

Research Article

The Architecture of Stability: Historical Continuity, Institutional Preparedness over Panic, and Strategic Contingency in the Age of Artificial Intelligence

Partha Majumdar* 

Department of Computer Science, Swiss School of Business and Management (SSBM), Geneva, Switzerland

Abstract

The rapid advancement of Artificial Intelligence (AI) has created a polarised public response, oscillating between utopian visions and existential fear. This analysis refutes these extremes, arguing that such anxiety is not a new phenomenon but a recurring historical cycle associated with the externalisation of human faculties. By examining historical precedents, the text reframes the AI challenge as a manageable variable rather than an uncontrollable force. The Socratic critique of writing, which warned of cognitive atrophy from externalised memory, is presented as an ancient parallel to modern concerns about AI fostering superficial competence. Similarly, the 19th-century Luddite movement is reinterpreted not as an irrational opposition to technology, but as a rational response by skilled artisans to the deskilling of their labour and the degradation of product quality—an analogue to fears that AI will devalue human expertise. Further parallels from the 20th-century "calculator wars" in education illustrate how curricula shifted from rote computation to higher-order problem-solving. These historical examples collectively argue that technological disruptions compel a redefinition of human competence and create new opportunities, rather than simply leading to cognitive or economic decline. Building on this historical perspective, the analysis proposes that the antidote to technological anxiety is a disciplined, managerial approach grounded in rigorous contingency planning. Through corporate case studies—contrasting Kodak's strategic inertia with Fujifilm's successful diversification, alongside the proactive pivots of Intel and Netflix—the text illustrates that organisational resilience depends on the ability to critically reassess core capabilities and cannibalise legacy models when necessary. This evidence informs a comprehensive, tiered risk management framework designed for professionals and institutions. "Plan A" focuses on Mitigation and Integration, advocating a "Centaur" model in which humans and AI collaborate as distinct but complementary agents, preserving human authority and using cognitive forcing functions to prevent automation complacency. "Plan B," a Contingency and Diversification strategy, involves developing multi-skilled, "M-shaped" professionals and cultivating an "analogue hedge" in roles requiring physical presence or legal accountability. Finally, "Plan C" outlines a strategy for Resilience and Sovereignty, serving as a last resort against systemic failure through the development of technological sovereignty via locally controlled AI and the willingness to execute a radical pivot to entirely new sectors. The overarching thesis is that by establishing these explicit contingency plans, organisations can transform existential risk into a manageable operational procedure, navigating the AI revolution with strategic preparedness rather than reactive panic.

Keywords

Historical Cycles of Technological Anxiety, Institutional Preparedness over Panic, Strategic Contingency Framework, Centaur Model, Capability Reassessment

*Correspondence: Partha Majumdar (partha.majumdar@hotmail.com), Partha Majumdar (partha.majumdar@icloud.com)

Received: 22 January 2026; **Accepted:** 30 January 2026; **Published:** 9 February 2026



Copyright: © The Author(s), 2026. Published by Science Publishing Group. This is an **Open Access** article, distributed under the terms of the Creative Commons Attribution 4.0 License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution and reproduction in any medium, provided the original work is properly cited.

1. The Genealogy of Technological Anxiety

To manage present risks effectively, it is essential to analyse past fears. The anxiety that Artificial Intelligence may render human cognition redundant or undermine labour dignity can be understood as a recurring psychological pattern, a recurrence relation seen throughout technological history. Every significant technological leap that externalised functions once done by humans has caused widespread social anxiety: the fear that the tool will weaken or diminish the user.

1.1. The Phaedrus Warning: The Externalisation of Memory and Wisdom

The roots of AI (Artificial Intelligence) anxiety can be traced back to 5th-century BCE Athens. In Plato's *Phaedrus*, we encounter one of the earliest and most influential philosophical critiques of mediated knowledge. The dialogue describes Socrates telling the myth of Theuth, the Egyptian god of invention, who offers the gift of writing to King Thamus. Theuth praises writing as a "recipe for memory and wisdom," a means for humans to preserve extensive knowledge outside the biological brain. [1].

King Thamus, however, dismisses this techno-optimism, issuing a warning that has been widely invoked in contemporary discussions of large language models (LLMs). He contends that writing does not enhance wisdom but instead erodes the innate capacity of memory.

Plato (1997, 274c-275b) reports Thamus's response:

"If men learn this, it will implant forgetfulness in their souls; they will cease to exercise memory because they rely on that which is written, calling things to remembrance no longer from within themselves, but by means of external marks."

Socrates elaborates on this point, suggesting that writing offers only an illusion of wisdom rather than genuine understanding. He warns that students may become mere listeners to many things without proper instruction, seeming knowledgeable to others while largely remaining ignorant and difficult to get along with, all the while developing a misplaced sense of epistemic confidence.

This ancient critique reflects the modern phenomenon known as the "Google Effect" or "Digital Amnesia," whereby individuals tend to remember where *information* is located rather than the information itself. [2] In contemporary AI-mediated contexts, this concern becomes even more pronounced: if a large language model (LLM) can produce an essay, legal brief, or software function, users may avoid the cognitive effort of synthesising that output themselves. Currently, critics argue-akin to Socrates-that this results in superficial competence, a "semblance of expertise" that proves fragile under critical scrutiny. [3].

However, the historical trajectory of writing offers a counter-narrative to Thamus's fear. While it is true that the oral tradition (the capacity to recite epic poetry from memory) declined, the invention of writing did not lead to a net loss of human intelligence. Instead, it enabled cognitive offloading (i.e., the externalisation of informational burdens), freeing the human mind to engage in higher-order analysis, cross-referencing, and the development of complex sciences that would be impossible to maintain in working memory alone. [4] The "risk" Socrates identified was real-the loss of mnemonic capacity-but the "contingency plan" of civilisation was to shift the definition of intelligence from *retrieval* to *synthesis*.

1.2. The Luddite Rationality: Quality, Dignity, and Labour Rights

Moving from the cognitive to the economic, the Luddite movement of the early 19th century stands as the primary historical analogue for fears of AI-induced job displacement. Popular culture has reduced the Luddites to irrational, anti-technology vandals. This caricature reflects a common misrepresentation of Luddite motivations. Historical analysis reveals that the Luddites were highly skilled artisans of their day-weavers, knitters, and croppers-who were not opposed to machinery in principle but to the *social contract* governing its use. [5].

The Luddite rebellion (1811–1816) was triggered not by the invention of the stocking frame or the shearing frame, but by the deployment of these machines to produce "cut-ups"-inferior, mass-produced textile goods that bypassed established quality standards and apprenticeship requirements. [5] The Luddites argued that this technology was being used to de-skill the workforce and drive wages below the subsistence level, effectively "robbing" the worker of their trade's dignity.

Binfield (2004) argues that the Luddites' demands centred on preserving product quality, apprenticeship systems, and wage standards, rather than opposing machinery per se.

The relevance to contemporary AI debates is significant. Modern professionals do not fear the algorithm because it is efficient; they fear the "cut-up" effect-the proliferation of "good enough" AI-generated content (code, copy, art) that de-values the meticulous, "artisanal" work of the human expert. [6] The Luddites essentially formulated a "Plan A" (petitioning Parliament for minimum-wage and quality-of-life laws), which was ignored. Lacking a "Plan B" (social safety nets or alternative employment in a war-torn economy), they resorted to an escalation to machine breaking, which led to military intervention and capital punishment. [6].

For contemporary organisations, this history suggests that resistance to technological systems is often a rational response

to the threat of *quality degradation* and *economic disenfranchisement*. Successful integration of AI requires proactively addressing these "Luddite" concerns—ensuring that the tool is used to augment the expert, not to replace them with a "cut-up" equivalent. [7].

1.3. The Calculator Wars: Re-baselining Competence

In the 1970s and 1980s, a significant pedagogical debate arose over the integration of handheld electronic calculators into mathematics education. This "Calculator War" can be understood as anticipating the current discussions surrounding AI in schools and workplaces. Math teachers and parents worried that the device would lead to a generation of "calcuholics"—students addicted to machines, lacking the ability to do mental math, and missing out on numerical intuition. [8].

Contemporary educators expressed concern that students increasingly accept calculator outputs uncritically, failing to recognise implausible results. [8] This observation foreshadows current issues with AI hallucination and automation bias, in which users tend to accept machine outputs as correct even when they are logically incorrect. [9].

However, the introduction of the calculator did not eliminate mathematics. Instead, it prompted a *reassessment of competence*, shifting the curriculum's emphasis from *manual computation*—such as long division and deriving square roots by hand—to understanding and *conceptualising* problems. The approach favoured by educators, known as "Plan A," was to integrate the calculator as a tool: it would handle the tedious calculations while increasing the focus on developing more rigorous problem-solving skills. [10].

This historical shift indicates that AI will not eliminate the need for human thinking but will elevate the importance of higher-level cognitive skills. Human expertise is likely to be valued less for routine execution and more for problem formulation and evaluative judgement. The shift can be understood as moving from writing the syntax of a Python loop (the calculation) to their capacity to understand *why* the loop is necessary and whether it is executing the correct logic (the formulation). [11].

1.4. The Printing Press: A Double-edged Disruption

The advent of the Gutenberg printing press in the 15th century provides a broad perspective on how information technology transforms labour and society. In ways analogous to contemporary AI systems, the press was a "generative" technology, enabling mass reproduction of information at almost no additional cost.

The short-term disruption initially led to the collapse of the scribal profession, as monks and professional copyists saw their value decline sharply. Additionally, it increased the cir-

culation of polemical, often unreliable pamphlets, which contributed to religious conflicts during the Reformation. [4].

However, the long-term impact was the creation of entirely new economies and the democratisation of knowledge. It broke the elite's monopoly on information, bolstered the middle class, and standardised languages. [12] For the "manager" of the 15th century (the Guild Master or the Bishop), the "Plan A" of banning books failed. The successful "Plan B" was to utilise the press to disseminate one's own ideas more effectively than the opposition.

Historical scholarship broadly agrees: technology eliminates certain *tasks* and *roles*, frequently causing significant disruption to the existing workforce (such as scribes, weavers, and calculators). Nevertheless, it seldom eradicates the entire *domain*. Instead, it broadens the domain by reducing entry barriers and boosting output volume. [12].

2. The Corporate Adaptation - Case Studies in Adaptation

The focus is on a managerial strategy: recognising risks and developing contingency plans. Corporate history offers the most detailed data for this, since companies face continuous threats of technological obsolescence. The key factor determining survival versus bankruptcy is often the strength of their "Plan B."

2.1. The Divergent Paths of Film: Kodak vs. Fujifilm

The fall of Eastman Kodak and the continued success of Fujifilm serve as a prime example of why contingency planning is crucial during a technological paradigm shift. Both firms encountered a similar crisis: the advent of digital photography, which made their main source of profit—silver halide film—obsolete.

Kodak: The Failure of Plan A

Kodak's decline cannot be attributed to a lack of technological innovation; the firm developed the digital camera in 1975. The firm's failure stemmed from managerial cognition and strategic framing. Their "Plan A" was to defend the high-margin film business—the "razor and blade" model in which cameras were inexpensive, but film and processing were profitable—for as long as possible. Senior management framed digital imaging primarily as a threat to the firm's established film-based business model, a cognitive commitment that inhibited strategic reorientation despite early technical competence. Even when they finally shifted to digital, it was too late, and their attempt to mimic the consumable model through printer cartridges failed to recognise that the new ecosystem centred on sharing images, not printing. They lacked a "Plan B" that was independent of their legacy identity. [13].

Fujifilm: The Triumph of Plan B and C

Fujifilm's leadership identified the crisis early and crafted

a multi-layered strategy, offering a useful analogue for contemporary AI-related disruption.

Fujifilm initially adopted a mitigation approach, leveraging residual value from its declining film business to fund organisational transition rather than immediately abandoning the legacy operation (Plan A). Simultaneously, the company invested heavily in digital imaging technologies to stay competitive, even though it recognised that profit margins in the digital market would be inherently lower than those of photographic film (Plan B). Most importantly, Fujifilm implemented a bold diversification strategy that extended beyond digital substitution (Plan C). Instead of defining itself solely by its end products, the company redefined its core identity based on its fundamental capabilities, aligning with models of dual business-model competition and capability reconfiguration. [14].

This capability-focused shift highlighted that the production of photographic film relied on transferable expertise in chemicals, materials, and processes. Specifically, the gelatin used in film is derived from collagen, a key protein related to human skin elasticity; methods for preventing film degradation closely align with those used to inhibit oxidative ageing; and the company's skill in microscopically layering chemical compounds directly correlates with its nanotechnological competencies. By repurposing these capabilities, Fujifilm successfully expanded into adjacent fields, notably launching the Astalift cosmetics line and branching into pharmaceuticals and healthcare-related materials. [15].

Managerial insight: Fujifilm's survival came from their "Plan C," which involved reducing their core identity to its chemical "source code" and reworking it for a different industry. This case suggests that adaptation to AI-driven disruption may hinge less on preserving specific occupational roles than on institutional capacities to reinterpret and redeploy underlying capabilities across domains, a pattern that may generalise beyond corporate settings.

2.2. Intel: Constructive Paranoia and the "Walk Away"

In 1985, Intel faced a crisis that threatened its existence. The company was founded on memory chips (DRAM), but Japanese competitors were dumping chips at prices Intel could not match. The "Plan A" of manufacturing better memory chips was failing.

Grove later described a pivotal moment in which Intel's leadership recognised that a rational external successor would immediately exit the memory business, prompting a deliberate decision to enact that break internally rather than wait for displacement. [16].

Intel executed a decisive "Plan B" that became its new "Plan A": it abandoned its core business (Memory) and reoriented the firm around a then-secondary product, the microprocessor (CPU). This required closing factories and laying off thou-

sands, but it saved the company. Grove characterised this approach as 'constructive paranoia', arguing that organisations must anticipate strategic inflexion points and be willing to cannibalise their own core businesses to survive them. [16, 17].

2.3. Netflix: Cannibalisation as Strategy

Netflix exemplifies a modern case of intentional self-cannibalisation as an organisational adaptation. Its leadership recognised early on that the DVD-by-mail business, though highly profitable, had a limited lifespan as internet bandwidth improved and streaming technologies developed. The company's initial approach focused on optimising its current logistics and distribution systems to maximise value from the DVD model while it was still viable (Plan A).

More importantly, Netflix adopted a strategic transition approach by introducing its streaming service as a free enhancement for existing DVD subscribers (Plan B). Instead of safeguarding its traditional business, this strategy actively motivated customers to shift from DVDs to streaming. In doing so, Netflix effectively cannibalised its own high-margin operations to establish an early, defensible position in the digital market, which at the time was characterised by lower margins and uncertainty. [18, 19] This approach emphasised long-term platform dominance over immediate revenue maximisation.

Anticipating a potential additional structural weakness, Netflix implemented a second strategic shift by heavily investing in original content creation (Plan C). Recognising that major content owners would eventually regain licensing rights and pose direct competition, the company aimed to reduce its reliance on external intellectual property by developing its own proprietary content, starting with prominent productions such as *House of Cards*. This transition transformed Netflix from a distribution platform into an integrated content producer, reducing supply-chain risks and strengthening its strategic independence. [20].

2.4. Nokia: The Burning Platform

Nokia's decline serves as a key example of organisational failure during major technological shifts. In 2007, the company held about 50% of the global mobile handset market; by 2013, this had dropped to under 5%. Senior leaders clearly acknowledged the strategic challenge posed by the shift from hardware-focused phones to software-based smartphone ecosystems, as outlined in the influential 2011 "Burning Platform" memo.

Although this recognition existed, action was slow to materialise. Nokia's response was influenced by sustained overconfidence in its hardware engineering and a fragmented organisational focus, which hindered the development of a timely ecosystem amidst competition from Apple and Android. The company lacked a strong alternative platform strategy, and its decision to sell its mobile handset division to Microsoft was more an exit than an adaptive measure. [21].

2.5. Nintendo: Radical Identity Exploration

Nintendo exemplifies a significant shift in organisational strategy. Established in 1889 as a producer of Hanafuda playing cards, it faced stagnation in the mid-20th century due to declining demand for traditional card games. To adapt, Nintendo embarked on various diversification initiatives into unrelated sectors such as transportation, food manufacturing, and hospitality, though many of these ventures were not commercially successful.

A significant breakthrough came from experimenting with electronic toys, especially the *Ultra Hand*, which helped transition into arcade games and later the home console market. Instead of merely expanding or protecting its original product area, Nintendo gradually reshaped its organisational identity through a process of trial and error outside its core competencies. This case indicates that, in some situations, organisational resilience might rely more on the ability to explore broadly and experiment during times of structural uncertainty than solely on optimising current capabilities. [22].

3. The Modern Risk Landscape

To effectively apply these lessons from history to the present, institutions and decision-makers must carefully assess and understand the landscape of the ongoing disruption caused by artificial intelligence. Unlike the risks faced during the industrial revolution, which primarily affected physical labour and muscle power, the current challenges are focused on cognitive functions, mental processes, and decision-making capabilities. Therefore, organisations need to consider these distinct aspects as they navigate the complexities of AI-driven changes in the modern landscape.

3.1. Cognitive Risks: The "Glass Cockpit" and the Jagged Frontier

The psychological risk of AI is automation complacency. Research by Ethan Mollick and others on the "Jagged Frontier" of AI capabilities reveals a critical danger zone. AI is excellent at some tasks (e.g., brainstorming, coding boilerplate) but unreliable at others (e.g., precise legal citation, nuanced ethical judgment). [23].

When users rely on AI for tasks beyond its competence, performance drops below the baseline of a human working alone. This is analogous to the "Glass Cockpit" problem in aviation: pilots who rely too heavily on autopilot can lose situational awareness and the manual skills required to recover the plane in an emergency. In the office, this manifests as the "calculolic" employee who accepts an AI-generated strategy without noticing the "absurd answer" buried in the logic. [9].

Furthermore, there is a risk that cognitive offloading will lead to atrophy. If the friction of work (the struggle to find the right word, the effort to debug code) is removed, the neural

pathways associated with deep learning may weaken. Neuroscientific research suggests that sustained effort and challenge contribute to myelination and skill acquisition; persistent removal of such effort may therefore undermine the development of deep expertise. [24].

3.2. Economic Risks: The White-collar Iceberg

The economic risk associated with technological advancement is often described as the "hollowing out" of the white-collar workforce. This phenomenon refers to the diminishing of middle- and upper-middle-level professional jobs as automation and AI become more prevalent. The "Iceberg Index" provides insight into this issue by indicating that, although the level of visible AI adoption appears quite high in the tech industry, the true extent of AI's impact is significantly broader than what is reflected in direct adoption metrics. Specifically, submerged or hidden exposure to AI and automation risks predominantly exists in sectors such as administrative services, the financial industry, and various professional services, including legal, consulting, and other specialised fields. These sectors, which often operate behind the scenes, are increasingly vulnerable to automation-driven disruptions that are not immediately apparent. This divergence between visible AI deployment and the hidden, yet significant, exposure underscores the importance of recognising the broader implications for the economy as a whole [25].

Unlike earlier technological shifts that disproportionately affected physically demanding or routine occupations, contemporary AI systems increasingly target roles traditionally regarded as secure and high-status. The potential risk associated with this trend does not necessarily point to widespread unemployment but rather to increased operational fragility and a breakdown in the traditional career progression ladder. Specifically, if AI systems become capable of performing tasks typically done by junior professionals, organisations might reduce or even eliminate hiring entry-level workers, thereby constraining the pipeline through which expertise is developed. Historical analysis of prior technological revolutions suggests that such disruptions primarily manifest as breakdowns in occupational progression systems when institutional training and credentialing fail to adapt to changing task demands, rather than as persistent joblessness. [26].

4. The Institutional Risk Framework - Plan A, Plan B, Plan C

Building on the historical patterns, case studies, and analyses of cognitive and economic risk dynamics from Sections 1–3, this section introduces an institutional risk framework with three contingency layers: Plan A, Plan B, and Plan C. This multi-layered system offers defensive and adaptive responses suitable for organisational, professional, and policy environments. Its goal is not to forecast precise results but to

establish a structured approach to handling uncertainty, safeguard institutional resilience, and support decision-making amid rapid technological change.

4.1. Plan A: Mitigation and Integration (the "Centaur" Strategy)

Plan A assumes that the potential risks associated with AI are present; however, these risks can be effectively managed and controlled through strategic adaptation efforts. The primary objective of this approach is to maximise the usefulness and benefits provided by the AI tool, ensuring it serves its intended purpose to the greatest extent possible. At the same time, it is crucial to implement measures that mitigate specific dangers, such as hallucinations—instances in which the AI generates inaccurate or nonsensical outputs—and cognitive atrophy, the gradual erosion of human evaluative and reasoning capacities when tasks are excessively offloaded to automated systems.

4.1.1. Strategy: The Centaur Model

For high-stakes cognitive and organisational tasks, professional communities should preferentially adopt the *Centaur Model* of human–AI collaboration rather than the *Cyborg Model*, which seeks seamless integration of human and machine cognition. In the Centaur configuration, humans and AI systems operate as distinct but complementary agents, with clearly delineated roles and decision rights. Crucially, this model preserves human epistemic authority and accountability in contexts where errors carry material, ethical, or strategic consequences.

Operationally, the Centaur Model can be implemented by structuring workflows into alternating phases of AI-supported generation and human-led evaluation. During AI-dominant phases, generative systems are employed for tasks such as producing alternative solution paths, drafting preliminary content, or identifying potential optimisations. In subsequent human-dominant phases, individuals must assess, select, verify, and contextualise AI outputs in light of domain knowledge, organisational norms, and situational judgment. This enforced alternation prevents the gradual erosion of human agency that can occur when AI outputs are consumed passively or integrated without reflection.

The value of this separation lies in its capacity to counteract what research on human–automation interaction identifies as *automation complacency*. When AI systems operate with uneven reliability—performing exceptionally well in some domains while remaining fragile or misleading in others—users may overgeneralise trust and fail to detect errors. By obligating explicit transitions between generative and evaluative modes, the Centaur Model introduces a structural pause that restores critical distance between the human decision-maker and the machine-generated artefact.

From a governance perspective, this approach aligns closely with established principles of *human-in-the-loop*

(HITL) system design. Rather than treating HITL as a nominal safeguard, the Centaur Model embeds it as a procedural requirement, ensuring that responsibility for judgment, verification, and final action remains unambiguously human. Empirical work on generative AI adoption suggests that such structured collaboration not only reduces error propagation but also supports learning and skill retention among users, mitigating the risks of cognitive atrophy associated with excessive offloading.

In this sense, the Centaur Model should be understood not as a rejection of advanced automation, but as an institutional architecture for its disciplined use—one that transforms AI from an opaque decision substitute into a controlled cognitive instrument. By preserving moments of human deliberation and evaluation, organisations can harness AI's generative strengths without surrendering the interpretive and ethical capacities that remain uniquely human.

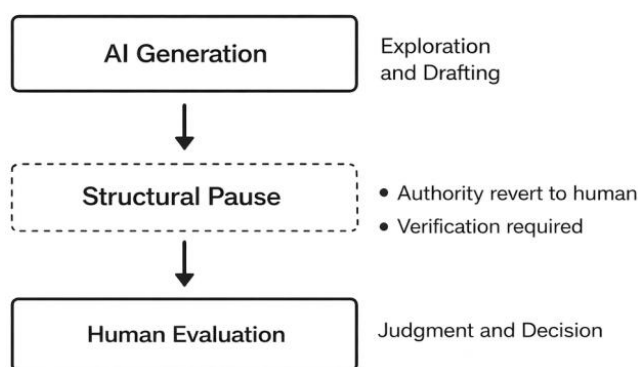


Figure 1. The Centaur Model of Human-AI Collaboration.

4.1.2. Tactics: Prompting as Computational and Analytical Reasoning

Organisational training in AI should shift focus away from superficial conceptions of “prompt engineering” as the identification of optimal or formulaic commands. Instead, effective interaction with generative AI systems can be understood as a form of computational and analytical reasoning. When users employ advanced prompting techniques, such as chain-of-thought prompting, they externalise and structure their reasoning in ways that align with established critical problem-solving practices.

In practice, this involves decomposing complex problems into explicit logical steps, articulating constraints and assumptions, and iteratively refining outputs in response to feedback. These actions extend beyond narrow technical skills and reflect fundamental cognitive capacities such as logical reasoning, abstraction, and communicative clarity. From this perspective, prompting serves as an indicator of how well a user understands the underlying problem being addressed.

Accordingly, organisations may view proficiency in prompting not merely as a technical skill, but as an indicator

of broader analytical and communicative capability. Difficulties in formulating effective prompts often reflect gaps in problem understanding rather than insufficient familiarity with the AI system. Training programmes that emphasise this interpretive and reflective dimension of human–AI interaction may strengthen, rather than erode, core cognitive skills while mitigating the risks of uncritical reliance on automation.

4.1.3. Control Mechanism: Cognitive Forcing Functions

To mitigate the risk of automation complacency associated with highly reliable but imperfect AI systems, institutions can introduce *cognitive forcing functions*—deliberate procedural constraints designed to interrupt uncritical reliance on automated outputs. Such mechanisms are well established in safety-critical domains, where they function as intentional “procedural interruptions” that compel human re-engagement at key decision points.

One such mechanism is a verification constraint requiring AI-generated factual claims to be independently corroborated by multiple primary sources before inclusion in formal outputs. By obligating users to seek external confirmation, this rule counteracts the tendency to treat machine-generated information as self-validating and reinforces standards of epistemic diligence.

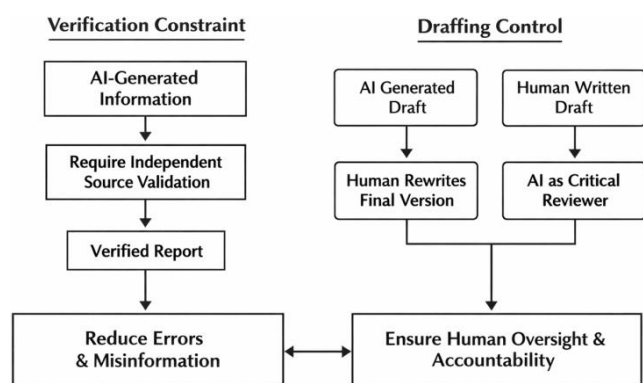


Figure 2. Cognitive forcing functions as institutional safeguards against automation complacency.

A complementary control involves structuring the drafting process to ensure active human reinterpretation of AI outputs. In one configuration, generative systems may produce an initial exploratory draft, which is then substantively rewritten by a human author prior to dissemination. In another, the human produces the primary draft while the AI system assumes a critical or adversarial role, identifying potential weaknesses, omissions, or inconsistencies. In both cases, the key principle is that AI-generated text should not be transmitted directly to decision-makers or stakeholders without human reformulation.

These forcing functions serve a dual purpose. First, they reduce the probability that errors or hallucinations propagate un-

checked through organisational processes. Second, they preserve human situational awareness and accountability by ensuring that users remain cognitively engaged with both the content and the reasoning underlying final decisions. Empirical research on automation bias and human–machine interaction suggests that such structured interventions are effective in sustaining vigilance and mitigating over-trust in automated systems.

4.2. Plan B: Contingency and Diversification (the "Fujifilm" Strategy)

Plan B is put into action whenever Plan A proves to be inadequate or insufficient. This typically happens when a particular role or task is essentially undervalued or dismissed due to automation, making it impossible for Plan A to handle it effectively. In such cases, the strategy involves capability diversification to manage potential risks and uncertainties.

4.2.1. Strategy: Capability Auditing and "M-Shaped" Skills Development

Analogous to the capability audit undertaken by Fujifilm during its strategic reorientation, individuals facing AI-driven disruption may benefit from systematically reassessing their own cognitive and professional capabilities. Rather than focusing narrowly on task automation risk, such an audit foregrounds the relationship between human skills and areas where AI systems remain limited or unreliable.

Existing research suggests that certain categories of human capability remain complementary to, rather than substitutable by, current AI systems. These include forms of complex communication, empathy and emotional intelligence, strategic negotiation, and tasks involving physical or spatial interaction with the environment. This pattern is consistent with observations often grouped under Moravec’s Paradox, which highlights that AI systems tend to outperform humans in abstract, formal domains while struggling with embodied, contextual, and socially nuanced activities.

From this perspective, individual adaptation strategies may involve a shift away from narrowly specialised “T-shaped” skill profiles—characterised by depth in a single domain—toward “M-shaped” skill configurations, in which individuals cultivate multiple areas of deep expertise. Such diversification does not imply superficial breadth, but rather the development of additional pillars of competence that draw on distinct cognitive, social, or embodied strengths.

For example, a software engineer whose work is increasingly automated may develop a second area of deep expertise in fields such as bioinformatics or hardware systems, where domain knowledge, physical constraints, or regulatory complexity limit full automation. Similarly, a professional writer facing displacement by generative text systems may cultivate a parallel specialisation in areas such as crisis communication or event coordination, where situational judgment, interper-

sonal sensitivity, and real-time decision-making remain central. In this way, individual resilience mirrors institutional strategies of capability redeployment, reducing dependence on any single task domain vulnerable to automation.

4.2.2. Tactic: The "Analogue Hedge"

A further adaptive tactic at the individual level involves what may be described as an *analogue hedge*: the strategic cultivation of skills and roles that require physical presence, embodied verification, or legally attributable human judgment. While AI systems can generate analyses, recommendations, or simulations, they currently lack the capacity to perform physical inspection or to assume legal responsibility for outcomes in high-stakes contexts.

Many professional roles retain durability precisely because they require a form of human attestation—such as a physical inspection, certification, or legally binding signature—that anchors responsibility to an identifiable individual. In fields such as construction oversight, medical practice, engineering certification, or regulatory compliance, errors carry legal, ethical, and financial consequences that cannot yet be delegated to automated systems. As a result, these roles remain insulated from full automation longer than positions that operate entirely within digital or advisory domains.

From a risk-management perspective, investing in competencies that intersect with physical reality or legal accountability functions as a hedge against displacement. Such roles combine technical expertise with responsibility-bearing functions that are difficult to externalise to machines, thereby preserving human agency within socio-technical systems. This does not imply immunity from technological augmentation, but rather a slower and more constrained trajectory of substitution, shaped by institutional, regulatory, and liability frameworks.

4.3. Plan C: Resilience and Sovereignty (the "Intel/Prepper" Strategy)

Plan C serves as the emergency fallback—the last resort when all other options have failed or are unavailable. It is specifically designed to serve as a safeguard against the most serious and detrimental risks that could threaten the entire system or industry. These risks include a widespread, systemic failure that could disable or crash the entire network or infrastructure, censorship that might suppress essential information or innovations, and a complete collapse of the industry itself, where all operations, businesses, and services could cease entirely.

4.3.1. Strategy: Technological Sovereignty (Sovereign AI)

At the level of long-term contingency planning, institutions may need to consider the strategic risks associated with de-

pendence on highly centralised, externally controlled AI platforms. As generative AI capabilities become increasingly embedded in organisational workflows, reliance on third-party providers introduces vulnerabilities related to pricing volatility, data governance, service availability, and shifting policy or regulatory constraints. In such contexts, technological autonomy becomes a component of institutional resilience.

One response to these risks involves developing *technological sovereignty* through the deployment of locally controlled AI systems. Rather than relying exclusively on cloud-based or proprietary platforms, organisations can maintain the ability to run open-weight language models on privately managed infrastructure. This approach enables greater control over data flows, ensuring that sensitive or regulated information remains within institutional boundaries. In this configuration, AI systems function less as rented services and more as owned cognitive infrastructure, embedded within existing governance and compliance frameworks.

Operationally, this strategy may involve investment in local computational resources and the internal expertise required to deploy and maintain inference capabilities. While such infrastructure does not seek to replicate the scale or performance of frontier commercial systems, it provides redundancy and optionality. Much like backup power systems in critical facilities, locally controlled AI capacity functions as a safeguard against external disruption rather than a replacement for primary services.

Importantly, technological sovereignty should not be understood as a rejection of large-scale AI ecosystems, but as a hedging strategy within a broader risk-management framework. By retaining the option to operate independently when necessary, institutions preserve decision-making autonomy and reduce exposure to external shocks in an increasingly centralised AI economy.

4.3.2. Strategy: The Radical Pivot

In certain scenarios, neither incremental adaptation nor technological hedging may be sufficient to preserve institutional or individual viability. Historical cases of organisational survival—such as Nintendo's transition away from playing cards or Intel's exit from memory manufacturing—illustrate that long-term resilience can sometimes require a radical departure from an established industry or professional identity. Within the context of AI-driven disruption, this corresponds to a third-order contingency strategy: deliberate sectoral reorientation.

If particular segments of the knowledge economy are subject to intense automation pressure, leading to commoditisation and declining returns to expertise, continued participation may evolve into an "existential threat." Under such conditions, persistence within the same occupational domain may no longer represent rational adaptation but rather path dependence driven by sunk costs, credentials, or professional inertia. A radical pivot involves recognising when these accumulated investments no longer yield protective value and re-evaluating

alternative economic domains.

Potential destinations for such reorientation include sectors characterised by high levels of physical interaction, embodied skill, interpersonal trust, or situational judgment-areas where AI systems currently lack operational footholds. These domains, often described as part of the “high-touch” or materially grounded economy, encompass forms of work where human presence, accountability, and contextual responsiveness remain central. Importantly, this strategy does not imply regression or de-skilling, but rather a reconfiguration of expertise toward domains with different vulnerability profiles.

From a psychological and institutional perspective, executing a radical pivot requires the capacity to disengage from established identities and to tolerate uncertainty during transition. This form of adaptive calmness-distinct from optimism or denial-enables actors to prioritise long-term viability over short-term status preservation. As such, radical reorientation functions not as a failure of adaptation, but as its most extreme expression within a comprehensive risk-management framework.

Conclusion: Planning, Preparedness, and Institutional Composure

This paper has argued that identifying technological risk should serve as a catalyst for structured planning rather than reactive anxiety. Historical and contemporary evidence suggests that periods of rapid technological change consistently provoke fears of displacement and loss of agency. However, these responses become most destabilising when uncertainty is experienced as a lack of available options. By contrast, the development of explicit contingency structures transforms uncertainty into a domain of strategic deliberation.

Within this framework, institutional composure is not synonymous with ignorance or denial of risk. Rather, it emerges from preparedness: the capacity to anticipate disruption, articulate alternative courses of action, and retain decision-making agency in the face of uncertainty. A “calm mind,” understood in this sense, is not an affective state but an organisational and cognitive posture grounded in foresight and optionality.

The historical cases examined in this paper reinforce this distinction. Classical critiques of writing in Plato’s *Phaedrus* illustrate that new cognitive technologies are routinely perceived as threats to existing forms of knowledge, even as they enable novel modes of reasoning and intellectual organisation. The Luddite movement demonstrates that resistance to technology often reflects contestation over labour conditions and dignity rather than opposition to technological change per se. More recent corporate trajectories-such as those of Fujifilm and Intel-show that even existential disruptions can be navigated when institutions are willing to critically reassess core capabilities and pursue strategic reconfiguration rather than preservation at all costs.

These cases suggest that resilience in the age of artificial intelligence will depend less on predicting specific technological outcomes than on cultivating adaptive capacity across

cognitive, economic, and institutional dimensions. By embedding contingency planning into organisational and professional practice, actors can replace panic-driven responses with deliberative strategies, preserving agency amid accelerating technological transformation.

Abbreviations

AI	Artificial Intelligence
BCE	Before Common Era
CPU	Central Processing Unit
DRAM	Dynamic Random Access Memory
DVD	Digital Versatile Disc
HITL	Human-in-the-Loop
LLM	Large Language Model

Author Contributions

Partha Majumdar is the sole author. The author read and approved the final manuscript.

Conflicts of Interest

The author declares no conflicts of interest.

References

- [1] Plato (1997) *Phaedrus*. In: Cooper JM, Hutchinson DS (eds) Plato: Complete Works. Trans Nehamas A, Woodruff P. Hackett Publishing Company, Indianapolis, pp 475–526.
- [2] Sparrow B, Liu J, Wegner DM (2011) *Google effects on memory: Cognitive consequences of having information at our fingertips*. Science 333(6043): 776–778. <https://doi.org/10.1126/science.1207745>
- [3] Ong WJ (1982) *Orality and literacy: The technologizing of the word*. Routledge, London and New York.
- [4] Eisenstein EL (1980) *The printing press as an agent of change: Communications and cultural transformations in early modern Europe*. Cambridge University Press, Cambridge.
- [5] Binfield K (ed) (2004) *Writings of the Luddites*. Johns Hopkins University Press, Baltimore.
- [6] Thomis M (1970) *The Luddites: Machine-breaking in Regency England*. David & Charles, Newton Abbot.
- [7] Jones SE (2006) *Against technology: From the Luddites to neo-Luddism*. Routledge, New York and London.
- [8] Fey JT (1979) *Calculators and mathematics education*. National Council of Teachers of Mathematics, Reston, VA.
- [9] Ward AF, Duke K, Gneezy A, Bos MW (2017) *Brain drain: The mere presence of one’s own smartphone reduces available cognitive capacity*. Journal of the Association for Consumer Research 2(2): 140–154.

- [10] National Council of Teachers of Mathematics (1989) *Curriculum and evaluation standards for school mathematics*. National Council of Teachers of Mathematics, Reston, VA.
- [11] Autor DH (2015) *Why are there still so many jobs? The history and future of workplace automation*. Journal of Economic Perspectives 29(3): 3–30.
- [12] Febvre L, Martin H-J (1976) *The coming of the book: The impact of printing 1450–1800*. Verso, London.
- [13] Tripsas M, Gavetti G (2000) *Capabilities, cognition, and inertia: Evidence from digital imaging*. Strategic Management Journal 21(10–11): 1147–1161.
- [14] Markides C, Charitou CD (2004) *Competing with dual business models: A contingency approach*. Academy of Management Executive 18(3): 22–36.
- [15] Tripsas M (2009) *Technology, identity, and inertia through the lens of the digital photography revolution*. Organization Science 20(2): 441–460.
- [16] Grove A (1996) *Only the paranoid survive: How to exploit the crisis points that challenge every company*. Currency Doubleday, New York.
- [17] Burgelman RA (1994) *Fading memories: A process theory of strategic business exit in dynamic environments*. Administrative Science Quarterly 39(1): 24–56.
- [18] Tryon C (2015) *TV got better: Netflix's original programming strategies and the on-demand television transition*. Media Industries Journal 2(2): 104–116.
- [19] McDonald R, Eisenhardt KM (2020) *Parallel play: Startups, nascent markets, and effective business-model design*. Administrative Science Quarterly 65(2): 483–523.
- [20] Cusumano MA, Gawer A, Yoffie DB (2019) *The business of platforms: Strategy in the age of digital competition, innovation, and power*. Harper Business, New York.
- [21] Vuori TO, Huy QN (2016) *Distributed attention and shared emotions in the innovation process: How Nokia lost the smartphone battle*. Administrative Science Quarterly 61(1): 9–51.
- [22] Sheff D (1993) *Game over: How Nintendo conquered the world*. Random House, New York.
- [23] Mollick E, Mollick L, Noy S (2023) *What can we learn from experiments with generative AI?* Journal of Economic Perspectives 37(4): 3–30.
- [24] Firth J, Torous J, Stubbs B, et al. (2019) *The “online brain”: How the Internet may be changing our cognition*. World Psychiatry 18(2): 119–129.
- [25] Felten EW, Raj M, Seamans R (2018) *A method to link advances in artificial intelligence to occupational abilities*. AEA Papers and Proceedings 108: 1–4.
<https://doi.org/10.1257/pandp.20181021>
- [26] Goldin C, Katz LF (2008) *The race between education and technology*. Harvard University Press, Cambridge, MA.