



WHAT MUST HUMANS REMAIN RESPONSIBLE FOR?

A Human Responsibility Framework for Learning & Development in the Age of AI

If AI can increasingly perform cognitive work itself, L&D; eventually has to ask not only “What should people learn?” but “What must humans remain responsible for?”

DR. RAVINDER TULSIANI
Enterprise Capability Executive
rtulsiani.com

ABOUT THIS REPORT

A strategy paper for L&D; leaders confronting a new division of labour

This report examines a question that becomes increasingly important as AI moves from assisting cognitive work toward performing more of it: what must humans remain responsible for?

The report combines current research on AI capability, work transformation, augmentation, expertise and human oversight with a practical operating framework for Learning & Development. It does not assume that every task AI can perform should be automated, nor that human participation is inherently superior. It asks where responsibility should sit and what capability is required to make that responsibility real.

Evidence and synthesis

Published research and governance guidance are identified as evidence. The Human Responsibility Framework, Oversight Paradox, Minimum Viable Human Capability, Human Responsibility Canvas and the principle “Automate capability freely. Delegate responsibility deliberately.” are synthesis developed for this report.

A practical limitation

The pace of AI capability improvement is uncertain. The framework is therefore designed to survive capability changes by separating what AI can do from what organizations decide humans should remain responsible for.

CONTENTS

The argument in one line

CORE ARGUMENT

As AI performs more cognitive work, L&D; must shift from asking only what people should learn to determining what humans must continue to understand, verify, decide, govern and own.

A strategy paper for L&D; leaders confronting a new division of labour	2
AI is forcing L&D; to answer a larger question	5
Three findings should change how L&D; approaches AI	6
The frontier of economically useful AI capability is expanding	7
Stop looking for skills AI will never possess	8
Technical capability does not determine authority	9
Seven responsibilities deserve explicit consideration	10
Verification and accountability become control capabilities	11
Values, relationships and intervention cannot be treated as afterthoughts	12
Humans can remain accountable after they lose the ability to judge	13
Human presence is not the same as human governance	14
AI can reduce the need for expertise while increasing the value of expertise	15
The objective should not be maximum automation	16
The task becomes more useful than the job	17
Two decisions determine the operating model	18
Each quadrant creates a different L&D; response	19
Every “Human Governs” task should trigger a readiness assessment	20
Not every task needs the same degree of human participation	21
Assess consequence before assigning control	22
Develop only the capability required by the responsibility	23
Before automating, test what expertise could disappear	24

Financial analysis: automate production without automating accountability	25
People managers: let AI reduce administration, not ownership of the relationship	26
Apply the framework to L&D; itself	27
A one-page operating tool for consequential AI-enabled work	28
A practical decision tree	29
L&D; begins with the architecture of work, not with learning	30
L&D;'s value moves toward capability architecture	31
AI allows capability requirements to become more selective	32
Some capability becomes an organizational resilience asset	33
What if AI becomes better at judgment and oversight too?	34
What the research supports and what this report adds	35
Twelve actions for Chief Learning Officers	36
The future-ready employee is defined by responsibility, not competition with AI	37
Selected research and governance sources	38

EXECUTIVE SUMMARY

AI is forcing L&D; to answer a larger question

Most L&D; teams are approaching AI from the wrong end of the problem.

They are asking how AI can help create courses faster, personalize learning, generate assessments, produce simulations or improve instructional design. Those questions are useful. They are also too small.

AI is beginning to change the work that L&D; exists to support. It can retrieve knowledge, produce professional work products, analyze information, recommend actions and increasingly execute multi-step work. The practical implication is not that all jobs disappear. It is that the boundary between human work and machine work becomes movable.

Once that boundary moves, a new strategic distinction becomes necessary:

THE DISTINCTION

Capability asks: What can AI perform?

Responsibility asks: What should humans continue to understand, decide, verify, govern and remain accountable for?

The central recommendation is that L&D; move upstream. The learning strategy should follow a deliberate decision about the human-AI division of responsibility, not precede it.

EXECUTIVE SUMMARY

Three findings should change how L&D; approaches AI

Finding	What it means for L&D
1 The problem is responsibility, not simply skills	Start with the work. Decide what AI should perform and what humans must own before designing learning.
2 Expertise can become more valuable as AI performs more work	Humans increasingly need enough domain expertise to judge, verify and challenge machine-produced work.
3 Human oversight can fail even when a human is “in the loop”	Formal approval is meaningless if people lack competence, authority, time or practice to intervene.

THE OVERSIGHT PARADOX

Humans can remain formally accountable after becoming practically incapable of exercising the judgment that accountability requires.

This third finding is the most important risk in the report. It means L&D; must treat human oversight itself as a capability that needs to be designed, practised and assessed.

01 WHY THIS QUESTION IS URGENT

The frontier of economically useful AI capability is expanding

The argument does not require predicting a technological singularity. It requires recognizing the direction of travel.

Stanford's AI Index documents continuing gains in reasoning, tool use and agentic performance while also emphasizing a jagged capability frontier: systems can be exceptional on sophisticated tasks and unexpectedly weak on others.

OpenAI's GDPval work evaluates AI on economically valuable, well-specified professional tasks drawn from dozens of occupations. Anthropic's Economic Index shows AI being used across a growing share of occupational tasks. The ILO takes a more cautious labour-market view and argues that task transformation is currently more plausible than wholesale job replacement for most exposed occupations.

IMPLICATION

The evidence does not justify a simple “AI replaces jobs” story. It supports a more useful conclusion: the boundary between human and machine work is becoming increasingly movable.

That makes work design, responsibility design and capability preservation strategic workforce issues, not merely technology-adoption issues.

02 THE WRONG QUESTION

Stop looking for skills AI will never possess

A common response to AI is to search for permanent human advantages: creativity, empathy, judgment, communication, critical thinking, leadership and ethics.

Those capabilities matter. But building strategy around the claim that machines will never become competent at them is fragile. AI capability is a moving target.

A STRONGER POSITION

Do not argue: “Humans must do this because AI cannot.”

Argue: “Humans should remain responsible for this because responsibility should not be surrendered merely because AI becomes capable.”

That distinction survives future capability improvements. It shifts the discussion away from human exceptionalism and toward organizational design, legitimacy, resilience and accountability.

03 CAPABILITY IS NOT RESPONSIBILITY

Technical capability does not determine authority

Suppose AI can identify which employees are most likely to leave. Three separate questions follow.

Question	Type
Can AI perform the analysis?	Capability
Should AI determine the action?	Authority / governance
Who owns the consequences?	Responsibility / accountability

The prediction might be technically excellent without determining whether the manager should promote the employee, redesign the role, change compensation, accept the departure or do nothing.

NIST, OECD and the EU AI Act all place increasing emphasis on explicitly defined human oversight, human agency and responsibility in consequential AI systems.

DECISION RULE

The organizational question is not merely “What activity can we automate?” It is “What authority are we prepared to delegate?”

04 HUMAN RESPONSIBILITY

Seven responsibilities deserve explicit consideration

1. Purpose

AI can optimize toward an objective. Humans still need to determine whether the objective is worth pursuing. A system may discover the fastest way to reduce call time, but it cannot legitimately decide whether speed should take priority over customer care, employee wellbeing, accessibility or long-term trust.

L&D; CAPABILITY

Problem framing, stakeholder understanding, systems thinking and questioning the objective before optimizing it.

2. Judgment

Prediction asks what is likely to happen. Judgment asks what should be done given competing goals, consequences and values. Better forecasting does not eliminate the need to choose among legitimate alternatives.

L&D; CAPABILITY

Interpretation, trade-off reasoning, scenario judgment and decisions under uncertainty.

04 HUMAN RESPONSIBILITY

Verification and accountability become control capabilities

3. Verification

When AI produces more first drafts, calculations, analyses and recommendations, humans increasingly become reviewers of machine-produced work. Review requires enough expertise to recognize when the output is wrong, incomplete, misleading or inappropriate.

A novice may be able to generate expert-looking work without being able to determine whether that work can be trusted.

4. Accountability

AI can recommend and increasingly act, but organizational accountability does not disappear. Employees need to know what decisions they own, what AI may do autonomously, when approval is required, what evidence should be retained and when escalation is mandatory.

KEY RISK

A human who technically appears in the workflow but lacks the capability or authority to disagree is not meaningful oversight.

04 HUMAN RESPONSIBILITY

Values, relationships and intervention cannot be treated as afterthoughts

5. Values and legitimacy

Many important decisions do not have objectively correct answers. Fairness, privacy, dignity, risk tolerance and competing stakeholder interests require normative choices that an organization must be prepared to defend.

6. Relationships

Work is also a social system. In coaching, leadership, conflict, caregiving, customer recovery and consequential conversations, the question is not only whether AI can communicate well. It is whether another human taking responsibility for showing up is itself part of the value.

7. Intervention

Someone must know when to stop the machine. As AI becomes more reliable, automation bias can increase and humans can become less prepared to recognize exceptions.

THE QUESTION

What capabilities must humans continue practising even if AI can normally perform them?

05 THE OVERSIGHT PARADOX

Humans can remain accountable after they lose the ability to judge

Imagine AI gradually takes over the routine reasoning in a financial-analysis workflow. Analysts still approve the result, but they perform less of the underlying modelling, assumption testing and scenario construction themselves.

The immediate effect looks positive: faster work, fewer manual steps, consistent output.

The long-term effect can be different. New analysts may never accumulate the experience through which previous generations developed pattern recognition. The organization can eventually reach a strange position:

- The AI performs the reasoning.
- The human approves the reasoning.
- The human remains accountable.
- The human may no longer possess enough independent expertise to challenge it.

THE OVERSIGHT PARADOX

Human responsibility can remain while human capability quietly disappears.

06 RUBBER-STAMP OVERSIGHT

Human presence is not the same as human governance

Meaningful oversight requires more than placing a person somewhere in the workflow.

MEANINGFUL HUMAN OVERSIGHT

Competence + AI literacy + Authority + Time + Practice + Accountability

A diagnostic model, not a mathematical equation.

Condition	Test
Competence	Does the person understand enough of the domain to evaluate the AI?
AI literacy	Do they understand relevant limitations, uncertainty and failure modes?
Authority	Can they actually reject, alter or stop the recommendation?
Time	Does the workflow permit genuine review rather than throughput pressure?
Practice	Will they retain enough exposure to maintain the required capability?
Accountability	Is responsibility for the outcome explicit and understood?

If a critical element is missing, the organization should question whether the process is genuinely human governed.

07 THE EXPERTISE PARADOX

AI can reduce the need for expertise while increasing the value of expertise

Automation can make sophisticated work easier to produce while making sophisticated errors harder for less experienced people to detect.

This creates a developmental problem. Traditional expertise often emerges through repeated exposure to routine work, mistakes, feedback, exceptions and progressively harder cases.

Traditional pathway	AI-enabled risk
Do basic work	AI removes much of the basic work
Make mistakes and receive feedback	Fewer opportunities to experience the reasoning process
Recognize patterns	Pattern recognition may develop more slowly
Handle harder cases	Experts are still expected to judge the hardest exceptions

APPRENTICESHIP QUESTION

If AI removes the developmental work, where will tomorrow's experts come from?

08 AUTOMATION VS. AUGMENTATION

The objective should not be maximum automation

Research on worker preferences and AI agents suggests that people do not uniformly want every technically automatable activity to be automated. Work provides more than output. It can also provide competence, autonomy, social connection, identity, mastery and meaning.

Erik Brynjolfsson and others have argued that replacement-focused AI strategies can miss the larger opportunity to augment human capability and redesign work.

STRATEGIC OBJECTIVE

Seek the strongest human-AI configuration for performance + safety + development + adaptability + human agency, not the highest possible percentage of automation.

This makes L&D; relevant to work redesign, because learning and development consequences are part of the operating-model decision, not something to address after the technology is deployed.

09 THE UNIT OF ANALYSIS

The task becomes more useful than the job

Jobs are bundles of tasks. AI rarely changes every task in a role at the same pace.

Task	AI capability	Human responsibility	Operating model
Information retrieval	High	Low	AI performs
First-draft analysis	High	Medium	AI leads, human reviews
Strategic recommendation	Growing	High	Human-AI collaboration
Ethical trade-off	AI can assist	Very high	Human decides
Compliance monitoring	High	High oversight	AI monitors, human accountable
Crisis response	Variable	Very high	Human command

L&D; QUESTION

What work actually happens inside this role, how is AI changing that work, and what should the human contribution become?

10 HUMAN RESPONSIBILITY FRAMEWORK

Two decisions determine the operating model

The Human Responsibility Framework starts with work, not learning. For every important task, ask two questions:

- How capable is AI of performing this task?
- How much human responsibility should remain?



The matrix deliberately separates capability from responsibility. High AI capability does not automatically imply low human responsibility.

11 FOUR OPERATING MODES

Each quadrant creates a different L&D; response

Mode	When it applies	Primary L&D response
HUMAN PERFORMS	Lower AI capability, lower human responsibility	Build sufficient task capability while monitoring automation trajectory.
HUMAN OWNS	Lower AI capability, higher human responsibility	Invest deeply in expertise, judgment, practice, coaching and experience.
AI PERFORMS	Higher AI capability, lower human responsibility	Teach system use, quality checks, exception handling and escalation.
HUMAN GOVERNS	Higher AI capability, higher human responsibility	Develop verification, judgment, authority, intervention and capability preservation.

MOST IMPORTANT QUADRANT

Human Governs is where organizations are most vulnerable to the Oversight Paradox because the machine may perform most of the cognitive work while the human retains responsibility.

12 MANDATORY OVERSIGHT CHECK

Every “Human Governs” task should trigger a readiness assessment

Do not assume oversight exists simply because a human is required to approve the decision.

Check	Required evidence
Competence	The person can independently assess the relevant reasoning and identify abnormal cases.
AI literacy	The person understands system limitations and likely failure modes.
Authority	The person can reject, alter, escalate or stop the output.
Time	The workflow provides enough time for meaningful evaluation.
Practice	The person has a mechanism to maintain underlying capability.
Accountability	The person knows exactly what they own.

DECISION RULE

If a critical requirement is missing, do not label the process human governed yet. Fix the operating model, work design or capability system first.

13 LEVELS OF HUMAN CONTROL

Not every task needs the same degree of human participation

Level	Human role	Operating logic
1 Human Performs	Performer	AI may assist, but the human does the work.
2 Human Decides	Decision-maker	AI analyzes and recommends. Human determines the action.
3 Human Approves	Approver	AI prepares the action. Human authorization is required.
4 Human Monitors	Supervisor	AI operates normally. Humans manage exceptions and intervene.
5 Human Audits	Governor	AI performs routine work. Humans periodically evaluate quality, controls and risk.

The correct level depends on consequence, reversibility, system reliability, legal requirements, values and the amount of human capability that must be retained.

14 SIX RESPONSIBILITY TESTS

Assess consequence before assigning control

Test	Question
Consequence	What happens if the AI is wrong?
Judgment	Does the task involve competing objectives rather than one correct answer?
Values	Does it affect fairness, dignity, rights, ethics or organizational values?
Relationship	Does human participation itself create important value?
Accountability	Would stakeholders reasonably expect a person to answer for the outcome?
Capability preservation	Would removing humans damage expertise the organization still needs?

RULE OF THUMB

The greater the combined consequence, judgment, values, relationship and capability-preservation requirements, the stronger the case for meaningful human control.

15 MINIMUM VIABLE HUMAN CAPABILITY

Develop only the capability required by the responsibility

Once responsibility has been assigned, L&D; should determine the minimum capability a person must retain to fulfil it safely and competently.

Requirement	Meaning
KNOW	Understand relevant facts, principles and boundaries.
INTERPRET	Understand what the AI output means in context.
VERIFY	Determine whether the output is credible and appropriate.
DECIDE	Exercise judgment and select the action.
PERFORM	Remain capable of performing the underlying activity independently.

IMPORTANT

Not every employee needs all five. The requirement is to retain enough human capability to discharge the assigned responsibility safely.

16 CAPABILITY PRESERVATION

Before automating, test what expertise could disappear

Automation plans should include a developmental question that is usually missing:

THE TEST

If humans stop performing this activity, how will future experts develop the capability required to supervise it?

Where the answer is unclear, the automation plan is incomplete.

Preservation mechanism	Use when
Simulation and failure scenarios	People need to practise rare but consequential exceptions.
Supervised practice / rotations	Expertise depends on experience across varied cases.
Case reviews / shadow decisions	Humans need to see expert reasoning even when AI executes.
AI-off practice	The underlying skill must remain available during system failure or crisis.
Apprenticeship pathways	The organization needs a reliable pipeline of future experts.

17 EXAMPLE

Financial analysis: automate production without automating accountability

AI increasingly performs	Human remains responsible for
Data gathering	Validating assumptions
Calculations	Understanding business context
Variance analysis	Recognizing unusual conditions
Scenario generation	Challenging implausible conclusions
Baseline forecasting	Interpreting implications
First-draft commentary	Recommending and owning action

The L&D; emphasis therefore shifts toward financial judgment, interpretation, verification, communication and decision-making rather than simply faster spreadsheet production.

OVERSIGHT CHECK

If analysts no longer build models or investigate assumptions themselves, what developmental mechanism ensures they can still challenge the AI five years from now?

18 EXAMPLE

People managers: let AI reduce administration, not ownership of the relationship

AI can assist with	Human should deliberately govern
Summarizing performance information	Understanding context
Drafting reviews	Evaluating mitigating circumstances
Pattern detection	Conducting consequential conversations
Meeting preparation	Making employment decisions
Option generation	Balancing competing considerations
Coaching prompts	Building trust and owning the relationship

L&D; should spend less time teaching managers how to compose polished review paragraphs if AI can do that reliably, and more time building judgment, coaching, difficult conversations, bias awareness and accountability.

19 EXAMPLE

Apply the framework to L&D; itself

The framework becomes uncomfortable when applied to the profession that created it.

AI can increasingly own	L&D should retain responsibility for
Research and first drafts	Problem diagnosis
Outlines, scripts and storyboards	Determining whether learning is needed
Assessments and scenarios	Understanding performance barriers
Visuals, translation and narration	Intervention selection
Routine course production	Evidence interpretation and quality standards
Administrative analysis	Capability strategy and human-AI work design

THE PROFESSION MOVES UPSTREAM

From “How do we create this course?” to “What capability does the organization require, where should it live, and what is the smallest intervention capable of producing reliable performance?”

20 THE HUMAN RESPONSIBILITY CANVAS

A one-page operating tool for consequential AI-enabled work

Canvas field	Decision
TASK	What outcome must be produced?
AI CAPABILITY	What can AI reliably perform today?
TRAJECTORY	Which parts are likely to become automatable?
CONSEQUENCE	What happens if this goes wrong?
HUMAN RESPONSIBILITY	What must remain human-owned?
CONTROL LEVEL	Perform / Decide / Approve / Monitor / Audit
MINIMUM HUMAN CAPABILITY	Know / Interpret / Verify / Decide / Perform
OVERSIGHT READINESS	Competence / AI literacy / Authority / Time / Practice / Accountability
CAPABILITY ATROPHY RISK	What expertise could disappear?
ESCALATION TRIGGER	When must a human intervene?
INTERVENTION	What is the smallest intervention required?
PERFORMANCE EVIDENCE	How will readiness be demonstrated?

21 DECISION PATH

A practical decision tree

Step	Question	Action
1	Can AI perform this reliably enough?	If no, human performs. If yes, continue.
2	Would failure create meaningful consequences?	If low, consider automation with proportionate controls.
3	Does the task involve consequential judgment, values, rights, relationships or accountability?	If yes, retain meaningful human authority.
4	What human role is required?	Perform / Decide / Approve / Monitor / Audit.
5	What human capability is required?	Know / Interpret / Verify / Decide / Perform.
6	Could automation erode that capability?	If yes, create a preservation mechanism.
7	Is oversight actually possible?	Run the mandatory readiness check.
8	What is the minimum effective intervention?	Training, simulation, workflow support, policy, coaching or redesign.
9	What performance evidence proves readiness?	Measure responsibility in practice, not course completion.

22 THE NEW L&D; OPERATING MODEL

L&D; begins with the architecture of work, not with learning

Traditional L&D	Human Responsibility Model
Training request	Business change
Needs analysis	Work decomposition
Course design	AI capability assessment
Development	Human responsibility decision
Delivery	Oversight-readiness assessment
Completion	Minimum human capability
Hope for transfer	Capability preservation
Post-program evaluation	Performance evidence and continuous reassessment

THE SHIFT

The learning strategy becomes a downstream response to work and responsibility design rather than the starting point.

23 FOUR STRATEGIC ROLES

L&D;'s value moves toward capability architecture

Role	Mandate
Capability Architect	Determine what capabilities the organization needs and where they should live: people, AI, workflows, teams or systems.
Human-AI Work Designer	Help define the division of labour and responsibility between people and machines.
Capability Risk Manager	Identify where automation could remove expertise, resilience or succession capacity the organization still needs.
Performance Engineer	Build the smallest intervention required to produce reliable performance and evidence.

Training remains one intervention. It is no longer the organizing principle of the function.

24 WHAT L&D; SHOULD STOP TEACHING

AI allows capability requirements to become more selective

If AI reliably performs a task and humans do not need the underlying skill for judgment, resilience or oversight, teaching everyone to perform that task manually may become wasteful.

FOR EVERY TRADITIONAL REQUIREMENT, ASK

Does the employee need to know this, understand it, interpret it, verify it, decide it, perform it, or simply know when to escalate?

These are different capability requirements. Treating them as one curriculum requirement is increasingly inefficient.

25 WHAT L&D; MUST PROTECT

Some capability becomes an organizational resilience asset

The opposite mistake is equally dangerous. Some capabilities must be deliberately maintained even if employees rarely use them.

- Domain expertise required to identify abnormal cases.
- Diagnostic reasoning and professional judgment.
- Ethical reasoning and high-consequence decision-making.
- Emergency and fallback capability when systems fail.
- Verification and independent decision-making.
- The ability to challenge, override and escalate AI.

CAPABILITY RISK

Critical capability can disappear accidentally. L&D; should identify that risk before automation removes the practice that sustains it.

26 THE STRONGEST OBJECTION

What if AI becomes better at judgment and oversight too?

The Human Responsibility Framework should not become a defence of human involvement for its own sake.

If AI eventually becomes materially better than humans at production, judgment, verification and monitoring, mandatory human intervention could itself reduce performance.

Human responsibility should therefore remain where there are compelling reasons such as legitimacy, accountability, legal requirements, human rights, human agency, relationships, resilience, social expectations or moral responsibility.

BOUNDARY CONDITION

The framework is not a defence of human work. It is a defence of deliberate responsibility.

27 EVIDENCE VS. SYNTHESIS

What the research supports and what this report adds

Research-supported observations	Synthesis developed here
AI capability is expanding into sophisticated professional and agentic tasks.	Human Responsibility Framework
Near-term labour evidence is more consistent with task transformation than universal job replacement.	Oversight Paradox
Workers do not uniformly prefer all technically automatable tasks to be automated.	Minimum Viable Human Capability
Augmentation is a meaningful alternative to replacement-focused design.	Human Responsibility Canvas
Governance frameworks increasingly emphasize human oversight, responsibility and agency.	Five levels of human control and the rule: Automate capability freely. Delegate responsibility deliberately.

Keeping this distinction explicit matters. The report uses evidence to support the direction of travel and then offers an L&D; operating model for decisions the research does not prescribe directly.

Twelve actions for Chief Learning Officers

#	Action
1	Break strategically important roles into meaningful tasks.
2	Assess where AI can already contribute.
3	Track where AI capability is advancing fastest.
4	Determine what responsibilities should remain human.
5	Assign the appropriate human control level.
6	Run the Oversight Readiness Check for every Human Governs task.
7	Define the minimum human capability required.
8	Identify capability atrophy and apprenticeship risk.
9	Protect expertise where necessary.
10	Build the minimum effective intervention.
11	Measure whether people can exercise responsibility, not merely whether they completed learning.
12	Reassess as technology, regulation and work change.

FIRST QUESTION

Do not start with “What AI training do employees need?” Start with “How is AI changing the division of responsibility between humans and machines in our organization?”

29 CONCLUSION

The future-ready employee is defined by responsibility, not competition with AI

AI may eventually perform a substantial portion of the cognitive work humans currently perform. That does not automatically make humans irrelevant. It makes the design of responsibility more important.

Organizations will increasingly need to decide what machines should perform, what humans should understand, what humans should verify, what humans should decide, what humans must remain capable of doing, what humans should be able to override and what humans remain accountable for.

THE CRITICAL QUESTION

If humans remain responsible, have we preserved enough human capability for that responsibility to mean anything?

Human presence is not enough. Human responsibility without competence creates risk. Human accountability without authority creates theatre. Human review without time creates rubber stamping. Human oversight without practice creates capability atrophy.

THE PRINCIPLE

Automate capability freely. Delegate responsibility deliberately.

L&D's corresponding mandate is to build and protect the human capability required to make that responsibility real.

SOURCE FOUNDATION

Selected research and governance sources

Stanford Institute for Human-Centered Artificial Intelligence. AI Index Report 2026. Evidence on capability progression, agentic systems and the jagged frontier of AI performance.

OpenAI. GDPval. Evaluations of frontier models on economically valuable professional work products across knowledge-work occupations.

Anthropic. Economic Index, March 2026. Evidence on how AI is being used across occupational tasks and patterns of augmentation and automation.

International Labour Organization. Generative AI and Jobs / Global Index research. Analysis of occupational exposure and task transformation.

Stanford Digital Economy Lab. Future of Work with AI Agents. Worker preferences and expert assessments of automation and augmentation potential across occupational tasks.

Erik Brynjolfsson and related Stanford research. Human-centred AI, augmentation, task redesign and expertise.

David Autor and related research. Automation, expertise and the changing economic value of human tasks.

NIST. AI Risk Management Framework and Playbook. Human-AI configurations, roles, responsibility and oversight.

OECD. OECD AI Principles. Human agency, oversight, accountability and responsible AI.

European Union. EU AI Act, including Article 14 on human oversight for high-risk AI systems.

UNESCO. Recommendation on the Ethics of Artificial Intelligence and related governance work on human accountability and AI literacy.

NOTE

The research base supports the report's evidence statements. The named frameworks and diagnostic concepts are the author's synthesis for L&D; and workforce capability decisions.