

Intelligence Amplified, Judgement at Risk.

Research Report on Artificial Intelligence, Human Behavior, Governance, and the Next Decade



About the Document

Intelligence Amplified, Judgment at Risk.

AI, Governance & the Next Decade • March 2026

SUSTAINABILITY

Sustainability Advisory

ABOUT THIS RESEARCH

This thought leadership paper was produced by Sajjeed at Sustainability, a sustainability advisory and technology firm. It synthesises peer-reviewed research, institutional reports, and real-time data from Epoch AI, Stanford HAI, McKinsey Global Institute, the World Economic Forum, Goldman Sachs, OECD, UNESCO, J.P. Morgan, Harvard Business School, MIT Media Lab, Nature Machine Intelligence, and the Council on Foreign Relations, among 30 other tier-1 sources. The analysis spans AI model scale and trajectory, human behavioral psychology, the employment impact of AI, training data provenance, five centuries of historical technology precedent, forward-looking scenario planning, and a stakeholder-specific governance action framework.

DEVELOPED IN PARTNERSHIP WITH

University of Lahore — Continuum Initiative

The University of Lahore (UOL) Continuum is an executive and professional development initiative designed to bridge academia and industry through accredited, applied learning programmes. The ESG Specialist Programme — delivered in partnership with Sustainability and Climena — represents Continuum's commitment to equipping Pakistan's professionals with globally recognised sustainability and AI governance competencies.

UOL is Pakistan's largest private-sector university, home to 35,000+ students from 20+ countries, with fully accredited programmes across Medicine, Engineering, Law, Business, and the Sciences.

#187

QS Asia 2026
Asia • #29 Southern Asia
#6 Pakistan • #2 Private

601–800

THE World 2026
#287 Research Quality globally
#2 Pakistan • #1 Private

#588

US News Global
#165 Asia • #4 Pakistan
#1 Private sector university

W4

HEC Category
Highest HEC category
#1 Private • Interdisciplinary
#114

QS Sustainability 2026 — Top 500 globally • #1 Private in Pakistan • #3 national

KEY CONTRIBUTORS

Lead Author & Research

Sajjeed Aslam
Partner, Sustainability LLC, USA

Academic Partner

University of Lahore
Continuum Initiative

ESG Programme

CLIMENA
Co-delivery Partner

AI COPILOT DISCLOSURE

This document was researched, designed, and developed under full human authorship and editorial governance. Claude (Anthropic) was used as an AI copilot to assist with data synthesis, structural drafting, and visualisation production. All arguments, conclusions, policy recommendations, and editorial judgements were made by the human author. Every AI-generated output was reviewed, verified against primary sources, and revised to reflect independent human reasoning before inclusion. No content was accepted without critical evaluation. This disclosure reflects the principle at the heart of this research: AI as an instrument of human capability — not a substitute for human judgment.

DISCLAIMER

Produced for informational, educational, and professional development purposes only. Views and recommendations represent the independent perspective of the authors and do not constitute financial, legal, regulatory, or investment advice. All data and rankings are cited from publicly available tier-1 sources; while every effort has been made to ensure accuracy at time of publication (March 2026), readers are encouraged to consult primary sources. University of Lahore rankings reflect QS, THE, and US News & World Report data as of 2025–2026 and are subject to change.



**Sajjeed Aslam, FCA
Partner**
Sustainadility LLC (USA)

MESSAGE FROM THE AUTHOR

Why This Research Was Undertaken And How It Is Intended to Be Used

Every era has a defining question. Ours is not whether artificial intelligence will change the world. **It already has.** The question is whether those changes will be equitable, governed, and directed by human values, or whether they will compound existing inequalities and erode the capacities we need to address them.

I undertook this research because the professional communities I serve, across sustainability, education, enterprise strategy, and public policy, are navigating AI at a moment when the volume of information is high and the quality of guidance is uneven. Too much of the current discourse on AI sits at one of two extremes: uncritical enthusiasm that treats every new model as a solution, or reflexive anxiety that treats capability as threat. Neither position serves the people who must actually make decisions.

This paper was designed to occupy the responsible middle: grounded in tier-1 data, anchored in historical pattern recognition, and explicit about both what AI makes possible and what it puts at risk. The five-century lens, from Gutenberg to the Industrial Revolution to the internet, was deliberate. Every major technology revolution has eventually become governable. The question AI poses is whether governance can arrive before the damage is structural, not after.

This report is intended for five audiences, each of whom faces a distinct version of the same underlying challenge. For government and policy professionals, it offers a structured framework for intervention before the default trajectory becomes irreversible. For investors and business leaders, it provides an honest assessment of where AI creates durable value and where it creates liability. For educators and institutions, it makes the case for building AI literacy as a public good rather than a competitive advantage for the few. For practitioners, including sustainability professionals, ESG specialists, and technology strategists, it offers the evidence base to make the argument internally. And for the intellectually curious reader, it offers the most important reassurance of all: this is a pattern we have seen before. We have governed it before. We can do so again.

This is not a document about fearing AI. It is a document about respecting it enough to govern it well, and about preserving the human judgment, curiosity, and ethical reasoning that no model can replicate and no algorithm can replace.

"The 5% of workers who are AI-fluent today earn 4.5x more than the 95% who are not. That single statistic is the organising imperative behind everything in this document."



Uzair Raoof
Deputy Chairman
University of Lahore

MESSAGE FROM THE UNIVERSITY OF LAHORE

Continuum: Beyond Degrees, Beyond Borders

A Lifelong Partnership for the People and Institutions We Serve

A university's obligation does not end at graduation. **That is the founding conviction of the University of Lahore Continuum Initiative.** We created Continuum because the world our graduates enter is not the world they were prepared for at enrollment. The pace of change, accelerated today by artificial intelligence at every level of professional life, demands that the relationship between an institution and its community be ongoing, adaptive, and genuinely useful.

Continuum is our answer to that demand. It is not a continuing education program in the traditional sense. It is a living ecosystem of knowledge, designed to serve four constituencies simultaneously: our students, who deserve preparation for the roles of 2030 and not merely those of today; our alumni, who carry the UOL name into organizations around the world and whose professional growth reflects directly on the institution; our faculty, who must engage with the frontier of practice to remain credible teachers of it; and our employer partners, who invest in our graduates and need to know that investment will compound over time.

The partnership with Sustainability on the ESG Specialist Program exemplifies what Continuum is designed to produce: rigorous, applied, globally relevant content delivered by practitioners who operate at the frontier of their field, credentialed through an institution that meets the standards of QS, Times Higher Education, and US News on the world stage. It brings Pakistan's largest private-sector university into direct partnership with the global sustainability and AI governance agenda.

This research represents the kind of thinking Continuum is committed to bringing into every engagement. The AI governance challenge is not an abstract academic question. It is the operational reality of every employer in our partner network, every professional in our alumni community, and every student deciding what skills to build before entering a labor market that is being structurally reorganized in real time.

Our guidance to all of them is the same: do not wait for certainty before acting. The organizations that will lead the next decade are those that invest now in human capability, including critical thinking, ethical judgment, and strategic adaptability, while deploying AI as a disciplined instrument rather than an unsupervised authority. That is the balance this research argues for, and it is the balance UOL Continuum is committed to making accessible: not only to Pakistan's top institutions, but to every organization that recognizes the urgency and chooses to act.

This is what beyond degrees means. And this is what beyond borders must come to mean for Pakistani higher education: not the export of talent, but the export of capability, judgment, and the kind of leadership the world actually needs.

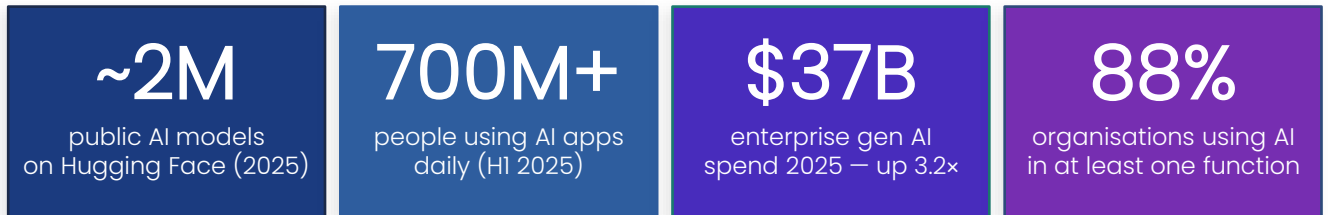
"The challenges facing contemporary industries demand more than qualified graduates, they demand thoughtful, ethically grounded professionals capable of leading through uncertainty. Continuum's academic model is designed precisely around that responsibility: to develop individuals who add genuine, sustained value to their disciplines."

At a Glance

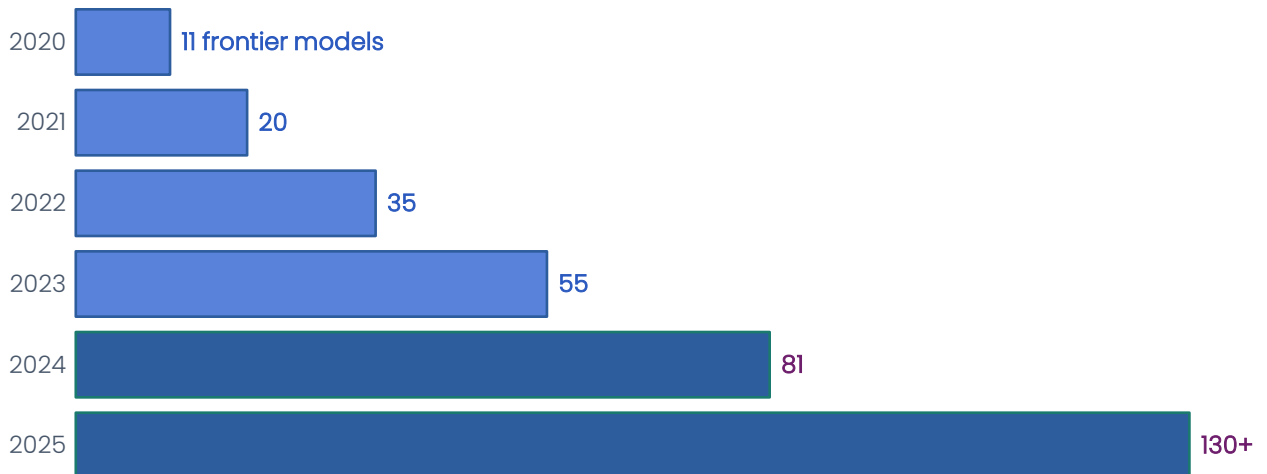


The AI Scale Explosion

From hundreds of models in 2020 to nearly 2 million today — AI is not a future trend. It is present infrastructure, reshaping how 700 million people work, decide, and understand the world.



Frontier model growth: Epoch AI tracks large-scale models (>10²³ FLOP training compute)



Seven AI model families — and their reach

Large Language Models GPT-5.4, Claude, Gemini, Llama	Agentic & Robotics Multi-agent, IEEE 2874-2025 Spatial Web
Image & Video Generation Midjourney, DALL-E 3, Sora, Veo 3	Embedding & Search Cohere Embed, text-embedding-3-large
Audio & Speech ElevenLabs, Whisper v3 — synthesis & cloning	Vertical Specialists Harvey AI (legal), FinGPT, AlphaCode
Scientific Domain Models AlphaFold 3, Med-Gemini, BloombergGPT	

Looking to 2036:

- AI agentic systems: \$8.6B (2025) → \$263B (2035)
- PwC: \$15.7T GDP contribution by 2030
- Goldman Sachs: 7% productivity gain after 10 years
- WEF: 170M new jobs created vs 92M displaced — net +78M
- Only 5% of workers AI-fluent today, earning 4.5x more

The Human: Behavior, Motivation & the Mirror

How humans really respond to powerful new tools

Four motivational drivers of technology adoption — proven across 500 years

Convenience & effort reduction

The most powerful driver. Humans consistently prefer tools that reduce cognitive load, even at cost of accuracy or autonomy. The printing press reduced copying effort. AI reduces reasoning effort. Gerlich (2025, n=666): frequent AI use correlates with critical thinking decline ($r=-0.68$, $p<0.001$).

Social proof & status signalling

AI fluency is now a professional status marker. The 5% of AI-fluent workers earn 4.5x more — creating powerful incentive gradients that accelerate adoption, including performative adoption for signalling rather than genuine capability.

Fear, anxiety & loss aversion

52% of US workers are worried about long-term AI impacts. Anticipatory anxiety mirrors Industrial Revolution mechanization fears (MDPI, 2025). Unaddressed anxiety drives two failure modes: over-adoption without judgment, or exclusionary resistance.

Curiosity & drive to explore

Intrinsic motivation to engage with novelty is both the greatest driver of AI's potential AND the capacity most at risk. MIT Media Lab EEG study: ChatGPT users exhibited least neural engagement of any tested group — far less than those working unaided.

The AI adoption spectrum — five distinct groups

5%	35%	30%	20%	10%
Power users	Pragmatic adopters	Anxious adopters	Resistors	Excluded
AI-fluent, earn 4.5x more. Treat AI as amplifier.	Productivity use without systematic verification. Highest hallucination risk.	Reluctant or performative use. Need institutional support.	Active rejection. Often in identity-anchored professions.	No meaningful access. Disproportionately Global South.

The cognitive offloading paradox

AI offloads reasoning, judgment, and creativity — the capacities that define human uniqueness and are hardest to recover once atrophied. Unlike calculators (arithmetic) or GPS (navigation), AI offloads the very capacities needed to evaluate, govern, and improve AI itself.

SBS Swiss Business School study (n=666): $r = -0.68$ between AI use and critical thinking. MIT Media Lab EEG: ChatGPT users showed least brain engagement of all groups tested.

"AI is a machine that can now understand us when we speak to it. What we need is ethics and anthropology and sociology and psychology and history and the humanities brought in."

— WEF Davos, Gutenberg Parenthesis Discussion, 2024

When Intelligent Machines Meet Human Psychology

Five collision dynamics with measurable consequences

01

PRECISION TARGETING OF COGNITIVE VULNERABILITY

100M+ views of AI-fabricated Iran-Israel war content (BBC Verify, 2026)

AI identifies personality type through self-reported online data, selects emotionally resonant messaging that bypasses rational scrutiny, and delivers it at machine speed before verification can occur. WEF 2026: disinformation is the top risk that 'catalyses or worsens all other risks.'

02

THE LABOR MARKET FRACTURE

55K jobs directly cut by AI in 2025 (Challenger, Gray & Christmas)

LLMs are disrupting high-wage, educated professions — the first time automation has reached the professional knowledge class at this scale. J.P. Morgan: cloud, web search, and computer systems design stopped growing at end of 2022, immediately after ChatGPT's release.

03

THE EPISTEMIC CRISIS

47% of enterprise AI users made major decisions on hallucinated content (2024)

700M+ people consult AI for information daily. When these systems fail — through hallucination, bias, or manipulation — they do so at population scale. The informational commons on which democratic deliberation depends is corrupted at the source.

04

CONCENTRATION OF POWER

5 companies control Western AI infrastructure — 3 from the same metro area

OpenAI, Google, Meta, Microsoft, and Anthropic control the foundation models on which virtually all Western AI applications are built. Their training data was extracted from the global commons. Their models reshape opinion, employment, and medical decisions — without democratic mandate.

05

EROSION OF HUMAN OVERSIGHT CAPACITY

HBR (Feb 2026): AI removing the 'messy work' that builds professional judgment

AI is simultaneously the system that most requires human oversight AND the technology most capable of eroding the capacities needed to provide it. As cognitive offloading becomes habitual, discernment, source verification, and ethical judgment — needed to govern AI — atrophy first.

500 Years of Technology Revolutions — One Pattern

Rapid capability → Concentration →
Disruption → Governance → Distribution



The critical difference: AI is the first cognitive revolution.

Every prior revolution automated physical Labor or bounded cognitive tasks. AI automates reasoning itself. The prior pattern of 'displaced workers retrain for new roles' depends on the availability of roles requiring cognitive skills automation cannot perform. When the automation performs those cognitive skills, the retraining pathway is structurally narrowed without historical precedent.

Three Futures for 2036

Where We Are Headed
— and Where We Choose to Go

Scenario A: Governed Acceleration

25% probability

Governance frameworks arrive before damage is structural. EU AI Act, OECD principles, UNESCO ethics, NIST AI RMF create binding standards. Open-weight models matched with open governance. AI fluency becomes a public education priority. Productivity gains redistributed through reskilling and progressive data taxation.

Projected outcomes:

- Net +78M jobs broadly distributed
- AI literacy as universal public good
- \$15.7T GDP — widely shared
- Democratic deliberation protected

Key indicator: G20 adopts binding AI transparency standards by 2027

Scenario B: Captured Acceleration

50% probability — current trajectory

Technology advances faster than governance. The AI-fluent 5% and their employers capture the majority of productivity gains. Displaced workers lack reskilling infrastructure. AI-generated disinformation degrades democratic processes. China and the US compete on capability rather than safety — a race to the bottom on governance.

Projected outcomes:

- GDP grows but inequality reaches historic highs
- 5% capture disproportionate gains
- Information commons corrupted
- Democratic erosion accelerates

Key indicator: No binding international AI data standards by 2027

Scenario C: Fractured Stagnation

25% probability

AI causes sufficient social disruption — mass unemployment, algorithmic election interference, AI-enabled security crisis — to trigger populist anti-technology backlash. Governance arrives as prohibition rather than regulation. The potential benefits of AI in healthcare, climate, and education are forfeited.

Projected outcomes:

- Prohibition rather than regulation
- Climate, health, education benefits lost
- Geopolitical AI competition hardens
- Democratic fragmentation on AI lines

Key indicator: 3+ major democracies pass AI prohibition legislation by 2028

The most important insight from scenario planning:

The 50% probability scenario (Captured Acceleration) is the default trajectory if governance does not intervene. It is not a failure of AI — it is the predictable outcome of applying AI in the same institutional environment that produced Gilded Age inequality, social media radicalization, and climate inaction.

The Action Agenda

Governance, Oversight & Stakeholder Responsibilities

7 interventions × 5 stakeholders × 3 non-negotiables

Three non-negotiables — the minimum viable governance architecture

1. DATA TRANSPARENCY BEFORE DEPLOYMENT	2. HUMAN OVERSIGHT BEFORE AUTONOMOUS ACTION	3. EQUITABLE ACCESS BEFORE CONCENTRATED POWER
No model deployed in high-stakes domains without full training data provenance, consent status, and demographic representation disclosed.	In healthcare, criminal justice, financial lending, weapons targeting, and elections — human review is not optional. Efficiency cannot justify opacity.	AI productivity gains must be distributed through AI literacy investment, reskilling infrastructure, and open-weight model ecosystems.

Five stakeholder responsibilities — why, what, and so what

Governments	NGOS & CIVIL SOCIETY	INVESTORS	BUSINESSES	ACADEMIA
WHY Set rules for whether AI serves many or few	WHY First to witness distributional impacts on marginalised communities	WHY Capital determines which apps get built and which governance gets funded	WHY Primary deployment layer — choices determine who benefits	WHY Generates the evidence base governance depends on
WHAT Bind training data transparency. Fund universal AI literacy. Mandate human-in-the-loop for high-stakes decisions.	WHAT Monitor impacts in real time. Advocate open multilingual training data. Build AI literacy in low-access communities.	WHAT Integrate AI governance into due diligence. Fund open-weight ecosystems. Require training data disclosure.	WHAT Human-in-the-loop for consequential decisions. AI literacy for all staff. Annual AI impact reports.	WHAT Build open demographically representative training datasets. Mandate AI literacy across all degrees.
SO WHAT Broadly distributed productivity gains; displaced workers with structured pathways	SO WHAT Communities as agents not subjects; evidence base policymakers cannot dismiss	SO WHAT Reduced regulatory risk; first-mover advantage in governance-compliant markets	SO WHAT Talent retention, liability reduction, customer trust in regulated markets	SO WHAT Evidence-informed governance ecosystem; professionals able to evaluate AI critically

THE VERDICT

AI deserves to be the first major technology revolution that is governed before the damage is done. The printing press took 80 years to develop governance. The Industrial Revolution the same. The internet took 20 years. AI is moving at 80 months, not 80 years.

The 5% of workers AI-fluent today earn 4.5× more than the 95% who are not. That single statistic determines whether the AI era is remembered as the decade humanity distributed the gains — or failed, again, to govern what it had built.

Primary sources: Stanford HAI AI Index 2025 • McKinsey Global Institute • WEF Future of Jobs 2025 • Goldman Sachs (Jan 2026) • OECD AI Principles & Governing with AI (2025) • UNESCO Ethics of AI • US Council of Economic Advisers (Jan 2026) • J.P. Morgan Global Research • Epoch AI Models Database • Nature Machine Intelligence Data Provenance Initiative • Gerlich (2025) Societies 15(1)8 • MIT MediaLab • CFR (Jan 2026) • RAND P-8014 • arxiv 2510.12859 • Menlo Ventures • AmplifiAI • RegenAIzer AI Job Market Dashboard • Challenger, Gray & Christmas • PwC • Gartner • EU AI Act (Aug 2025) • NIST AI RMF • IEEE 2874-2025

Intelligence Amplified, Judgement at Risk.

Report

EXECUTIVE SUMMARY

The question is no longer whether AI will transform civilization. It will. The question is whether that transformation serves humanity as a whole or concentrates its rewards among the few who control the models, the data, and the compute.

Artificial intelligence is not a future threat. It is the present infrastructure of **700 million daily users, \$37 billion in enterprise investment**, and nearly **2 million publicly available models** spanning language, vision, audio, scientific discovery, and autonomous action. The technology is no longer arriving. It has arrived, and it is compounding.

Yet the governance frameworks to contain it, the institutions to distribute its benefits equitably, and the human capacities to oversee it are all lagging behind the pace of deployment. That gap, between what AI can do and what our systems are designed to manage, is the subject of this research.

The central question of our era is not whether AI will transform civilization. It will. The question is whether that transformation serves humanity as a whole, or concentrates its rewards among the few who control the models, the data, and the compute. History suggests the answer depends almost entirely on whether governance arrives before or after the damage is structural.

The evidence is unambiguous on three fronts. **First**, cognitive science research now documents a measurable inverse relationship between AI reliance and critical thinking capacity. A 2025 study of 666 participants found a correlation of negative 0.68 between AI tool usage and independent reasoning ability. MIT Media Lab neuroimaging data corroborates this: ChatGPT users showed the lowest neural engagement of any group studied, including those who used no digital tools at all. The convenience of AI is real. So is its cognitive cost.

Second, geopolitical actors are already weaponizing generative AI at scale. The Iran-Israel conflict in early 2026 produced over 100 million views of AI-fabricated war content within days, documented by BBC Verify as the first AI-native military conflict in recorded history. Russia's cognitive warfare operations in Ukraine have synchronized AI-generated disinformation with physical missile strikes, using fabricated evacuation orders and synthetic video to paralyze civilian decision-making. Globally, AI-enabled disinformation campaigns have increased between 250 and 400 percent since 2023. The same technology that accelerates drug discovery and climate modeling is being deployed to destabilize democracies and manufacture consent at industrial scale.

Third, the labor market is bifurcating faster than policy has registered. Only 5 percent of the global workforce is AI-fluent today, and that 5 percent earns 4.5 times more than the 95 percent who are not. The World Economic Forum projects a net creation of 78 million jobs over the next five years, but that net figure conceals a gross displacement of 92 million roles and a gross creation of 170 million, concentrated almost entirely in high-skill, high-access economies. The workers most at risk are those least likely to receive retraining, and the institutions designed to serve them, including universities, regulators, and workforce development agencies, are operating on timelines that do not match the pace of disruption.

Every prior technology revolution, from **Gutenberg's printing press to the steam engine to the internet, eventually became governable**. The printing press took roughly 80 years to regulate. The Industrial Revolution took a similar period before labor protections, public education, and safety standards caught up. The internet's governance debate arrived two decades after deployment, by which point surveillance capitalism and algorithmic radicalization were already embedded in the architecture of public life. **AI's governance debate has begun before the most capable systems are fully deployed**. That is, by historical standards, the most favorable possible starting position.

But favorable is not sufficient. This report identifies **seven governance interventions**, ranging from mandatory training data transparency to international disinformation protocols, and maps a specific action agenda for five stakeholder groups: governments, investors, educators, practitioners, and individual professionals.

The scenarios section models three plausible futures for 2036, with current trajectory pointing toward Captured Acceleration, a world in which AI capability grows rapidly but its benefits accrue to an increasingly narrow group of actors.

The verdict is not pessimistic. It is urgent. The window to govern well is open. The question is whether those with the authority to act will use it before the window closes.

TABLE OF CONTENTS

01	The Technology: Scale, Scope & the Next Decade	15
02	The Human: Behavior, Motivation & the Mirror	17
03	The Collision: How AI Shapes Human Action	19
04	The Historical Lens: Patterns Across 500 Years	21
05	The Scenarios: Three Futures for 2036	23
06	The Governance Imperative: Policy & Oversight	24
07	The Action Agenda: Recommendations by Stakeholder	27
08	The Verdict: A Call to Govern	30



FROM HUNDREDS TO NEARLY TWO MILLION — AND ACCELERATING

The scale of AI model proliferation has no historical parallel in the history of software. In 2020, Epoch AI tracked just 11 frontier models exceeding the computational threshold that regulators use to monitor large-scale AI systems. By 2024, that number had risen to 81. Across all scales and types, approximately 1.98 million public models exist on Hugging Face alone — and the number grows by thousands every week.

This is not an exponential trend in a narrow technical domain. It is the emergence of **cognitive infrastructure**: a layer of machine intelligence now embedded in legal analysis, drug discovery, financial trading, education, military targeting, and the management of public opinion at scale.

~2M

Public AI models
(Hugging Face, 2025)

\$37B

Enterprise gen AI
spending in 2025 — up
3.2x year-on-year

700M+

People using AI
applications (H1 2025)

88%

Organizations using AI in
at least one business
function

The Seven Families of AI — and Their Reach

AI models are not a single technology. They span seven distinct capability domains, each with its own risk and opportunity profile:

- GPT-5.4, Claude Opus 4.6, Gemini 3.1, Llama 4, DeepSeek V3.2 — reasoning, writing, code, multimodal synthesis. Large Language Models (LLMs)
- Midjourney, DALL-E 3, Sora, Veo 3 — content creation, propaganda, deep-fake production. Image & Video Generation
- ElevenLabs, Whisper v3 — voice cloning, transcription, accessibility and manipulation. Audio & Speech
- AlphaFold 3 (proteins), Med-Gemini, BloombergGPT — accelerating discovery and concentrating scientific advantage. Scientific Domain Models
- Multi-agent orchestrators, IEEE 2874-2025 Spatial Web agents — autonomous action in the physical world. Agentic & Robotics Systems
- Cohere Embed 3, text-embedding-3-large — powering the information retrieval layer beneath every AI response. Embedding & Search
- Harvey AI (legal), AlphaCode (software), FinGPT (finance) — automating professional domains at scale. Specialized Vertical Models

The Data Provenance Crisis: Whose Knowledge Is This?

Every AI model is, at its core, a reflection of the data it consumed. Understanding provenance—where the data came from, who consented to its use, what biases it encodes—is central to understanding what the model knows, what it distorts, and whose interests it serves. The picture here is profoundly troubling: **transparency has decreased as capability has increased.**

A landmark audit published in Nature Machine Intelligence examined over 1,800 training datasets and found licence omission rates exceeding 70% and error rates above 50%. The Data Provenance Initiative documented that in under one year, ~5% of tokens in major corpora became restricted by robots.txt opt-outs — and multiple developers were accused of bypassing those restrictions to continue scraping.

Core training data sources (publicly documented across frontier models): Common Crawl (a petabyte-scale web scrape of the public internet, dominated by English-language Western content); Books and digitized libraries; Wikipedia and academic papers; GitHub code repositories; licensed publisher content (varying by developer); and increasingly, synthetic data generated by earlier model versions. Undisclosed third-party data slices can add 10–40% of training tokens (Nature Machine Intelligence / Data Provenance Initiative).

The bias implication is structural: **models trained predominantly on English-language, Western, high-income-population data will systematically underserve the 6.5 billion people who live outside that cultural and economic context.** This is not a technical glitch. It is an architectural choice with distributional consequences.

The Next Decade: What Is Coming by 2036

The trajectory of AI capability over the next ten years is uncertain in its precise timing but directional in its thrust. Based on the convergence of Goldman Sachs, Stanford HAI, McKinsey, and the US Council of Economic Advisers (January 2026):

2026–2028: Agentic Scale-Up

- 40% of enterprise applications will incorporate task-specific AI agents by end of 2026 (Gartner)
- AI agents completing software tasks double in length every 7 months (METR research, 2025) — by 2027 they may handle full developer work-weeks autonomously
- Personal AI agents arrive: systems that book flights, reschedule meetings, manage finances without explicit instruction (Goldman Sachs, January 2026)
- AI economic dashboards track task-level displacement in real time (Stanford HAI, Erik Brynjolfsson)

2029–2036: Structural Transformation

- AI agentic systems market: \$8.6B (2025) → \$263B (2035)
- PwC projects AI contributes \$15.7 trillion to global GDP by 2030 — 14% of global GDP
- Goldman Sachs: 7% productivity increase after 10 years; BIS Academic Working Paper scenarios range 20–45%
- 170M new jobs created vs 92M displaced — net +78M, but skills gap creates structural inequality (WEF)
- AGI consensus timeline: 2030s (Stanford), with 50% probability of key milestones by 2028

“In 2026, arguments about AI’s economic impact will finally give way to careful measurement. Early-career workers in AI-exposed occupations are already experiencing weaker employment and earnings outcomes.” — Erik Brynjolfsson, Stanford HAI, January 2026

WHAT WE KNOW ABOUT HOW HUMANS BEHAVE WITH POWERFUL NEW TOOLS

To understand AI's impact, we must first understand the human. Not the abstract rational agent of economics textbooks, but the actual human: cognitively limited, emotionally driven, socially anchored, and remarkably consistent in the patterns of behavior exhibited across technological transitions.

The Four Motivational Drivers of Technology Adoption

Human engagement with technology — across centuries and cultures — is governed by four primary motivational drivers, each of which AI now exploits at unprecedented scale:

Convenience and Effort Reduction	Social Proof and Status Signalling
The most powerful driver of technology adoption across history is the elimination of friction. Humans consistently prefer tools that reduce cognitive load, even at the cost of accuracy, autonomy, or long-term capability. The printing press reduced the Labor of copying manuscripts. The calculator reduced arithmetic effort. AI reduces the effort of writing, coding, reasoning, and deciding. A 2025 study (Gerlich, SBS Swiss Business School, n=666) found a significant negative correlation between AI tool usage and critical thinking scores ($r=-0.68$, $p<0.001$). The convenience trap is not a bug in AI adoption—it is a feature of human psychology.	Humans adopt technologies that signal competence and belonging to valued groups. In professional contexts, AI fluency has become a status marker: the 5% of workers who are AI-fluent earn 4.5x more, creating powerful incentive gradients. This accelerates adoption—but also performative adoption, where tools are used for signalling rather than genuine capability augmentation.
Fear, Anxiety and Loss Aversion	Curiosity and the Drive to Explore
An MDPI 2025 meta-study found that “anticipatory anxiety” significantly influences technology adoption, mirroring historical patterns: Luddite fears in the Industrial Revolution, 1990s internet skepticism, and now AI displacement anxiety (52% of US workers concerned per survey data). Anxiety, when not addressed by institutions, drives two distinct failure modes: over-adoption without judgment, or resistance that leaves communities behind.	Human curiosity—the intrinsic motivation to understand and engage with novelty—is both the greatest driver of AI's beneficial potential and the capacity most at risk from its convenience architecture. MIT Media Lab EEG research found that participants using ChatGPT exhibited the least neural engagement of any group tested—less than those using search engines, far less than those working unaided. Convenience, when systematically provided, suppresses curiosity.

The Cognitive Offloading Paradox

Cognitive offloading—delegating mental tasks to an external tool—is not new. We offloaded arithmetic to calculators, navigation to GPS, memory to smartphones. Each delegation was initially contested as a threat to human capacity, then normalized. AI represents the same pattern, but with a critical difference: prior tools offloaded specific, bounded tasks. AI offloads **reasoning, judgment, and creativity**—the capacities that have defined human uniqueness and that are hardest to recover once atrophied.

The research is unequivocal: SBS Swiss Business School (2025, n=666): significant negative correlation between AI tool usage and critical thinking ($r=-0.68$). Microsoft CHI 2025 survey of 319 knowledge workers: AI enhanced productivity but impaired critical thinking and created long-term technological dependency. MIT Media Lab EEG study: ChatGPT users showed least brain engagement of all tested groups. The long-term implication is not that AI makes individuals dumber — it is that systematic offloading erodes the cognitive infrastructure that future generations will need to evaluate, govern, and improve AI itself.

The Behavioral Spectrum: How Different Groups Engage with AI

Human adoption of AI is not uniform. Motivational research and demographic data reveal a consistent spectrum:

- Treat AI as amplifier. Earn 56% premium. Already operating as hybrid human-AI workers. Power users (AI-fluent, **5% of workforce**):
- Use AI for specific productivity tasks. Accept outputs without systematic verification. Most vulnerable to hallucination-driven decisions. Pragmatic adopters (**~35%**):
- Use AI reluctantly or performatively. High displacement anxiety. Require institutional support to convert anxiety into productive engagement. Anxious adaptors (**~30%**):
- Active or passive rejection. Often in occupations with strong professional identity (medicine, law, teaching) where AI threatens expertise-based status. Resisters (**~20%**):
- No meaningful access to AI tools due to cost, connectivity, or language barriers. Disproportionately in the Global South and rural communities. Excluded (**~10%**):

"AI is a machine that can now finally understand us when we speak to it. What we need is ethics and anthropology and sociology and psychology and history and the humanities brought into the conversation." — WEF Davos, Gutenberg Parenthesis Discussion, 2024

WHEN INTELLIGENT MACHINES MEET HUMAN PSYCHOLOGY

The preceding two chapters establish the machine and the human separately. Chapter 3 addresses what happens when they interact at scale. The evidence documents five critical collision dynamics, each with measurable consequences.

Collision 1: Precision Targeting of Cognitive Vulnerability

AI's most potent behavioral impact is not in what it produces, but in what it targets. Modern AI systems can identify personality type through self-reported online data, select emotionally resonant messaging that bypasses rational scrutiny, and deliver it at machine speed before verification can occur. The WEF Global Risks Report 2026 placed mis- and disinformation among the top short-term global risks, noting it "catalyzes or worsens all other risks on the list."

The mechanism is well-documented in the Russia-Ukraine cognitive warfare campaign (Frontiers in Artificial Intelligence, 2025, peer-reviewed): psychological manipulation in cognitive warfare operates at a subconscious level, exploiting emotional salience to bypass rational scrutiny. Russia's disinformation operations in 2025 synchronized with missile strikes, using AI to generate narrative surges within 3–6 hours of every military event — each designed to reframe civilian casualties as precision engagement before independent verification could arrive. In the Iran-Israel conflict (February 2026), BBC Verify found that AI-fabricated war content accumulated over 100 million views before fact-checks could catch up.

Collision 2: The Labor Market Fracture

AI is now actively reshaping who can earn a living and on what terms. The pattern differs critically from prior automation waves, which primarily targeted blue-collar and routine manual work. Large language models are disrupting high-wage, highly educated professions — the first time automation has reached the professional knowledge class at this scale and speed.

J.P. Morgan Global Research documents a mildly negative correlation between employment trends and AI usage in exposed sectors. **Cloud computing, web search, and computer systems design employment — three of the fastest-growing sectors of the 2010s — stopped growing at the end of 2022, immediately after ChatGPT's release.** The US Council of Economic Advisers' January 2026 report projects GDP gains of 2.4–4.1% (McKinsey baseline) to 7% (Goldman Sachs) over 10 years—but these aggregate gains mask catastrophic displacement for specific cohorts.

55K

Jobs directly attributed to AI cuts in 2025 alone

80%

US workforce with 10%+ of tasks AI-affected (OpenAI/MIT research)

56%

Salary premium for AI-fluent workers (WEF/Ragenaizer 2026)

\$15.7T

Projected AI contribution to global GDP by 2030 (PwC)

The inequality arithmetic: WEF projects 170M new jobs vs 92M displaced (net +78M) by 2030. But only 5% of current workers are AI-fluent, and they earn 4.5x more. The new jobs being created require technical expertise that mid-career displaced workers in legal support, data analysis, and content creation cannot readily acquire. This is not a short-term adjustment. It is a structural bifurcation of the Labor market.

Collision 3: The Epistemic Crisis — When AI Becomes the First Filter

Over 700 million people now routinely consult AI models for information, news synthesis, and decision support. When these systems fail—through hallucination, bias, or deliberate manipulation—they do so at population scale. 47% of enterprise AI users made at least one major decision based on hallucinated content in 2024 (AmplifAI data). During the Iran-Israel conflict, Google's AI-powered search summaries actively repeated false claims when users performed reverse image searches on fabricated war footage.

The deeper risk is epistemic: when populations routinely accept AI outputs without scrutiny, and when those outputs are themselves shaped by training data that lacks provenance transparency, the informational commons on which democratic deliberation depends is corrupted at the source.

Collision 4: The Concentration of Power

AI is not distributing power. It is concentrating it. Five companies — OpenAI, Google, Meta, Microsoft, and Anthropic — control the foundation models on which virtually all Western AI applications are built. Three of these are headquartered in the same metropolitan area. Their training data was extracted from the global commons. Their models are being used to reshape public opinion, determine employment, allocate credit, and guide medical decisions — without democratic mandate and with minimal accountability.

The geopolitical counterweight — China's DeepSeek, Alibaba's Qwen, state-affiliated models — introduces a parallel concentration with its own sovereignty concerns. Neither axis of this power structure was chosen by the populations it governs. As the Council on Foreign Relations documented in January 2026: "Who bears responsibility for an AI system's actions? Which governance models will fill the vacuum while democracies deliberate?"

Collision 5: The Erosion of Human Oversight Capacity

The most dangerous long-term collision dynamic is this: AI is simultaneously the system that most requires human oversight and the technology most capable of eroding the human capacities needed to provide that oversight. As cognitive offloading becomes habitual, the skills of discernment, source verification, probabilistic reasoning, and ethical judgment—the very skills needed to govern AI—are the first to atrophy. HBR's February 2026 research documented that AI is removing the 'messy work' that traditionally built judgment in junior professionals: the manual document review, the probabilistic estimation, the uncertain client conversation. These were not inefficiencies. They were the curriculum of professional development.

FIVE PRECEDENTS — AND WHAT HISTORY ACTUALLY TEACHES

The most common error in the AI debate is treating it as unprecedented. The specific technology is new. The patterns are ancient. Five historical moments offer both analogy and warning.

THE PRINTING PRESS (1440s)	Gutenberg's press democratized information and spurred the Renaissance, the Scientific Revolution, and the Reformation. It also enabled the industrialization of propaganda (Luther's 95 Theses spread across Europe in weeks), the concentration of publishing power in a small number of cities, and 150 years of religious warfare partly fuelled by competing information ecosystems. RAND's 1998 analysis documents that <i>"changes in the information age will be as dramatic as those in the Middle Ages, and the future will be dominated by unintended consequences"</i> . AI is Gutenberg at machine speed.
THE INDUSTRIAL REVOLUTION (1760–1840)	Mechanization automated physical labor, created unprecedented wealth, and triggered mass displacement of skilled artisans. The Luddite movement was not irrational: handloom weavers who had spent careers building expertise correctly identified that the machine would destroy their livelihoods. What they could not know was that governance—child labor laws, the Factory Acts, trade unions, universal education—would eventually distribute the gains. That governance took 80 years to arrive. AI is moving at 80 months, not 80 years.
THE TELEGRAPH AND MASS MEDIA (1840–1900)	The telegraph compressed information transmission from weeks to seconds, enabling coordinated financial markets, coordinated military campaigns, and coordinated disinformation. Yellow journalism — the deliberate fabrication of emotionally charged stories to drive newspaper sales and public opinion — contributed to the Spanish–American War of 1898. The mechanism was identical to AI-generated conflict propaganda in 2026: speed of emotional content outrunning verification. The difference is cost. A 19th-century yellow journalism operation required presses, journalists, and distribution networks. An AI operation in 2026 requires a laptop and a subscription.
THE INTERNET AND SOCIAL MEDIA (1990s–2010s)	The internet's promise was universal access to knowledge and connection. Its delivery was surveillance capitalism, algorithmic radicalization, and the weaponization of attention. Facebook's internal research (2021, the Wall Street Journal) documented that the platform's algorithms <i>"amplify hate, misinformation and political unrest"</i> —and that executives chose engagement over safety. Social media's governance framework arrived 20 years after the technology.
THE GENERATIVE AI REVOLUTION (2022 ONWARDS)	ChatGPT reached 100 million users within two months of launch in November 2022, faster than any technology in recorded history. Unlike every prior revolution, AI does not automate a bounded category of labor. It automates reasoning, synthesis, and judgment across virtually every professional domain at once. The promise is real: accelerated discovery, democratized expertise, and historic productivity gains. The risk is equally real: a bifurcating labor market, measurable cognitive atrophy under AI dependence, and industrial-scale disinformation at near-zero cost. For the first time, the governance debate has begun before the most capable systems are deployed. Whether that window is used is the defining choice of this generation.

The Critical Difference: AI Is the First Cognitive Revolution

Every prior technological revolution automated physical labor or specific bounded cognitive tasks. AI automates reasoning itself. This distinction matters for governance because the prior pattern of ‘displaced workers retrain for new roles’ depends on the availability of roles that require cognitive skills the automation cannot perform. When the automation performs the cognitive skills, the retraining pathway is structurally narrowed in ways without historical precedent.

As the diginomica analysis (April 2025) documents: “The textile worker could become a factory worker with transferable mechanical skills. The factory worker could transition to service work with transferable interpersonal skills. But what transferable skills will help the accountant, paralegal, or content creator displaced by AI? The new jobs require specialized technical expertise or advanced degrees not readily accessible to most displaced workers.” **This is not pessimism. It is a design problem.** And design problems have solutions—if they are identified before the damage is structural.

The Consistent Historical Pattern (RAND / arxiv.org Three Lenses, 2025)

Every major technological revolution follows a five-stage cycle:

- (1) Rapid capability growth outpacing governance;
- (2) Early concentration of power and profit among innovators;
- (3) Broad social disruption generating resistance and political pressure;
- (4) Institutional response creating governance frameworks;
- (5) Eventual distribution of benefits

At the cost of decades of inequity and, often, conflict. The question AI poses is whether we can compress Stage 4, because Stage 3 is already underway.

FIVE PRECEDENTS — AND WHAT HISTORY ACTUALLY TEACHES

Scenario planning is not prediction. It is structured preparation. The following three scenarios are grounded in documented trends and represent genuinely plausible alternative futures, not optimist/pessimist bookends. Each has early indicators that decision-makers can monitor today.

✓ Governed Acceleration	! Captured Acceleration	X Fractured Stagnation
Probability: 25%	Probability: 50%	Probability: 25%
- Governance frameworks arrive before damage is structural. - The EU AI Act, OECD principles, UNESCO ethics framework, and NIST AI RMF create binding interoperability standards. - Open-weight models are matched with open governance. AI fluency becomes a public education priority. - The productivity gains of AI are redistributed through reskilling investment and progressive data taxation. Net benefit: +78M jobs, \$15.7T GDP, distributed. Who benefits: broadly. Indicator: G20 adopts binding AI transparency standards by 2027.	- Technology advances faster than governance. - The 5% of AI-fluent workers and the companies they serve capture the majority of productivity gains. - Displaced workers lack reskilling infrastructure. AI-generated disinformation degrades democratic deliberation. - China and the US compete on capability rather than safety, creating a race to the bottom on governance. Net impact: GDP grows but inequality reaches historic highs. it is the present trajectory. Indicator: No binding international AI data standards by 2027.	- AI causes sufficient social disruption, mass unemployment, algorithmic disinformation in elections, an AI-enabled security crisis to trigger populist anti-technology backlash. - Governance arrives as prohibition rather than regulation. - Democratic societies fragment on AI policy lines. - Geopolitical AI competition hardens. - The potential benefits of AI, in healthcare, climate, education are forfeited. Indicator: 3+ major democracies pass AI prohibition legislation by 2028.

The most important insight from scenario planning: **the 50% probability scenario is the default trajectory if governance does not intervene.** Captured Acceleration is not a failure of AI. It is the predictable outcome of applying AI in the same institutional environment that produced Gilded Age inequality, social media radicalization, and climate inaction. The technology reflects the governance context it operates in.

SEVEN STRATEGIC POLICY INTERVENTIONS

The question before every government, institution, and organization is not whether to engage with AI governance. The technology is already deployed. The question is whether governance arrives before the damage is structural — or after, as an apology. The following seven interventions represent a minimum viable governance architecture, grounded in OECD, UNESCO, NIST, IEEE, and EU AI Act frameworks.

Intervention 1 — Mandatory Training Data Transparency

What: Require all general-purpose AI model developers to publish, in machine-readable format, a full audit trail of training data sources, licensing status, consent mechanisms, and demographic representation.

Why: Nature Machine Intelligence's audit of 1,800 datasets found license omission rates >70%. California's AB 2013 (2024) and the EU AI Act (August 2025) represent early mandates — but both allow 'sufficiently detailed summaries' that remain vague. Binding specificity is required.

How: Adopt the 12-point disclosure framework of California AB 2013 as a global minimum standard via OECD/G20 mechanisms. Require independent third-party audits before model deployment at scale. Create accredited knowledge bases — curated, consent-based, demographically representative corpora — for training models used in high-stakes domains.

Intervention 2 — Human-in-the-Loop Mandates for High-Stakes Decisions

What: Establish legally enforceable requirements for human review of AI-generated outputs in: healthcare diagnosis and treatment recommendation, criminal justice (bail, sentencing, parole), financial lending and credit scoring, electoral systems and voter communications, weapons targeting and military engagement rules.

Why: The EU AI Act classifies many public-sector AI use cases as 'high-risk'; South Korea's AI Basic Act (2024) designates 'high-impact' use cases requiring enhanced accountability. The principle is correct. The implementation is incomplete. IDC projects that by 2027, half of all AI applications will require dedicated oversight positions—but these roles do not yet exist at scale.

How: Establish sector-specific oversight standards through existing regulators: health authorities for medical AI, judicial commissions for criminal justice, and financial regulators for lending and credit. Require annual audits of human review procedures for all high-stakes AI deployments. Create a certified AI oversight practitioner designation, accredited through national professional bodies and recognized across OECD member states.

Intervention 3 — Universal AI Literacy as a Public Good

What: Fund universal AI literacy programs equivalent in ambition to post-WWII universal literacy campaigns—targeted at the 95% of workers who are not AI-fluent and the communities most at risk of exclusion.

Why: The 56% salary premium for AI fluency is documented. The gap between 5% AI-fluent and 95% not is the defining inequality metric of the next decade. The WEF explicitly frames AI literacy as a prerequisite for inclusive growth in its 2025 Workforce Blueprint.

How: Integrate AI literacy (understanding how models work, how to evaluate outputs, how to identify bias and hallucination) into national curricula at primary and secondary levels. Fund adult reskilling at scale: the OECD recommends 'training programmed along the working life' with social protection for those displaced. Establish micro credential frameworks validated by industry and accredited by national higher education bodies.

SEVEN STRATEGIC POLICY INTERVENTIONS

Intervention 4 — Open, Accredited, Multilingual Knowledge Bases

What: Create publicly governed, demographically representative, multilingual training corpora—the equivalent of public libraries for AI training data—accessible to all developers and mandated for use in AI systems deployed in public services.

Why: Current frontier models are trained predominantly on English-language, Western, high-income-population data. Multilingual models like Alibaba's Qwen (119 languages) and Cohere's Aya (101 languages) demonstrate the feasibility of broader representation—but commercial models will not build this infrastructure without mandate or incentive. The 6.5 billion people outside English-language dominance deserve AI systems trained on their knowledge, their languages, and their contexts.

How: Establish an UN-sponsored International AI Data Commons governed by a multi-stakeholder board representing governments, civil society, and academia. Fund initial development through G20 contributions as public infrastructure investment. Mandate commons-sourced data for AI systems deployed in public services, education, and healthcare. Adopt an extended Creative Commons licensing framework as the default consent mechanism.

Intervention 5 — Data Sovereignty and Portable Rights

What: Establish the legal principle that individuals have portable data rights: the right to know when their data has been used in AI training, the right to opt out prospectively, and the right to data deletion from training corpora.

Why: Meta has confirmed using public Facebook and Instagram posts to train Llama 3 and 4. Google's Gemini draws on Gmail patterns and YouTube data. None of these users provided informed consent for their data to become the training corpus of commercial AI systems generating \$37 billion in enterprise revenue. The EU's GDPR provides a partial framework—but covers personal data, not the collective data commons.

How: Extend GDPR-equivalent frameworks to cover collective data use in AI training. Require developers to maintain auditable registries of user-generated content used in each model version. Establish a right of deletion at training corpus level as a condition of operating license. Create international data sovereignty agreements, modeled on existing intellectual property treaties, that apply these rights across jurisdictions.

Intervention 6 — International AI Disinformation Protocols

What: Establish binding international protocols, negotiated through the UN Global Dialogue on AI Governance (launched 2026), requiring AI-generated content to carry machine-readable provenance markers, and requiring platforms to apply independent fact-checking at military-conflict disinformation speed.

Why: BBC Verify documented 100M+ views of AI-fabricated Iran-Israel war content before verification caught up. The Atlantic Council's DFRLab found Grok gave contradictory authenticity verdicts on the same fake video 300+ times. Meta's Oversight Board ruled the platform failed to flag AI-generated war content despite six user reports. Platform self-regulation has demonstrably failed at the required speed and scale.

How: Adopt the C2PA provenance standard as a binding international requirement for all AI-generated media, enforced through platform operating licenses. Require real-time provenance verification at recommendation algorithm speed. Fund an independent UNESCO/OECD disinformation monitoring body with authority to issue public compliance reports. Mandate 24-hour takedown timelines for verified AI-fabricated conflict content during active military engagements.

SEVEN STRATEGIC POLICY INTERVENTIONS

Intervention 7 — AI Governance as a Competitive Advantage, Not a Compliance Cost

What: Reframe AI governance—for policymakers, investors, and businesses—as a source of sustainable competitive advantage and stakeholder trust, not as regulatory friction.

Why: Organizations that implement governance now—before compulsion—build institutional trust, reduce liability exposure, attract AI-fluent talent, and develop the internal capabilities needed to deploy AI effectively. The WEF's Workforce Blueprint 2026 is explicit: 'The most successful organizations will invest in building human capabilities—such as critical thinking, creativity, and discernment—alongside AI fluency.' These are not competing priorities. They are the same priority.

How: Develop an independently verified AI Governance Certification standard administered through ISO or an equivalent body. Create tax incentives for early adoption: AI literacy programs, human oversight infrastructure, and open training data contributions should qualify as R&D expenditure. Publish an annual World AI Governance Index jointly through WEF and OECD. Make governance certification a prerequisite for public sector AI procurement across G20 member states.

WHY, WHAT, HOW, AND SO WHAT FOR EVERY RESPONSIBLE ACTOR

Governance is not a government function alone. It is a distributed responsibility across five stakeholder groups, each with distinct leverage, accountability, and urgency. The following framework maps specific interventions to each group

	Why It Matters	What Must Change	How To Act	So What? (Outcome)
 GOVERNMENTS	Governments set the rules that determine whether AI serves the many or the few. Current regulatory frameworks are fragmented, slow, and outpaced by capability. The cost of inaction is structural inequality and democratic erosion.	Enact binding AI transparency mandates. Fund universal AI literacy. Create human-in-the-loop requirements for high-stakes decisions. Establish international disinformation protocols. Build accredited public training data commons.	Adopt California AB 2013's 12-point framework as national minimum. Commission annual AI inequality impact assessments. Fund adult reskilling at 0.5% of GDP (comparable to post-war Marshall Plan as % of economy). Ratify UN Global AI Governance framework when available.	A governed AI ecosystem where productivity gains are broadly distributed, where democracies can maintain information integrity, and where displaced workers have structured pathways to new roles.
 NGOs & CIVIL SOCIETY	Civil society organizations are the first to witness AI's distributional impacts on marginalized communities — and the institutions with the trust infrastructure to bridge communities and governments.	Monitor and document AI's distributional impacts in real time. Advocate for open, multilingual, consent-based training data. Build AI literacy in communities with lowest access. Hold governments and companies accountable to stated commitments.	Establish AI impact observatories tracking employment, misinformation, and health outcomes in communities. Partner with Cohere (Aya), Mistral, and open-model communities to build multilingual, culturally representative training data for Global South contexts. Train community leaders as AI literacy educators.	Communities equipped to engage with AI as agents rather than subjects. Evidence base for policy advocacy that policymakers cannot dismiss.

07 THE ACTION AGENDA | RECOMMENDATIONS BY STAKEHOLDER

	Why It Matters	What Must Change	How To Act	So What? (Outcome)
INVESTORS	Capital allocators determine which AI applications get built and which governance frameworks get funded. ESG-aligned investment in AI governance is now a fiduciary issue, not a values question: regulatory risk, litigation exposure, and reputational liability are material.	Integrate AI governance into investment due diligence. Fund AI safety research and governance infrastructure. Divest from AI applications that systematically violate consent or exacerbate inequality. Support open-weight model ecosystems.	Develop AI Governance Scoring frameworks for portfolio companies (analogous to ESG scoring). Require disclosure of training data provenance in investment mandates. Allocate 2-5% of AI portfolio to governance infrastructure. Back open-weight model initiatives (Mistral, Llama) that distribute capability without concentrating power.	Reduced regulatory and reputational risk. Portfolio alignment with sustainable AI regulation trajectory. First-mover advantage in governance-compliant AI markets, particularly in the EU and Global South.
BUSINESSES	Enterprises are the primary deployment layer for AI. Their choices about which models to use, how to apply them, and what oversight to maintain determine whether AI's productivity gains benefit employees and customers or merely shareholders.	Implement human-in-the-loop requirements for consequential decisions. Invest in AI literacy for all staff — not just technical teams. Publish annual AI impact reports covering employment, bias, and data provenance. Treat AI governance as a board-level strategic priority, not IT compliance.	Audit all AI applications against the EU AI Act's high-risk classifications before August 2026 deadline. Establish AI ethics oversight committees with employee and civil society representation. Create internal reskilling academies for roles at displacement risk. Publish model cards for all internally deployed AI systems.	Reduced liability. Talent retention through demonstrated responsible practice. Regulatory compliance advantage. Customer trust in markets where AI governance is becoming a purchasing criterion (financial services, healthcare, public sector).

07 THE ACTION AGENDA | RECOMMENDATIONS BY STAKEHOLDER

	Why It Matters	What Must Change	How To Act	So What? (Outcome)
ACADEMIA	Universities and research institutions generate the evidence base that governance depends on, train the professionals who will deploy and oversee AI, and create the open-weight models that can distribute AI capability without concentrating power.	Build and publish open, demographically representative training datasets. Conduct longitudinal research on AI's cognitive, labour, and democratic impacts. Train the next generation of AI ethicists, governance specialists, and AI-literate professionals across all disciplines. Engage policymakers with actionable evidence.	Establish AI Governance Centres at leading universities as permanent infrastructure (analogous to nuclear safety research post-1945). Publish training data openly under Creative Commons. Mandate AI literacy modules across all degree programmes. Create interdisciplinary AI ethics programmes integrating humanities, social science, law, and computer science. Partner with governments on accredited knowledge base development.	An evidence-informed governance ecosystem. Professionals equipped to evaluate AI outputs critically rather than accepting them as authoritative. Training corpora that reflect humanity's full diversity rather than English-language internet dominance.

THE CHOICE IS NOT BETWEEN AI AND NO AI IT IS BETWEEN GOVERNED AND UNGOVERNED AI

This research has presented a large body of evidence. The headline finding is simple: artificial intelligence is not a threat to be feared or a savior to be celebrated. It is the most powerful general-purpose technology in human history, arriving in an institutional vacuum, with governance frameworks developing at the pace of democratic deliberation while capabilities advance at the pace of machine computation.

The five historical precedents examined — the printing press, the Industrial Revolution, the telegraph, social media — share a common structure: transformative technology, unchecked by governance for one to eight decades, generating concentration of power and distributional harm before institutional response eventually arrived. AI is following this pattern at 10-100x the speed. The feedback loop from capability to consequence is measured in months, not generations.

The 50% probability 'Captured Acceleration' scenario — where AI's benefits concentrate among the AI-fluent 5% while the 95% experience displacement without structured support — is not inevitable. It is the current default trajectory. Changing it requires neither technophobia nor techno-utopianism. It requires governance: the deliberate, evidence-based design of rules, institutions, and incentives that shape how powerful technology is deployed.

"Past technological transitions — from writing and the printing press to industrialization and the internet — were disruptive, yet they ultimately became governable as new norms, standards, and institutions emerged. The question AI poses is whether we can compress the governance timeline." — arxiv.org, *Three Lenses on the AI Revolution*, October 2025

The Three Non-Negotiables

1. No model that cannot disclose the provenance, consent status, and demographic representation of its training data should be deployed in high-stakes domains. This is not a restriction on innovation. It is a prerequisite for trustworthiness. **Data transparency before capability deployment.**
2. In healthcare, criminal justice, financial lending, weapons targeting, and democratic processes, human review of AI-generated outputs is not optional. The efficiency argument for removing human oversight in these domains is structurally identical to the efficiency arguments that produced every major governance failure in modern history: financial regulation, food safety, pharmaceutical approval, environmental protection. **Human oversight before autonomous action.**
3. The productivity gains of AI must be distributed through investment in AI literacy, reskilling infrastructure, and open-weight model ecosystems — or the net gain of 78 million jobs will be celebrated by the 5% who needed it least. **Equitable access before concentrated power.**

Final finding: The 5% of workers who are AI-fluent today earn 4.5x more than the 95% who are not. This single statistic is the most important number in the AI landscape more important than benchmark scores, model parameters, or quarterly revenue. It is the number that will determine whether the AI era is remembered as the decade when humanity finally distributed the gains of its most powerful technology, or as the decade when it failed, again, to govern what it had built.

SOURCES & REFERENCES

- Stanford HAI — Artificial Intelligence Index Report 2025; “Stanford AI Experts Predict What Will Happen in 2026” (January 2026)
- McKinsey Global Institute — Workforce transformation projections; 2.4-4.1% GDP increase estimates; 78% AI adoption rate (2025)
- Goldman Sachs — “What to Expect from AI in 2026” (January 2026); 7% GDP increase after 10 years; labour market analysis
- World Economic Forum — Future of Jobs Report 2025: 170M new jobs, 92M displaced; Workforce Blueprint 2026; Global Risks Report 2026
- OECD — AI Principles (2019, updated 2024); “Governing with Artificial Intelligence” (2025, 200 government use cases); “Thriving with AI” horizontal project
- UNESCO — Recommendation on the Ethics of Artificial Intelligence (2021, endorsed by 194 member states)
- US Council of Economic Advisers — “Artificial Intelligence and the Great Divergence” (January 2026 White House Report)
- J.P. Morgan Global Research — “AI’s Impact on Job Growth”; mildly negative correlation between employment and AI usage
- Epoch AI — Data on AI Models (3,200+ notable models, updated March 15, 2026); tracking of large-scale models
- Nature Machine Intelligence — Data Provenance Initiative: 1,800+ dataset audit; licence omission rate >70%; error rate >50%
- Gerlich, M. (2025). AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking. *Societies*, 15(1), 6. DOI: 10.3390/soci15010006 (n=666)
- MIT Media Lab — Kosmyrnq, N. (2025). Your Brain on ChatGPT — EEG study showing reduced neural engagement
- MDPI Systems (2025) — It’s Scary to Use It, It’s Scary to Refuse It: Psychological Dimensions of AI Adoption (DOI: 10.3390/systems13020082)
- Springer Nature (2025) — Digital Psychology: Conceptual Impact Model and the Future of Work
- Council on Foreign Relations — “How 2026 Could Decide the Future of Artificial Intelligence” (January 2026)
- RAND (1998) — “The Information Age and the Printing Press: Looking Backward to See Ahead” (P-8014)
- arxiv.org (2025) — “Three Lenses on the AI Revolution: Risk, Transformation, Continuity” (October 2025)
- diginomica (April 2025) — “AI is a Cognitive Revolution: Why History May Not Repeat Itself”
- WEF Davos (2024) — “What is the Gutenberg Parenthesis and how does studying 500 years of the printing press help us tackle the era of AI?”
- Menlo Ventures — 2025 State of Generative AI in the Enterprise: \$37B spending, 10+ products at \$1B+ ARR
- AmplifAI — 90+ Generative AI Statistics 2026: 47% enterprise users made major decisions based on hallucinated content
- Ragenaizer AI Job Market Dashboard (March 2026): 5% AI-fluent, earning 4.5x more; 56% salary premium
- Challenger, Gray & Christmas (2026): 55,000 job cuts directly attributed to AI in 2025
- PwC — \$15.7 trillion AI contribution to global GDP by 2030; 8-15% GDP increase scenarios
- Gartner — 40% enterprise AI agent adoption by end 2026; 15% autonomous decisions by 2028; 80% customer service resolution by 2029
- EU AI Act (August 2025) — Binding governance framework; prohibited practices active February 2025
- NIST AI Risk Management Framework (2023) — Practical governance structure for AI lifecycle
- ISO/IEC 42001 — International standard for AI governance systems certification
- IEEE 7000-2021 — Standard for ethical AI system design; METR research: AI task length doubles every 7 months
- Bradley Law (2025) — “Global AI Governance: Five Key Frameworks Explained”
- OECD GPAI — Futures of Global AI Governance: Scenarios for inclusive and fragmented trajectories
- Data Provenance Initiative — Consent in Crisis (2025): robots.txt bypassing documented; 5% of tokens restricted within one year

Intelligence Amplified, Judgement at Risk.

Research Report on Artificial Intelligence, Human Behavior, Governance, and the Next Decade

KEY CONTACTS

Hina Nasir

hina@sustainadility.com

Faraz Khan, MBE

faraz@sustainadility.com

Sajjeed Aslam, FCA

Sajjeed@sustainadility.com

sustainadility.com

Uol.edu.pk/continuum