples, and output criteria
Verify	Checks claims, calculations, sources, completeness, and downstream impact
Protect	Handles data, privacy, intellectual property, and access within policy
Integrate	Places AI inside the workflow with clear review and ownership
Improve	Uses evidence to refine instructions, controls, process, or tool choice

OMIE LEVEL	EVIDENCE STANDARD
0: Unexposed	Cannot yet identify the allowed workflow or basic risk
1: Aware	Explains the concept but needs a guided example
2: Assisted	Completes a bounded task with review and support
3: Independent	Handles routine variation and documents verification
4: Enabling	Improves the workflow and coaches others within policy

Omie local 0 to 4 rubric. Set the required level separately for each capability and role.

CURRENT LEGAL CONTEXT

Treat literacy as an operating responsibility

The Council adopted the 2026 Digital Omnibus legislative text on 29 June. Its amended Article 4 would require measures supporting AI literacy development without guaranteeing a specific level for every person. The Council said Official Journal publication would follow. Status checked 11 July 2026. Translate applicable obligations with legal and risk owners. This guide is not legal advice.

Evidence: AI6: European Parliament and Council (2026) / AI8: Council of the European Union (2026)

DECISION A role standard should say what the person does, under which conditions, and what evidence demonstrates safe independence.

RISK ARCHITECTURE

Turn governance into teachable behavior

Policy becomes useful when each risk has a moment of recognition, an expected action, a review control, and an escalation path.

RISK	LEARNING BEHAVIOR	CONTROL EVIDENCE
Confabulation	Verify material claims against approved sources	Source record and reviewer
Privacy	Classify data before entering or retrieving it	Approved environment and data log
Intellectual property	Check rights, attribution, and permitted reuse	Provenance and approval
Bias	Test affected groups and challenge proxy assumptions	Test cases and decision record
Information integrity	Inspect origin, manipulation, and audience risk	Authenticity and release check
Security	Reject suspicious instructions and minimize access	Security review and incident route
Human oversight	Know when to review, override, or stop	Named accountable decision maker

Risk categories are adapted from the NIST Generative AI Profile.

Evidence: AI7: National Institute of Standards and Technology (2024)

Omie default risk tiers

- 1 Tier 1: bounded**
Low-consequence internal assistance with approved data and normal review.
- 2 Tier 2: controlled**
Material external or employee impact requiring documented verification and owner approval.
- 3 Tier 3: restricted**
High-consequence, sensitive, or regulated use requiring specialist review, stronger controls, or prohibition.

Risk-to-learning card

Risk event and earliest visible signal

Required user action and prohibited shortcut

Review control, accountable owner, and evidence retained

Scenario that tests recognition, response, and escalation

Omie working template. Risk owners must approve the local tier and response.

DECISION Teach the decision and escalation behavior that makes a policy operational, then test it under realistic pressure.

PILOT PLAN

Build a role pilot across 90 days

A small pilot should establish a baseline, build shared foundations, practise one workflow, and decide whether evidence supports scaling.

39%

of core worker skills expected to change by 2030

Employer forecast reported by the World Economic Forum.

AI4: World Economic Forum, 2025

59 / 100

workers expected to need training by 2030

Eleven may not receive the training they need.

AI4: World Economic Forum, 2025

47%

of leaders ranked upskilling existing employees as a top strategy

Microsoft leader survey finding.

AI5: Microsoft WorkLab, 2025

51%

of managers expected AI training to become a key responsibility

Expectation for the next five years, not an observed outcome.

AI5: Microsoft WorkLab, 2025

Omie recommended pilot cadence

- 1 **Days 0 to 15: baseline**
Select one role and workflow, map risk, freeze measures, and collect current work samples.
- 2 **Days 16 to 30: foundations**
Teach policy, tool limits, data handling, verification, and the shared capability language.
- 3 **Days 31 to 60: workflow labs**
Run varied scenarios, compare artifacts to the rubric, and repair workflow friction.
- 4 **Days 61 to 90: coached pilot**
Apply in live work, monitor controls and outcomes, then decide stop, revise, scale, or study longer.

Pilot charter

Role, task, cohort, owner, and affected stakeholder

Baseline behavior, quality, time, risk, and volume

Allowed tool, data boundary, review, and escalation

Learning sequence, practice evidence, and manager support

Decision thresholds, confounders, and scale or stop date

Ninety days is an Omie implementation pattern, not a universal time-to-impact claim.

DECISION Pilot one role workflow deeply enough to observe behavior, quality, risk, and operating ownership before expanding access.

PRACTICE DESIGN

Practise judgment, not prompt memorization

Good practice recreates ambiguity, incomplete evidence, sensitive data, and review pressure so the learner must choose and explain a safe action.

ILLUSTRATIVE SCENARIO

A rushed workforce recommendation

A leader asks for an AI-generated summary of employee comments before a meeting. The file contains identifiers, the deadline is short, and the first output makes a claim unsupported by the source. Decide what to withhold, where to work, what to verify, how to communicate uncertainty, and when to escalate.

OMIE RUBRIC AREA	EVIDENCE IN THE RESPONSE
Task judgment	Chooses an appropriate AI contribution and preserves human ownership
Instruction quality	Supplies purpose, context, constraints, and acceptance criteria
Verification	Tests material claims, coverage, source fidelity, and calculation
Protection	Handles data, access, rights, and affected people within policy
Escalation	Recognizes uncertainty or impact that requires another owner

Score each area from 0 to 4 for an Omie local total of 20. Set local pass rules by risk tier.

Scenario loop

- 1 Attempt**
Let the learner decide and produce the work without hidden cues.
- 2 Inspect**
Compare the artifact, reasoning, and control evidence to the rubric.
- 3 Adjust**
Choose one changed behavior and explain why it matters.
- 4 Vary**
Repeat with different data, urgency, stakeholder impact, or model behavior.
- 5 Transfer**
Use the behavior in approved live work and review the evidence.

DECISION A competent user can explain the decision, detect a weak result, protect the boundary, and know when not to proceed.

MEASUREMENT

Measure absorption, not attendance

Use a connected measurement chain from access to behavior to work outcomes. Keep quality and risk beside speed so adoption cannot hide harm.

MEASURE	DEFINITION	INTERPRETATION
Activation	Eligible users completing the approved target action	Access and onboarding signal
Practice quality	Rubric points earned across scenario attempts	Learning and judgment signal
Transfer	Users demonstrating the behavior in approved live work	Workflow adoption signal
Cycle-time change	Baseline time minus pilot time, divided by baseline time	Efficiency signal with workload context
Quality change	Pilot quality minus baseline quality on the same rubric	Performance signal
Incident rate	Confirmed control failures per 1,000 approved uses	Risk signal requiring exposure context

Omie measurement dictionary. Confirm local definitions, windows, and owners before the pilot.

TRANSFER RATE

people with verified live-work evidence ÷ eligible pilot participants × 100

Define the evidence standard and observation window before launch. Do not substitute self-reported confidence for demonstrated behavior.

Interpretation guardrails

- Freeze the baseline, cohort, workflow, and quality rubric before training.
- Report numerator, denominator, exposure, missing data, and sample size.
- Track rework, review time, incidents, and exception handling beside cycle time.
- Separate tool effects from process, manager, staffing, and demand changes.
- Use finance-approved assumptions before assigning monetary value.

DECISION AI readiness improves only when safer behavior appears in real work and the resulting performance can withstand scrutiny.

FIRST 30 DAYS

Create the conditions for a credible pilot

The first month is for narrowing scope, assigning ownership, collecting baseline evidence, and making the learning and risk system ready.

Omie 30-day launch sequence

- 1 **Week 1: align**
Name the business decision, workflow, role, sponsor, accountable owner, and non-goals.
- 2 **Week 2: map**
Document tasks, data, affected people, risks, controls, and current work evidence.
- 3 **Week 3: design**
Set the capability standard, scenarios, rubric, support path, and measurement dictionary.
- 4 **Week 4: approve**
Review policy, tool, data, pilot thresholds, communications, and the scale or stop decision.

Launch gate

- ☐ The workflow and accountable human decision are explicit.
- ☐ Tool, model, data, retention, and access boundaries are approved.
- ☐ Role capability and scenario evidence are observable.
- ☐ Managers can protect practice time and review live-work evidence.
- ☐ Baseline, quality, risk, value, and stop rules are frozen.
- ☐ Learners know how to question, report, appeal, and escalate.

ACTION REGISTER FIELD	ENTRY
Decision	What must be decided and by when
Evidence	Artifact, measure, source, and current confidence
Owner	Person accountable for action and approval
Risk	Failure mode, control, and escalation route
Status	Open, active, blocked, accepted, or stopped
Review	Date, attendees, and next decision

Omie working template. Keep the register with the operating team, not only L&D.

DECISION Readiness is a shared operating system. L&D can enable it, but business, technology, legal, risk, and managers must own their decisions.

SOURCES AND METHOD

Named evidence, visible caveats

Every external numeric claim in this guide points to one of these 2024–2026 sources. Forecasts, self-reported surveys, associations, and illustrative examples are labelled so they are not mistaken for causal proof or universal benchmarks.

AI1

AI Index Report 2026, Chapter 4: Economy

Stanford Institute for Human-Centered Artificial Intelligence / 2026

Section 4.3, including Figure 4.3.1. Organizational adoption is self-reported and directional.

hai.stanford.edu

AI5

Work Trend Index Annual Report: Executive Summary

Microsoft WorkLab / 2025

Leader and manager expectations. Survey findings are directional, not guaranteed outcomes.

cdn-dynmedia-1.microsoft.com

AI2

Work Trend Index Annual Report: Agents, Human Agency, and the Opportunity for Every Organization

Microsoft WorkLab / 2026

AI-user study across ten markets. Reported relationships and model importance are not causal estimates.

microsoft.com

AI6

Digital Omnibus on AI, PE-CONS 30/26

European Parliament and Council / 2026

Adopted 29 June 2026. The Council notice said Official Journal publication would follow. Status checked 11 July 2026. Confirm the current text, applicability, and obligations with legal counsel.

data.consilium.europa.eu

AI3

Generative AI and Jobs: A 2025 Update

International Labour Organization / 2025

Exposure estimate. Occupational transformation is more likely than full redundancy.

ilo.org

AI7

Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile

National Institute of Standards and Technology / 2024

Sections 2 and 3 describe generative AI risks and risk-management actions.

nvlpubs.nist.gov

AI4

The Future of Jobs Report 2025

World Economic Forum / 2025

Employer forecasts in Section 3.1 and the workforce training outlook.

reports.weforum.org

AI8

Artificial Intelligence: Council Gives Final Green Light to Simplify and Streamline Rules

Council of the European Union / 2026

Adoption notice dated 29 June 2026. It said Official Journal publication and entry into force would follow.

consilium.europa.eu

HOW TO USE THIS GUIDE

Use the worksheets as decision scaffolds, then replace every recommended default and illustrative value with your own verified data. Recheck time-sensitive claims before publishing them outside your organization.