



AI-Assisted Incompetence and Fraud in IT–Separating Real Talent from Artificial Credibility in Organizations

Dr.A.Shaji George

Independent Researcher, Chennai, Tamil Nadu, India.

Abstract – The perception of expertise and its reality have been completely transformed by generative artificial intelligence (AI). Nowadays, these systems are utilized by qualified specialists for research, composing, and assessing alternatives in a more comprehensive manner. Meanwhile, unskilled, and dishonest people can use the same tools to craft beautiful resumes, strategies, architectures, security policies, code, and interview responses that they don't know, and can't execute. This article explores how AI can be used to become incompetent and fraudulent at every level of the information technology hierarchy, from the developers and administrators to the chief information, technology, and security officers. It categorizes AI users based on government guidance on how to protect against ID fraud by remote workers, AI risk-management frameworks published by various organisations, international labour research, and empirical productivity studies, into three groups competent professionals using AI to improve their legitimate work, developing professionals using AI to build their skills, and individuals using deception to appear competent. It's only the third category that is the problem. The article examines the impact of artificial credibility on the individual, team, enterprise, economy, and societal levels and suggests ways to detect it, governance mechanisms to prevent it, due process safeguards to protect against it and a phased action plan to take place at the organizational level. It ends by saying that organizations need to move away from the focus on artifacts and presentation and towards verified identity, demonstrated reasoning, measurable outcomes, and personal accountability.

Keywords: Generative artificial intelligence, Competence verification, Identity fraud, IT governance, Cybersecurity workforce, Artificial credibility, Accountability, Talent assessment.

1. INTRODUCTION

The output of a believable intellectual product was a good indicator of competency for most of the history of the information technology profession. In almost all instances, the person who created a coherent enterprise architecture diagram, a well-structured security policy, a persuasive digital transformation strategy or a working block of code had to have both domain knowledge and experience, as well as the analytical discipline to organize the knowledge into a usable form. This correlation was used by the hiring managers, boards of directors, procurement committees, and clients. The amazing document was a proxy for the competent person, and although the proxy was never to be perfect, it was sufficiently reliable to underpin decades of professional evaluation practice.

Generative artificial intelligence has broken this correlation. Now, anyone with an internet connection can use the systems that provide fluent, organized and superficially authoritative text, diagrams, plans, and code. Anyone who doesn't really know about cloud migration can ask for a three-year cloud transformation plan and get one in a few minutes. A candidate that has no security experience can create a zero-trust roadmap, an incident-response policy, or a list of confidence interview answers. The

manager that is under pressure can create a presentation for the board that is full of the terminology of the day, but has no idea of the assumptions, risks, or requirements for implementation that are contained within. The artefact is still a representation of competence but no longer displays competence.

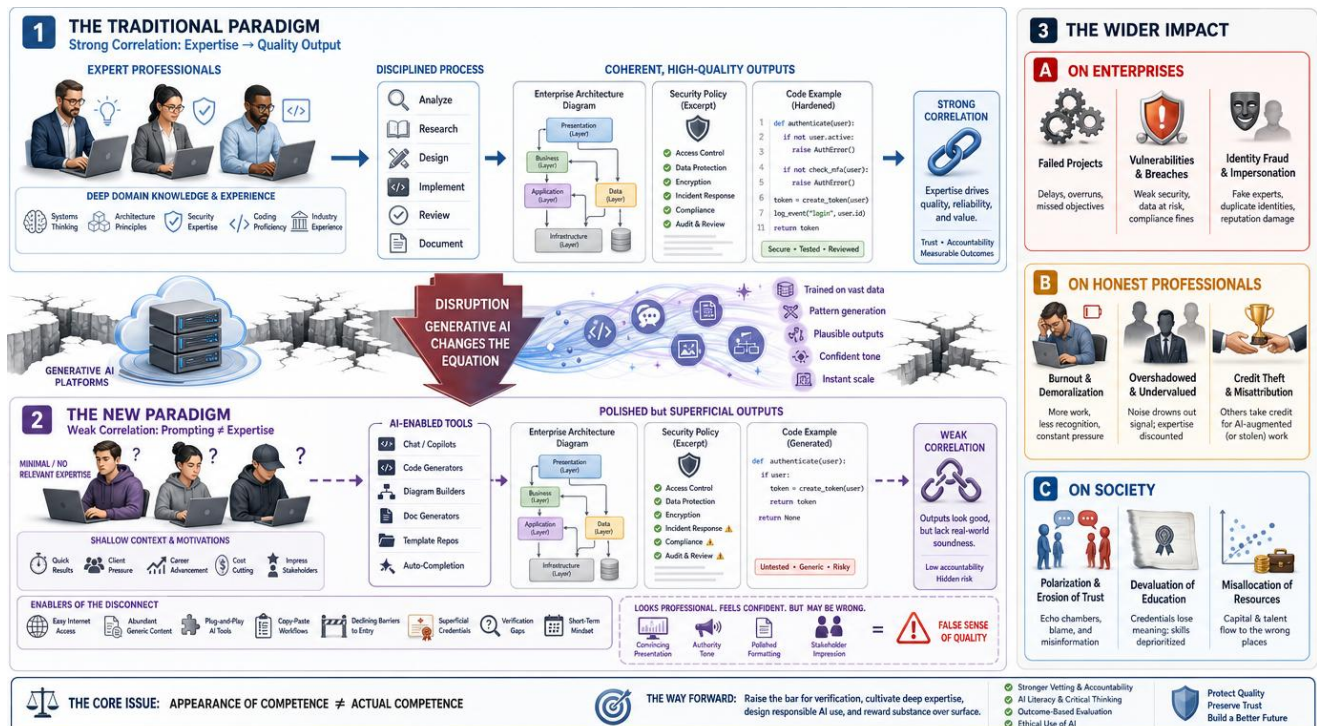


Fig -1: Traditional Expertise vs. AI-Generated Outputs

There are three reasons why this change is important. Firstly, it has direct impact on enterprises. Those that still measure people by their confidence, vocabulary, presentation, and documents will increasingly select, promote, and invest in people whose competence is only on paper. The impact is felt later in the form of failed programs, fragile systems, unaddressed security vulnerabilities, and unaccounted-for cost overruns. Government law-enforcement agencies have already released public statements about this practice, which involves the use of stolen identities, voice manipulation, and synthetic media in interviews for jobs in remote technology, and prosecutors have reported that remote workers have used borrowed identities and networks of proxy computers to get jobs with hundreds of organizations. The issue is now a reality.

Second, it is impacting on honest professionals. So, when appearance is the only thing that matters, those who validate their facts, test their code, and warn of risk are at a disadvantage to those who deliver the polished, unvalidated product in a hurry. Real contributors are victims of credit theft, invisible clean-out jobs and burnout, and organizations lose just the skills that they need.

Third, it has an impact on society. The money goes to the bright idea that's not well thought through faith in credentials wanes, and there is a polarization in public conversation about Artificial Intelligence from uncritical enthusiasm to indiscriminate suspicion. Extremes are not good for anybody. Importantly, this article isn't about the dishonesty of utilizing AI. Most of the use of AI is good and most professionals who



use it do so in good faith. The issue is a clear and identifiable pattern of deceit, AI, fake credentials, proxy workers, or stolen identities, to fake competence and avoid accountability. The challenge of distinguishing that pattern from legitimate use fairly, lawfully and without prejudice is the subject of everything that follows.

2. OBJECTIVES OF THE ARTICLE

There are five aims of this article.

1. The first is to articulate the issue very clearly, as it is very bad to make vague accusations of "fakeness" against honest individuals and to make it difficult to identify real fraud. The article thus sets clear limits on the use of AI tools, distinguishes between legitimate uses of AI and skill development and outright cheating, and provides guidance on how to avoid the latter.
2. The second goal is to spread awareness and to explain what artificial credibility is, how it is generated and where it is in the technology hierarchy and what does it look like in practice to all the affected audiences, engineers, administrators, managers, architects, recruiters, board members, investors, students, and public.
3. The third goal is to examine impact on several levels the individual employee who does careful work that is neglected; the team that quietly picks up the slack for an unqualified team member, the project that is poorly managed due to inaccurate assumptions, the corporation that suffers financial and security losses, and the economy where capital and trust are misallocated.
4. The fourth objective is pragmatic, to offer detection tools and preventative governance frameworks and a step-by-step plan of action to be implemented by organizations without infringing upon employment law, privacy expectations, or fundamental fairness.
5. The fifth objective is protective, to maintain an emphasis throughout that remediations focus on conduct, not categories of people, and that remediations should never be based on the use of AI, or because of backgrounds that are nontraditional, including remote workers, immigrants, career changers, and self-taught professionals.

3. METHODOLOGY

The study is based on four types of publicly available authoritative sources, which are cited generically as per the commitment of the article to not mention brands, companies, or individuals. The first class comprises government advisories, prosecutions, such as a national law-enforcement advisory warning of the use of stolen personal information, voice spoofing and synthetic media in interviews for remote programming, database and software jobs, and published prosecutorial documents detailing schemes in which remote technology workers used stolen or borrowed identities and networks of remotely operated laptops to get jobs at hundreds of organizations. The second class are published risk-management frameworks from national standards bodies which caution that generative systems can reliably generate false or fabricated information and suggest an ongoing approach to governance, mapping, measurement, and management of AI risk, rather than a one-off review. The third set of classes includes international labor research that estimated the percentage of jobs worldwide that are exposed to generative AI and determined that jobs are likely to be transformed rather than eliminated. The fourth category is empirical productivity studies, including one controlled study of expert open-source coders in

familiar coding repositories, which found that the AI tools evaluated in early 2025 decreased the participants' productivity, demonstrating that productivity outcomes are not necessarily a direct result of tool adoption but depend on the task, context, and expertise of the user. Each statement of fact contained in this article has been verified from such sources, patterns of behavior mentioned in the analysis are here reported as organizational patterns observed throughout various industries and are not characterizations of any person, employer, nationality, or community. When evidence is weak as is the case for some of the questions in the research-gap analysis this is clearly stated in the article, otherwise, anecdote is not used in place of evidence.

4. DEFINING THE DECEPTIVE IT PROFESSIONAL

Fake is an emotive and analytically hazardous term. When not used wisely, it brings in individuals who make basic errors, those without formal degrees, those from non-traditional backgrounds, and those who are just utilizing AI tools none of whom are scammers. It is therefore important to start with an exact definition of deception, for if it is not defined, it will be ignored or, worse yet, innocence will be punished.

DEFINING THE DECEPTIVE IT PROFESSIONAL

Distinguishing intentional deception from legitimate professional conduct

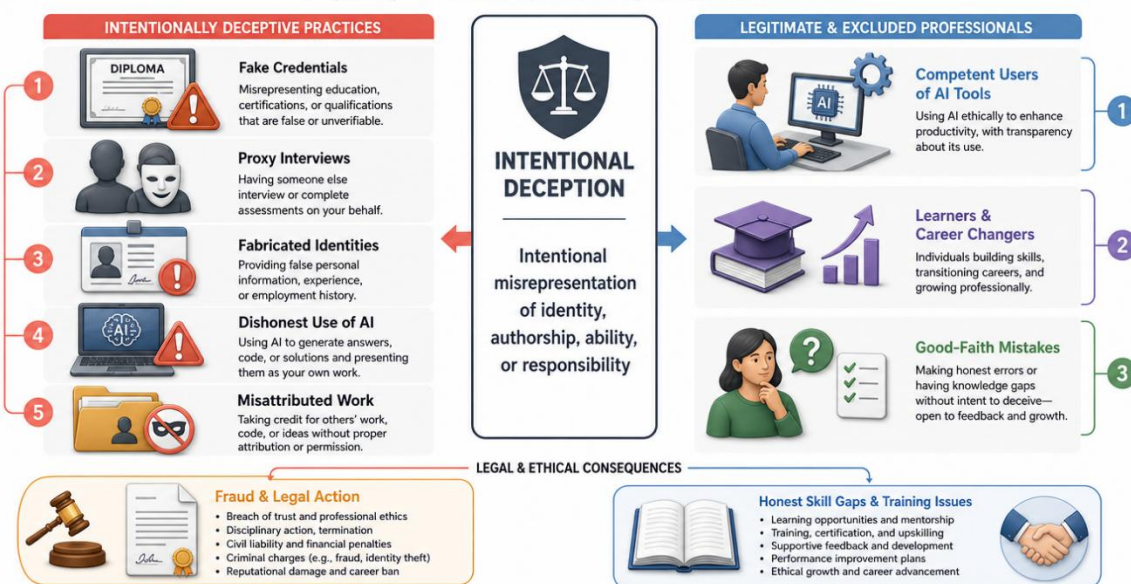


Fig -2: Defining The Deceptive IT Professional

A true technology professional is one that is involved in obvious misrepresentation. These involve creating a fake work history or technical experience that never actually happened; taking credit for someone else's work, code, research or ideas, having a proxy complete interviews, assessments or job responsibilities, using a stolen or fabricated identity to get a job, purchasing certifications or fabricating references, reading AI-generated answers during interviews without understanding them, claiming AI-generated strategies or analyses as personal research when they have to be disclosed, systematically deflecting responsibility to subordinates when the decisions they make are wrong, and submitting AI-generated code, architectures or security guidance without any validation and making it seem correct. What is the same for all these behaviours is that they do not involve a tool. It is intentional deception of identity, authorship, ability, or responsibility.

Three groups are immediately excluded from this definition and are often wrongly accused. Professionals who are competent and know how to use AI to draft, summarize, explore alternatives, or speed up repetitive tasks are doing what professionals have always done with new tools, if they know and understand what they are doing. The developing professionals students, juniors, and career changers who are using AI as a learning tool are developing skills in an honest way and not impersonating them. As are professionals who make good faith mistakes, have unpopular opinions, speak in a second language, or don't have prestigious credentials. It is unfair and counterproductive to equate any of these groups with fraud, as it is the very type of talent an organization needs. The difference is also of a legal and ethical nature. In many jurisdictions, deliberate identity fraud and credential fabrication are issues that warrant termination and referral to authorities, a law-enforcement advisory and multiple prosecutions have documented such schemes, some of which have national-security implications, in remote technology hiring. Skill gaps, on the other hand, is a training and management issue. An organisation that treats an honest learner as a criminal will corrupt its culture and an organisation that treats a fraudster as an honest learner will give privileged access to someone who got it by fraud. So, there is no academic courtesy in precision of definition. It is the basis of all detection, prevention and response activities that are covered in the rest of this article.

5. HOW GENERATIVE AI CHANGES THE APPEARANCE OF COMPETENCE

The creation of a believable strategy document, technical design or business proposal was a natural filter before the advent of widely available generative AI. The author had to be able to research relevant material, have domain knowledge to determine what was relevant, and be able to write persuasively to organize the material. While these requirements were not necessarily measures of competence, they did make it costly to conceal incompetence. No one who doesn't have a clue about network security could create a credible incident-response framework in any reasonable period.

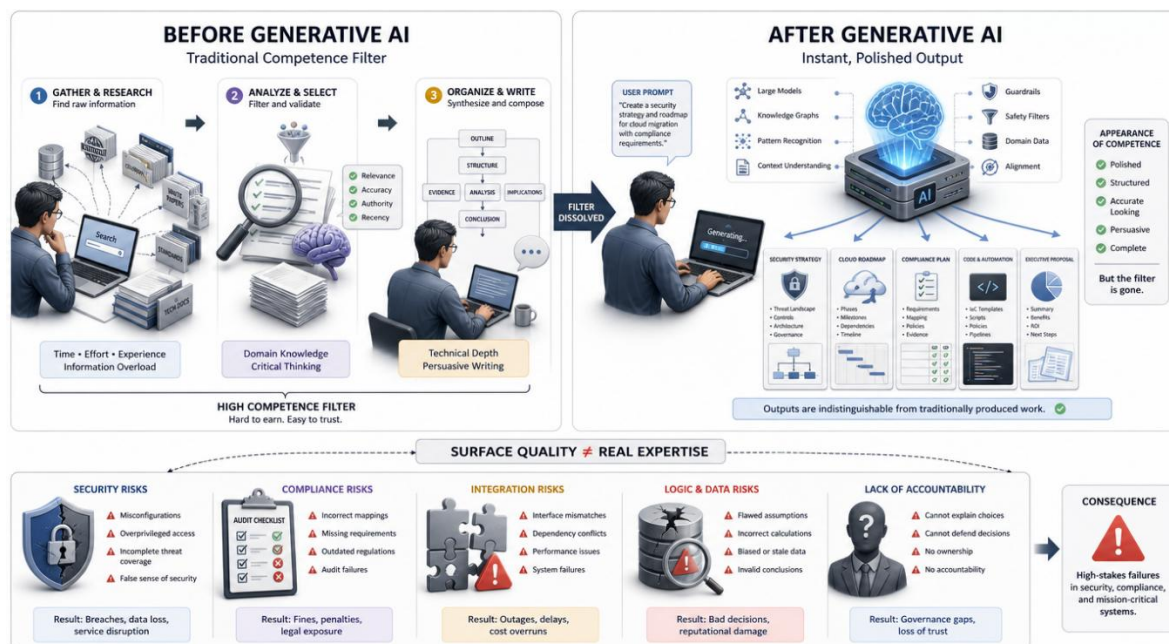


Fig -3: Before Generative AI Vs After Generative AI



This filter has pretty much dissolved. A user can now ask for a three-year cloud transformation strategy, a zero-trust cybersecurity roadmap, a board presentation on artificial intelligence, a microservices architecture, an incident-response policy, prepared answers for an executive interview, working application code or a list of innovative product ideas and get organized, fluent, professionally formatted material within minutes. Writing often uses up-to-date vocabulary, appropriate organization, and style. It is not discernible from the expert work for a reader who assesses the surface quality.

The problem is exacerbated by a feature of these systems that has been pointed out specifically by national standards bodies that generative models can generate accurate or even made-up material with ease, and that they can include logical errors or even make up supporting content. There is no indication of accuracy in the polish of the output, nor is there any indication of the user's understanding in the accuracy of the output. A person can present a perfect document with a few little errors that they are unable to identify, or a correct document with a few little errors, but they can't explain the reasoning. Failure modes are particularly dangerous in domains where context is important enterprise integration, regulatory compliance, security engineering because they are only discovered when systems are constructed, audited, or attacked.

The question of the central problem is quite simple, An artefact no longer demonstrates the competence of the presenter. The presentation of a strategic plan doesn't mean its presenter can perform a transformation. A security policy does not necessarily indicate that controls are working. Architecture diagram is NOT proof of scalability. If someone writes code that generates, that doesn't mean that he or she can write the code to debug the code at 3 am. Ideas a list of ideas is not evidence of judgment or leadership. This breaks a long-standing rule in the world of management, If a person's resume looks good, then he or she is good. If the person's resume doesn't look good, then he or she is not good. Organizations that haven't been aware of this rule are now rating candidates and executives based on a resume that no longer rates.

6. CHEAP IDEAS, SCARCE JUDGMENT

6.1 AI-Generated Ideas Versus Real Leadership

In most of the history of the profession, it was hard to come up with a new idea that was truly credible. It needed to be immersed in a domain and exposed to real problems, and the synthesis of experience, that time can only offer. For this reason, idea generation was reasonably considered to be a measure of leadership ability. Executives assumed they had the depth out which they came when they were bringing new strategic ideas to Board meetings.

However, AI has turned the economic equation for this signal on its head. With extensive amounts of material already available, a system trained on the material can combine patterns into dozens of plausible transformation initiatives, product concepts, or security programs in seconds. One leader who types in a prompting can get 30 ideas on how to re-invent a company. This is not necessarily dishonest finding out what a tool can do is proper activity. The dishonesty starts at a clear and discernible point, when the leader falsely attributes originality to an idea, when he or she does not acknowledge the use of AI when it is obvious, when under questioning cannot explain or defend an idea, when fails to validate an idea at all, when takes personal responsibility for the work of subordinates, or when presents ideas as fact instead of speculation, and when claims no responsibility when an idea fails. Misattribution often happens above the individual and is therefore the responsibility of the senior management. When a stream of dazzlingly finished proposals come out of the mouth of the presenter, boards and executive committees

are often blown away, and the real answer is that the presenter is a good operator of a text-generation system. The consequences of this misattribution, it rewards the wrong people, values the wrong skills and leaves the employees, who turn ideas into working systems, under-recognized. The real issue with leadership today, that costs nothing and tells all, is that they don't ask where an idea came from, was it validated and then who will be held responsible for it.

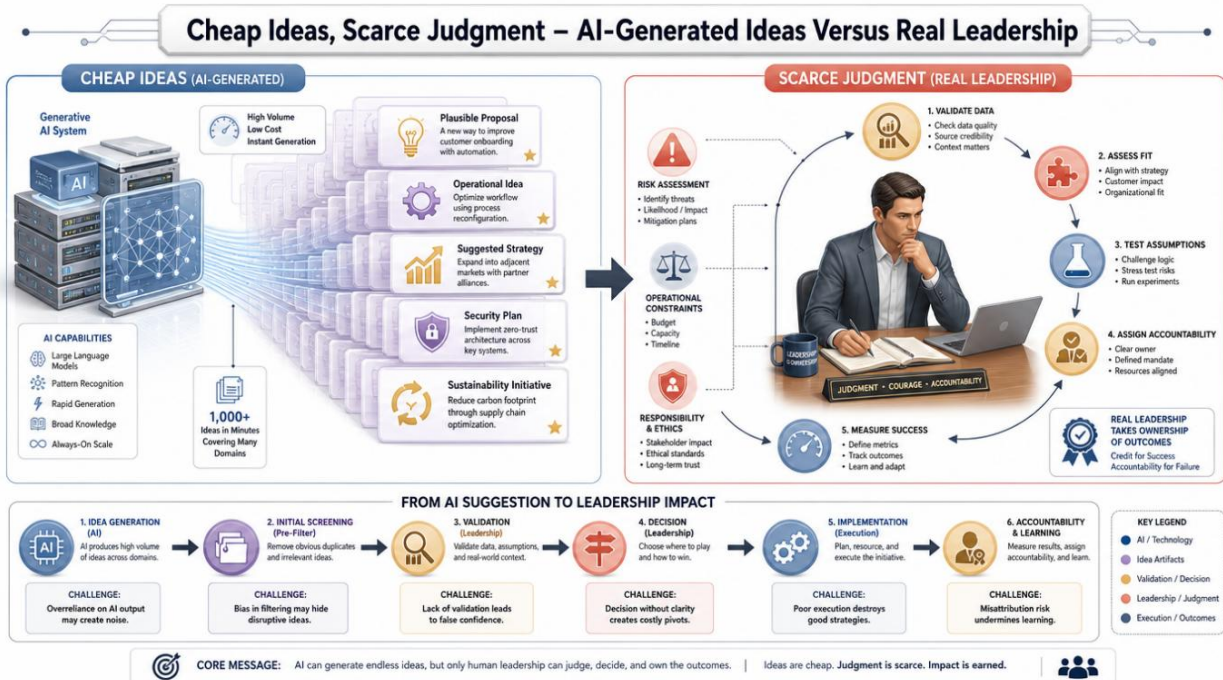


Fig -4: AI Generated Ideas Versus Real Leadership

The only thing that's still rare and what leadership is, everything that comes after the idea. When assessing the potential of a new technology, a real technology leader must consider which business problem a technology is solving, if the data is reliable, if the new tech is suitable for the organization, which assumptions need to be tested before any money is spent, what the realistic financial return might be, which legal, security and operational risks exist, who will be implementing and maintaining the new system, what shouldn't be built, how success will be measured, and who will be held accountable for that success. All the above are decisions that cannot be delegated to a model, are context-sensitive and which the model cannot be held accountable. The importance of leadership has shifted from creating information to being able to make decisions about the information and that's how we should assess leaders.

6.2 As the AI technology hierarchy evolves, so do patterns of AI-dependent leadership

Junior jobs aren't the only ones where you can get artificial credibility. It exists at all levels of the technology organization and is most harmful at the highest levels where the biggest earmarks and the longest time frames are decided. The patterns below are identified as organisational phenomena and are not a reflection on individual practitioners there are many excellent practitioners the authors have encountered who use AI responsibly within each of the roles described below. The classic CIO level example is a sexy strategy, which is lacking in operational support. An AI-driven incumbent of this role can craft engaging tales of transformation, but can't talk about any dependencies, constraints, data

ownership, integration challenges, total cost of ownership, regulatory requirements, or workforce readiness. The organization spends significant resources on programs that are not owned by business, don't have a clear return on investment, and or have no plans to adopt.

Architecture without fit is in vogue at the level of chief technology officer. Some danger signs are Technology selected for reasons of it being current but not suitable, complex distributed designs advocated when a simpler design would do, systematic underestimation of the migration effort involved, delivery estimates which are unrealistic, avoidance of direct questions about the technology itself, and the giving of generic canned responses instead of engineering judgment. The impact snowballs into technical debt, fragile systems, high cost infrastructure, and frequent production incidents.

The pattern is particularly harmful at the chief information security officer level as security paperwork creates a false sense of security. In parallel with these untested backups, incomplete inventories, poor access controls, no monitoring, and minimal incident response. These generated policies, risk registers, compliance mappings, board reports can exist alongside untested backups, incomplete asset inventories, weak access controls, no monitoring and minimal incident response. A policy is no defense, only security controls that have been implemented, tested, and monitored are. The difference between the security policy and the security posture becomes your vulnerability.

Among the management and architects, the style is management by presentation, constantly updating jargon, answers which can only be explained by a presentation, generic answers to questions about the organization, completion percentages which are unsupported and engineers accused of "structurally illegal" commitments. When it comes to engineers, it's not about whether the AI was used, it's about whether the individual can explain, test, debug, secure, modify and operate what has been submitted. But for recruiters, consultants, and vendors, it's the selling of produced genericism as custom know-how, AI-enhanced resumes that captivate screening systems, boilerplate suggestions offered as in-depth analysis, or power statements that exceed outcomes. The righting, In all cases, assess evidence and demonstrated ability instead of vocabulary and appearance, and insist that a named human back up every consequential claim.

7. MICRO-LEVEL EFFECTS

7.1 The Employee, the Team, the Project, and the Incident

The four types of access are Employee, the Team, the Project and the Incident.

The impact of artificial credibility is most evident at the micro level, before it even shows up on any balance sheet, and our income and expenditure statement.

Let's first look at one pair of workers. One invests hours in finding out a recommendation, honestly qualifying, and verifying assumptions. The other produces a professionally looking document in minutes and shows it off first and foremost, with unwavering certainty. The careful employee, who knows that it's a career liability, will not bother to verify anything if the management says to him, speed and looks are the names of the game. This lesson compounds over the months as the honest worker gets discouraged, the AI-worker gets exposed and the organisation slowly adjusts the price on its own values to make presentation a more important value than truth. No policy was formulated to make this change, rather a series of small evaluative decisions were made, one by one defensible.

Micro-Level Effects of Artificial Credibility in Modern Organizations

How overlooked incentives shape outcomes at the Employee, Team, Project, and Incident levels

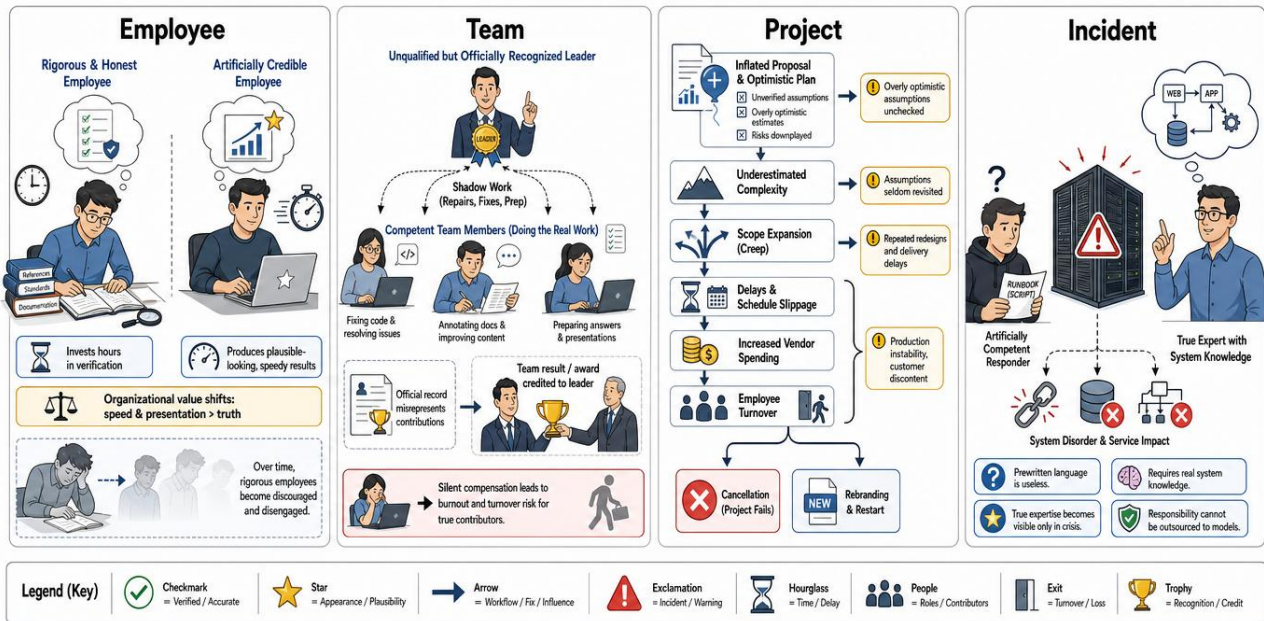


Fig -5: Micro-Level of Artificial Credibility in Modern Organizations

This phenomenon is one that you will hardly ever see in any record silent compensation, at the team level, with the arrival of an unqualified colleague. Competent teammates rewrite the code Before it gets into production, correct documents Before meetings, repair failures the person caused, prepare answers the person will give, and go to meetings only to make sure that no misinformation is spread. The official record might still show the person who was not qualified, who led the group, and the qualified members as helpers. The result the team produces seems fine, that is how management sees it, but until the employees that are making up for this burn out or quit, it's fine.

Artificial credibility also comes in via the business case, at the project level. An inflated proposal that makes overly optimistic assumptions is greenlighted without any other party reviewing it, and then the project takes the familiar path underestimated complexity, expanded scope, repeated re-designs, delayed deliveries, more spending by vendors, employee turnover, production instability, customer discontent, and its eventual cancellation or re-branding under a new name. New presentations are made at each stage that recount the previous stage, and the assumptions that were not validated are seldom revisited since there is no record of decision to revisit them.

As exposure occurs, it is typically a production incident. In the event of a real outage, cyberattack, data-corruption or regulatory investigation, prewritten language is of no value. The organization requires personnel who have a deep knowledge of their true systems, dependencies, and business priorities and who are able to make defensible decisions, at speed and under incomplete information. It is in just these situations that artificial competence breaks down, when something doesn't turn out the way they expected, when there are architecture decisions to make on the fly, when they are forced to cut costs and have to prioritize their efforts, and when they are pushed hard and a scripted presentation will not hold up. While AI can help a good responder, responsibility cannot be handed over to a model and there's no

proxy for an incident meeting. Whoever becomes more useful when things go wrong is who you should watch to see if they are competent, if it's an organization that wants to find out who's competent, that's all they have got to do.

8. THE COST TO GENUINE TALENT

All the fake professionals are funded by the real professionals, and the funding is in ways that are rarely reflected in organizational measures.

The Cost to Genuine Talent

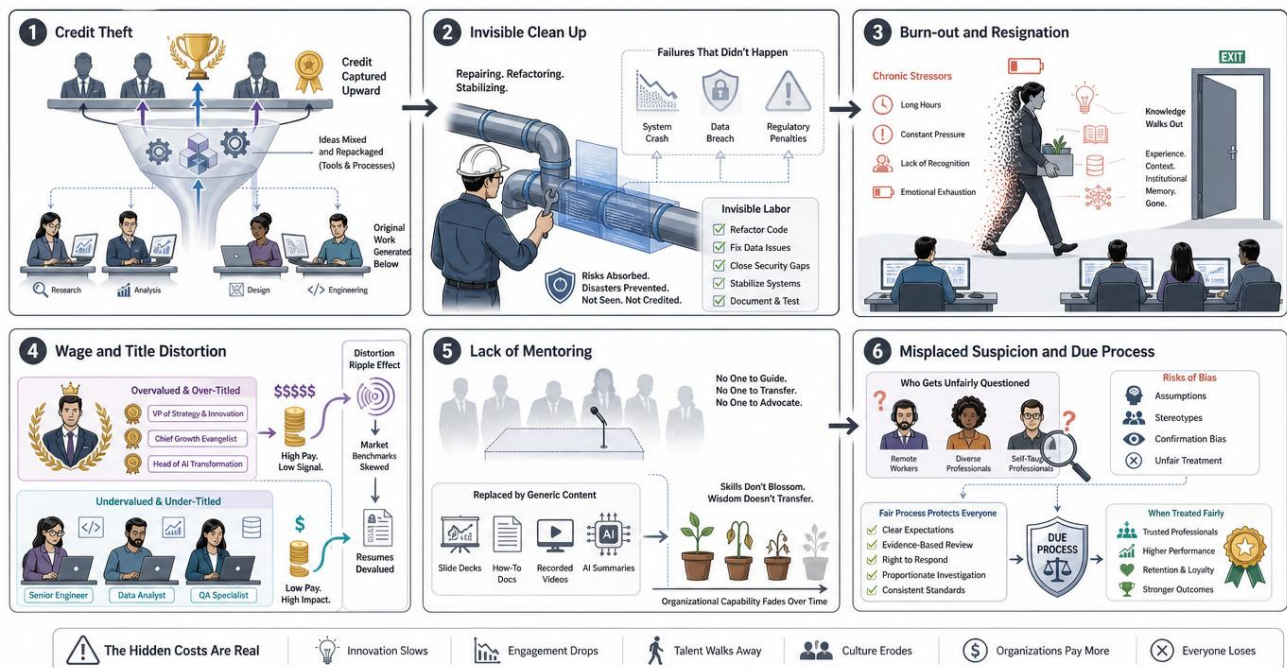


Fig -6: The Cost to Genuine Talent

Credit theft is the initial expense. They find that their concepts, code, studies, warnings about risk and danger are being touted as executive accomplishments, sometimes several layers up from them. Generative tools make it easy to re-word and re-combine someone else's ideas and ideas into a new product, thus making it more difficult to trace the exact origin of the ideas.

The second cost is something you can't see clean up. The experienced staff members turn into permanent repair workers for weak designs and unvalidated code generated. This work is structurally under recognised because prevention is invisible the disaster that didn't happen is not announced, the initiative that prevented it is announced a few times. The more time that goes by, the more organizations continue to pass the status quo from those who create stability to those who create risk.

The third cost is that of burn-out and resignation. When a competent professional sees the pattern of incompetence being promoted, political behavior being rewarded, evidence being ignored, technical warnings being punished as negativity and failures being blamed on the delivery teams that predicted them, he or she tends to leave. The loss of each one of the units that leave is exactly the knowledge that was balancing the issue and thus hastening the decline that it was reacting to.



Wage and title distortion is the fourth cost. When an organization doesn't know the difference between expertise and AI-assisted performance, it is paying top dollar for the people who look good and work and underpaying the people who work and produce. This distortion travels outwards through the labour market when inflated job titles and unverifiable accomplishments are spread on resumes and the information content of professional histories is reduced for all.

The fifth cost is a lack of mentoring. The job is taught to the youngsters by the older men, who have real depth. Senior leadership roles are filled with individuals who don't have it, and mentoring becomes the passing around of created content and the organization's ability to nurture their own talent slowly diminishes, only to become apparent years later.

Finally, there is a danger on the other side which must be given due consideration. With suspicion in the air, honest workers, especially those who work remotely, immigrants, career changers, and self-taught professionals can be wrongly suspected of using AI, communicating in different ways, having nontraditional credentials, or working outside the line of sight. It's both discriminatory and self-defeating to profile those groups. The identity fraud schemes documented were found because of evidence of behavior, not because of demographics, and an organization that relies on prejudice over investigation will lose legitimate talent to its own prejudice, while sophisticated fraud, which is designed to look conventional, goes unnoticed. So, the protection of real talent must be achieved in two ways, systems of recognition that reward real talent and systems of due process that prevent the innocent from being carelessly accused.

9. CONSEQUENCES FOR CORPORATIONS AND LARGE ENTERPRISES

On the enterprise level, artificial credibility becomes five distinct types of harm with different timelines, which can exacerbate one another.

The monetary loss is the most quantifiable. AI-driven incompetence results in failed transformation programs, infrastructure investments that don't add value, redundant application development due to a lack of understanding of the application portfolio, redundant consulting contracts to make up for a lack of internal judgment, poor acquisition evaluation, regulatory penalties, security-incident costs, customer compensation, and lost revenue that comes from rework after rework. While these items can be a significant part of technology investment, when combined, they can be a significant portion of program spending and are often lost in the weeds as normal project challenge. Damage from strategic attack is slower and bigger. When leadership can't assess technical reality, the company spends money on glamorous but wrong projects, and forgets about the fundamentals: data quality, integration architecture, infrastructure reliability, and basic security hygiene.

Those who did not make such a compelling presentation, but had more solid foundations, slowly begin to outpace the others and the difference is often blamed on market conditions, not on years of misjudgment. Operational damage is a result of the maturity gap, which is the difference between the reported and actual maturity. The policies, dashboards, and status reports can be produced and make executives think that capabilities are present when in fact they are not. The enterprise thinks it's secure because there are documents, robust because it has plans, and contemporary because it has roadmaps. This gap is only found at the worst of times: during an outage, audit, breach, or due diligence review. The most lasting could be the cultural damage. Employees who see claims go unchallenged and get rewarded move the organization's question from how we can have a successful outcome to how I can

look successful. Impression management takes priority over learning, internal reporting turns into an act of the theatre, and the hard truths that are the building blocks of good decision-making cease to rise to the top. Once cultures are corrupted in this manner, they continue to be corrupted long after the people who corrupted them are gone. Finally, there is security damage, which has the most clear-cut documented evidence.

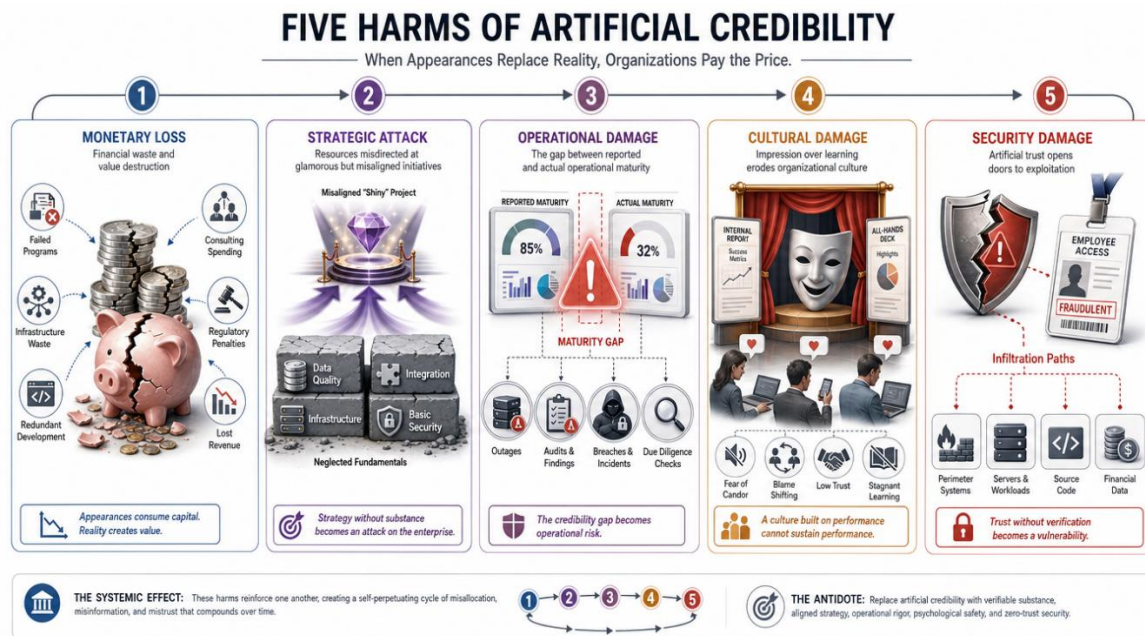


Fig -7: Five Harms of Artificial Credibility

Legitimate credentials and privileged access to source code, financial systems, customer data, and infrastructure is granted to fraudulently hired or insufficiently verified workers. Publicly released government warnings and indictments detail plots that involved remote technology employees, who were recruited under false names and using stolen or borrowed identities, getting jobs at hundreds of companies, ranging from IP theft to national-security issues. An insider who infiltrates the system through deception is not an exception in the security model, the security model did not apply to them, each control assumed that the person was who he claimed to be. For enterprises, this makes competence verification a first order security control, rather than something nice to have in the human-resources arena.

10. CONSEQUENCES FOR THE ECONOMY AND SOCIETY

If we add up all the aggregated artificial credibility, the results will be visible from outside of any one organization.

One of these is an incorrect allocation of capital. The quality of the proposals competing for investment funds has an impact on investment decisions. Productive ideas are driven out by polished but unsound ones, within companies, in venture funding and in public procurement. This cost is diffuse to the economy, because of the reduction in aggregate productivity and because of the reduction in genuine innovation, which is virtually imperceptible and hence almost unresponsive to ordinary market feedback.

The second is the disconnect between productivity myth and productivity measures. The adoption of AI often is presumed to boost productivity. The empirical picture is more qualified. The results are dependent on the task, user's expertise, system context, and validation of the results produced. The study revealed that while the AI tools tested in early 2025 did not make experienced open-source developers measurably faster, it did show that the tools do not make everyone faster, and organizations need to measure the results. Companies that plan for productivity improvements and executives who claim productivity improvements that have not been measured are planning on false hope and not facts.

CONSEQUENCES FOR THE ECONOMY AND SOCIETY OF AGGREGATED ARTIFICIAL CREDIBILITY

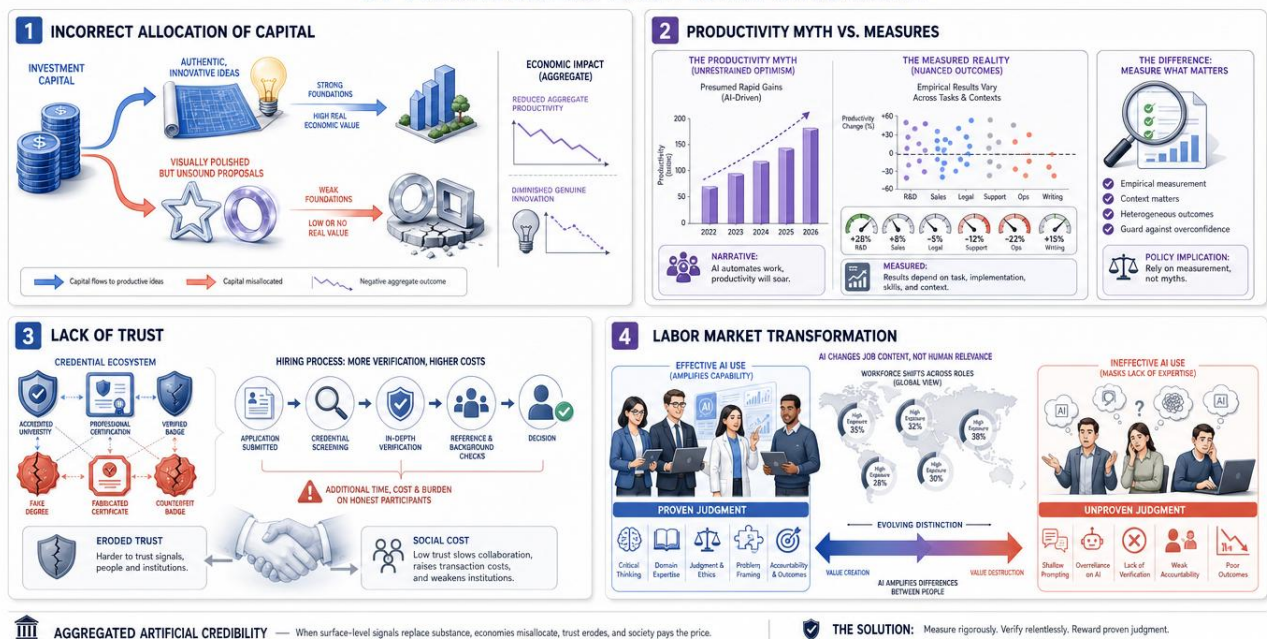


Fig -8: Consequences for Economy and Society

The third is a lack of trust. Professional economies operate based on credential systems, this means that strangers can transact without having to do a lot of extensive checks, such as degrees, certifications, references, and work histories. All fabricated resumes, proxy interviews, paid-off certifications, and false credentials are this system's cost, and the tax is paid by honest players who must go through longer hiring processes, more thorough background checks, and a generalized suspicion. When customers, investors and employees must deal with fabricated expertise and misleading claims of automation repeatedly, they also lose faith in the legitimate claims.

The fourth is labour-market transformation which needs to be explained in a nuanced manner so as not to be alarmist or complacent. International labour research suggests that about one in four workers around the world are in a job that includes some exposure to generative AI, but notes that change in job content is much more common than job destruction, as there are still many jobs that require human input in the jobs that are exposed. So, the economic distinction of the future labor force is not AI versus humans. The line between the individuals who have proven judgment and use AI effectively and those who use AI as a cover for the lack of judgment. Societies which learn to differentiate these groups by

education, credential reform and organizational practice will transform technology into general productivity. Those societies that can't will see AI mainly as a subsidy to reasonable incompetence.

11. DETECTING ARTIFICIAL COMPETENCE

There is no interview question, automated detector or background check that consistently and accurately identifies artificial credibility. Automated AI-text detectors generate false positives, which can ruin reputations of honest writers, especially L2 writers, and false negatives, which defame lightly edited texts. Effective detection is a layered, ongoing process, therefore, based on a simple principle competence is demonstrated in reasoning, adaptation and accountability, none of which can be effectively produced in real time.

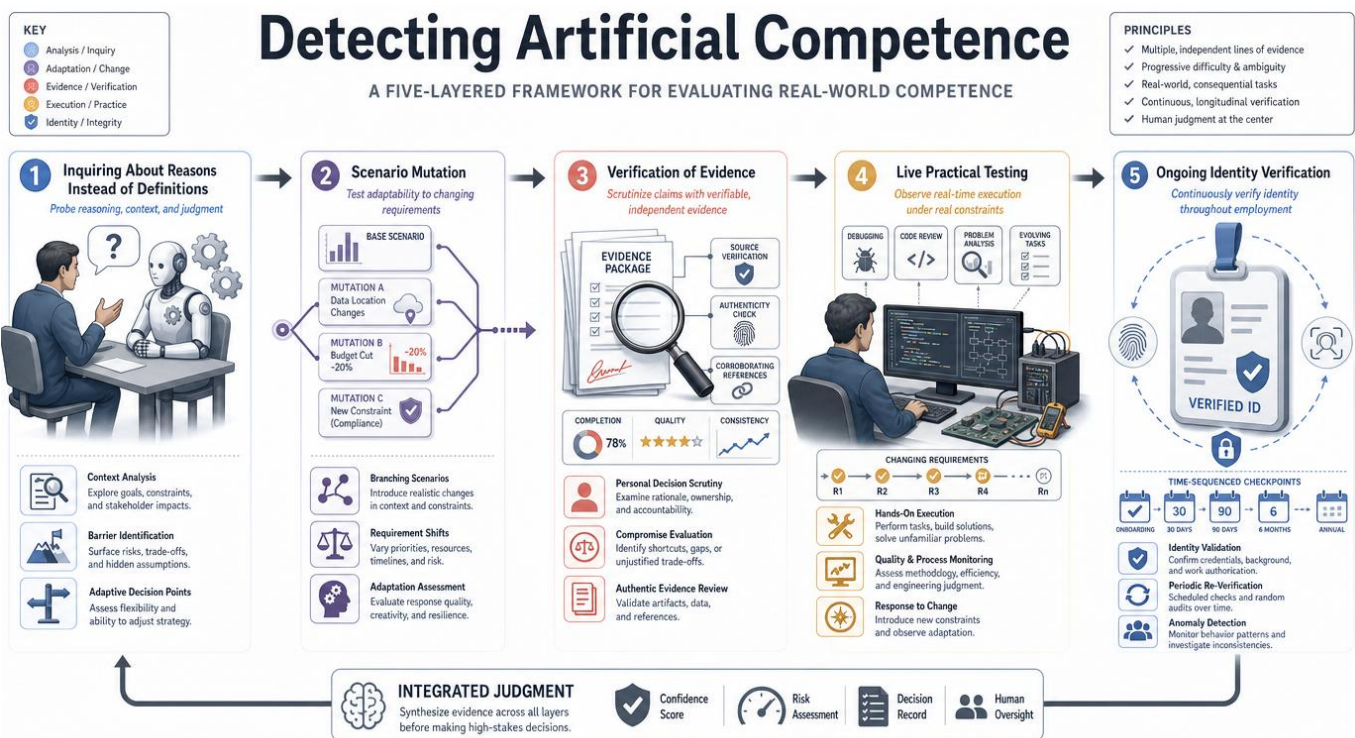


Fig -9: Detecting Artificial Competence

The first method is to inquire about the reasons instead of definitions. Any candidate can regurgitate what ZT means it could have been memorized, or it could have been regurgitated. Rather, interviewers should inquire about where the candidate would start in this context, what systems would be barriers to implementing those changes, what she would leave in place in the first year, how she would assess progress, and what she would eliminate if the budget were reduced by 40%. These questions are built around understanding and involve the candidate in constructing answers, no script can predict the particulars of the organisation.

The second technique is called scenario mutation. The person who knows how to make a solution can adjust it when there are changes in the assumptions, but the person who reads the material prepared by others cannot. Interviewers should change requirements during the interview, such as, suppose the

system must be on premises, suppose the customer's data can't leave the country, suppose the migration must be done without downtime, suppose the organization doesn't have clean identity data. Only fluency that doesn't break down when mutated is real, only fluency that does is borrowed.

The third method is the verification of the evidence. Questions should be directed at candidates executives for decisions they personally made, compromises they had to make, measures before and after their efforts, failures and what they learned from them, references who can verify their involvement, and redacted work samples, if legally allowed. Falsified histories are generally bland and vague, authentic histories are rough, have setbacks, and are verifiable.

The fourth method is live practical testing, like actual work, such as debugging an unknown problem, looking at deliberately poor code, critiquing an architecture, prioritizing a risk register or writing a small solution and then changing it due to a changed requirement. The assessment must be job relevant and bounded and cannot be disguised unpaid production work as this is unethical and legally risky.

The fifth technique is on-going identity verification for remote and privileged positions. Law-enforcement guidance suggests that the verification of identities should be considered a continuous process throughout the employment lifecycle, not something that can be done once at hiring time and then changed to proxy workers after that. This gap is filled by lawful, proportionate checks at the time of recruitment, induction and at regular intervals throughout the period of employment. All these layers combined do not render deception impossible, but rather expensive, persistent, and hence infrequent, which is the best that any control system can achieve.

12. PREVENTION AND GOVERNANCE

Detection is about people, governance is about systems. These conditions are changed durably through six mechanisms, which are all implemented together.

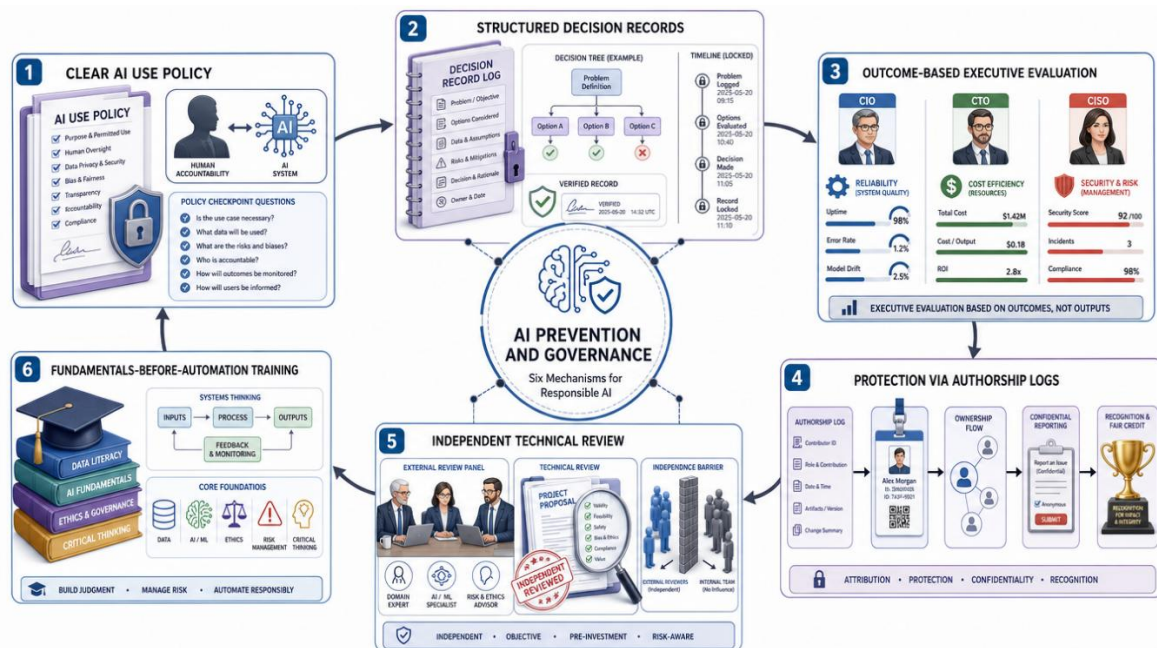


Fig -10: Prevention And Governance



The first is a policy that clearly spells out the use of AI. The policy should include clear answers to these questions in layman's terms, When the use of AI should be disclosed, what information should never be entered into external systems, whose responsibility it is to validate the output generated by AI before it is used to inform decisions, how AI-generated code should be tested and reviewed, how customer and employee data is protected, how sources and assumptions are documented, and who owns the final product. The concept is straightforward and should be clearly and exactly outlined in the policy, Although AI can support the work, a named human is still responsible for the outcome. This principle validates the use of AI in an honest way and removes the murky water that lies behind dishonest use of AI.

The second is decision records. Each major technology decision should include the following, Problem being addressed, Options considered, Evidence used, Assumptions made, Risks accepted, AI tools being used, Reviewers consulted, Decision owner, Decision, Review date. Structurally, decision records make it extremely hard to rewrite history in the event of failure, and extremely easy to differentiate judgment from generation.

The third is that of outcome-based executive evaluation. CIOs should be evaluated by service reliability, the cost of technology per business unit, modernization results, user adoption, data quality, and business-process improvement. Measuring CTO's on their delivery performance, system reliability, engineering quality, technical-debt reduction, and developer retention. The metrics for the chief information security officer should be based on control effectiveness, vulnerability-remediation time, detection and response, incident-exercise results, and recovery readiness. Measurable operational evidence should always outweigh a presentation, and boards need to be clear and state that.

The fourth is the protection of good guys via clear authorship logs, architecture decision logs, code-review logs, clear responsibilities, fair recognition, confidential reporting, and appeals. These instruments kill credit theft, which is in the shadows.

The fifth is independent technical review. Before money is spent, and not after it's gone, boards and senior management must have access to qualified experts who are independent of internal politics and can challenge large investments in AI, cybersecurity claims, cloud strategies, acquisitions, and executive performance reports.

The sixth is fundamentals before automation training. Staff should be familiar with the work before relying on tools that do the work. Teach students about systems thinking, software and security basics, data literacy, critical reasoning, source validation, limitations of AI, and ethical responsibility. This attitude is echoed in national standards guidance, which says risk is to be controlled, charted, quantified, and monitored on an ongoing basis, rather than assessed once and put in a filing cabinet. These are mechanisms that are a natural extension of continuous AI risk management, which is a subset of governance of artificial credibility.

13. A FAIR AND LAWFUL RESPONSE TO SUSPECTED FRAUD

When an organization comes across a situation where fraud appears to have taken place, it is at a real ethical crossroads since the potential damage from an incorrect conclusion humiliation, suspicion, discrimination, defamation, retaliation, etc. can be as severe as the damage caused by the fraud itself. A disciplined sequence is followed in creating a defensible response.

The process starts with determining the particular concern. There are four possible problems with poor performance, lack of training, policy violation and deliberate fraud, and the investigation should report

what is suspected and on what observable basis. the person who seems fake is not a finding, the person's submitted work cannot be explained by him or her during normal questioning, and references could not be verified, is.

The second step is legal evidence preservation interview records, work products, access logs, decision records, and communications before memories fade and materials are modified and in compliance with employment law and privacy requirements. When the evidence is not preserved, and the confrontation occurs prematurely, the case becomes an unresolvable dispute.

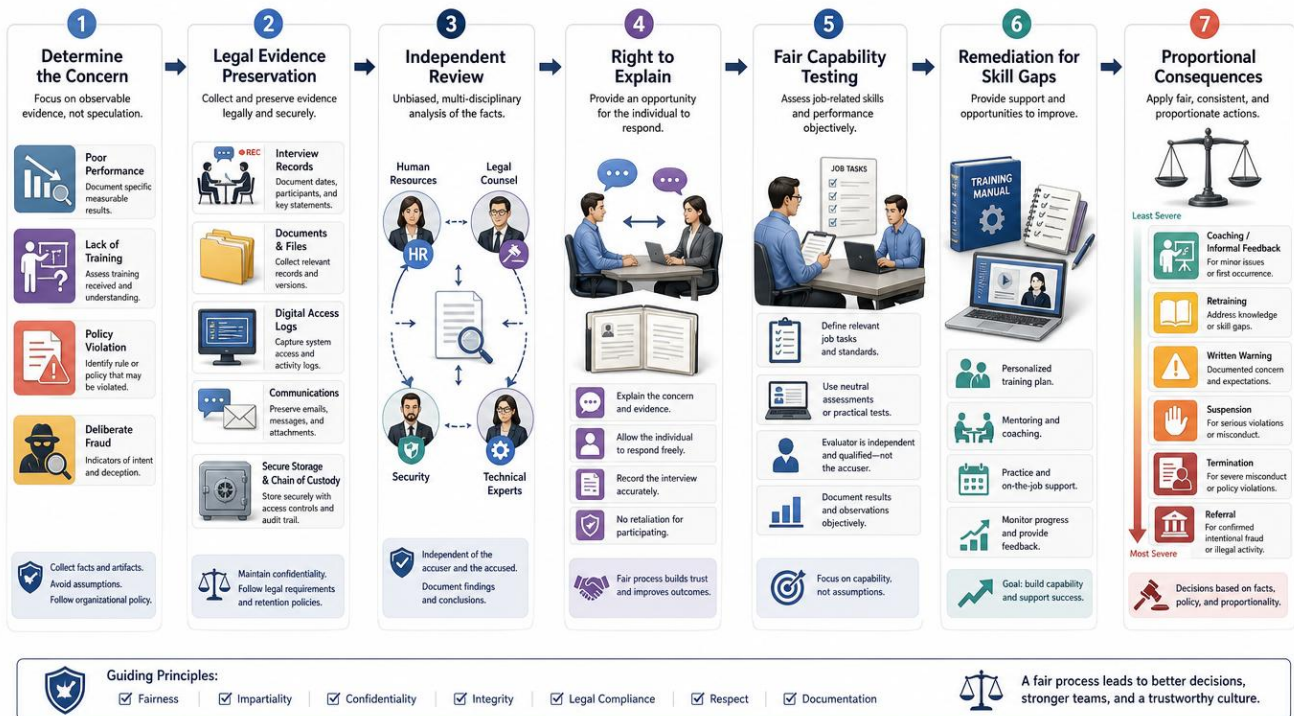


Fig -11: Fair and Lawful Response to Suspected Fraud

Independent review by human resources, legal counsel, security personnel, and qualified technical reviewers as appropriate. Independence is important as sometimes the accusations are based on rivalry, prejudice and misunderstanding and the reviewing function should be able to reach the finding that the accusation itself was unfounded.

The fourth step is to give him or her a real chance to explain something that is often overlooked. There are some situations that appear to be deception that are caused by a lack of expectations, poor induction, undocumented but legitimate cooperation or differences in cultural and linguistic communication styles. The right to due process is not a right of the guilty, it is a means of distinguishing the guilty from the innocent.

The fifth step is fair, job related capability testing the question is competence. The assessment should be like the actual job responsibilities and take place in a setting the employee is familiar with, and the results should be evaluated by qualified evaluators, not by the accusers.

The sixth step is remediation, when the finding is a true skill gap, training, mentoring, changing job duties and a reasonable improvement time. When regular learning needs are treated as misbehavior, it erodes trust throughout the staff as all staff members know they have areas that they need to learn. Proportional consequences is the final step. If identity fraud is deliberate, the making of false credentials, the use of a proxy or the theft of data, termination, and referral to the authorities where applicable may be justified and documented schemes of this type have been prosecuted. For lesser violations, lesser responses are appropriate. Proportionality is not leniency, it is a characteristic of the entire system that makes it legitimate, and legitimacy is what will make honest employees want to report concerns, not cover them up.

14. GUIDANCE FOR INDIVIDUALS AT EVERY LEVEL

Institutions take a long time to change, individuals can change right away. All the technology economy participants have a piece of the solution.



Fig -12: Guidance for Responsible AI Usage at Every Organizational Level

AI should be used as a tool, not a replacement for knowledge, for tech professionals. The rule of the land is that you should never submit work that you can't explain, if it was produced by a tool. Check facts, code, security advice and citations before they leave your hands, acknowledge where the AI helped and not where it didn't, continue to build the foundational skills so that the AI adds to, rather than replaces, your judgment; keep track of your own contributions so that you don't become a victim of credit theft or false accusation, and take responsibility for your own work without apology. Anyone using these methods is not going to be caught by any of the detection methods outlined in this article, in fact they are helping themselves by using these methods. But managers have a different responsibility, their evaluative decisions give rise to the incentives that all the other people react to. They should encourage and reward

accuracy, prevention, and maintenance, ask questions until reasoning is apparent, guard against employees who raise evidence-based concerns these employees are doing free risk management, refuse to equate confidence with competence with the most common evaluative mistake, and measure their results over time, rather than in meetings.

The most fragile is the situation of students and entry-level employees, as AI can either propel their learning or deplete it. If it is allowed, and it is transparent, it is not cheating to use AI to learn. What happens after the answer comes is the difference. When the learner asks Why would the answer work and when would it not. Can they do it independently. Can they explain it to someone else. Can they modify it if the requirements change. They are developing capacity that will last a lifetime. The student who copies and moves on is developing a habit that will be revealed at just the right career time when it can be most expensive. Finally, boards and investors should not green-light large technology investments because proposals are replete with trendy words, such as artificial intelligence, zero trust, digital transformation, cloud-native, etc., and these words are used in an appealing manner. They need baseline measurements, clear named accountable owners, independent validation, staged funding with demonstrated success, explicit risk analysis, explicit exit criteria, and post-implementation measurement. If a cost proposer's proposal is good, he or she will not be charged for these requirements, and that's the idea they're a filter that only artificial credibility can't pass. The common discipline is the same at all levels, Claims must be related to demonstrated understanding, starting with one's own.

15. The Business and Economic Forces Behind Artificial Credibility

Artificial credibility is not only a moral lapse of individuals, but the rational result of certain market incentives and will continue to exist as long as they exist.

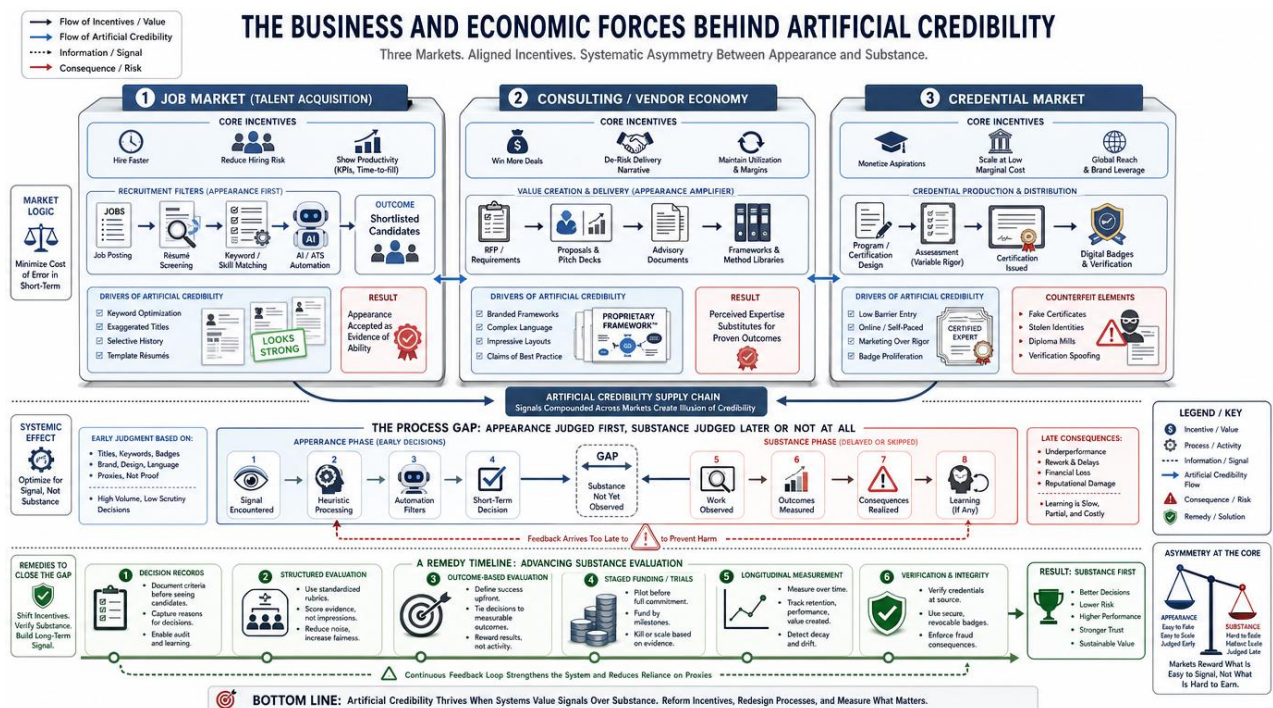


Fig -13: The Business And Economic Forces Behind Artificial Credibility



It is structurally rewarded in the job market. High volume recruitment is based on resume screening, keyword matching, and short interviews tools that assess the artefact not the person, and ones that can now be optimised for directly by generative tools. A candidate who invests in presentation gets more out of each hour invested than does a candidate who invests in substance, for the screening process sees only presentation. Recruiters are typically paid based on placement, not on the placement's performance over time, so there's little reason to dig deeper, and in most organizations there's hardly any feedback loop to correct that: measuring the outcomes of hires back to the screening process is very rare.

There are parallel incentives in the consulting and vendor economy. Many documents proposals, frameworks, roadmaps, maturity assessments are used in advisory work and the cost of creating these documents has been reduced to almost nothing with generative AI. A provider can now create seemingly custom-made deliverables by using templates for a negligible amount of money, and the client has not been able to evaluate depth in the same fashion. There is an additional incentive for vendors to sell AI, in a market where buyers are often unable to validate technical claims, it's competitively advantageous to overstate capabilities, until they don't work out, which can be years after the contracts are inked.

The credential market makes up the rest of it. Trust infrastructure, such as certifications, endorsements and professional profiles, are what these are, and the more valuable they are, the more they are subject to counterfeiting. Because there is an economic pay-off to a fake credential, and the risk of discovery is relatively small, purchased certifications, fabricated references, ghost-written exams and as government prosecutions have revealed even outright credential substitution all exist. Generative AI drastically reduces the cost side of that equation, and the enforcement side has not yet increased the penalty side.

Under all three markets is the more fundamental asymmetry that appearance is judged first, substance is judged later or not at all. Promotions, funding decisions, and contracts are made on presentation time scales, and the effects of incompetence come to light on project and system time scales, often after the people who made the decisions have moved on. Wherever there is a pay-for-appearance policy, there is a consequent pay-for-substance policy. The remedies discussed throughout this article decision records, outcome-based evaluation, staged funding, longitudinal measurement are all best thought of as ways of closing this gap, to move the evaluation of substance forward in time until deception becomes unprofitable. There is a need to have no false hope about awareness, it's incentives, not exhortations, that drive action at scale.

16. Implications for the Future of IT, Automation, Application Development, and Cybersecurity

The path outlined in this article demands a set of adaptations that can be predicted, and those that can make the leap will be better off than those who make the same errors as the present generation of organizations.

The Artifacts Era of Assessment in Hiring and Workforce Management is coming to an end. Resumes, portfolios and eloquence in interviews will continue to lose evidential value and the professions will move haltingly, with some resistance towards demonstration-based assessment, live problem solving, verified histories of outcomes and ongoing identity assurance for privileged professions. The public will learn what it doesn't know yet, that a well-put-together document, a well-read presentation or an impressive profile is nothing on its own. This understanding, which is now widespread, is a social asset, as it enables people to defend themselves as consumers, patients, investors, and voters against the produced power in all

areas, not just technology. The amount of AI-generated code in the application development will continue to increase, shifting, not diminishing, the need for skills. The limited skills are validation, testing, security review, and architectural judgment of what the generated code does, when it's incorrect, and whether it should be in the system or not. Such organizations will end up with defects at the machine rate if they do not validate while staffing. The key pitfall for automation, in general, is to automate over unconfirmed skills. It is an automated process, which means that what its designer knows is automated, and if the designer didn't know it, then it's an industrialization of the error. Preventing disciplines are fundamentals first training and independent review prior to automation.

IMPLICATIONS FOR THE FUTURE OF IT, AUTOMATION, APPLICATION DEVELOPMENT, AND CYBERSECURITY

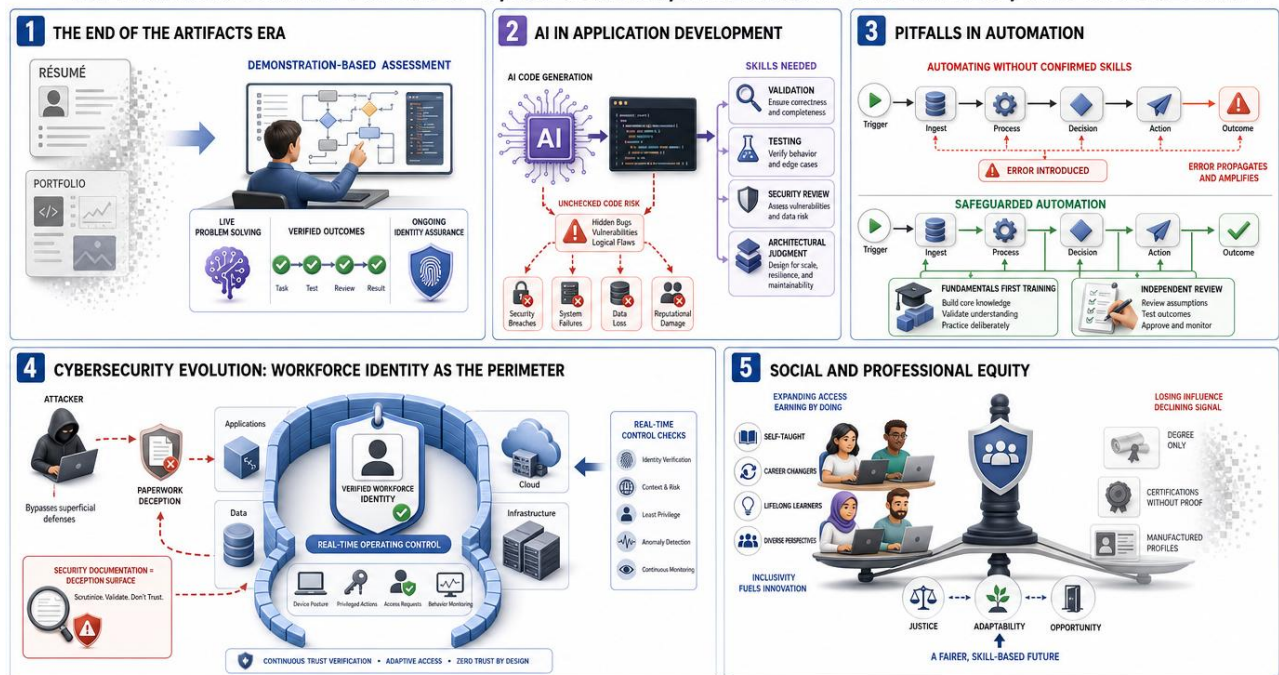


Fig -14: Implications for the Future of IT, Automation, Application Development and Cybersecurity

The impact on cyber security is most pronounced. Workforce identity becomes a first rank security control, as documented fraud schemes have demonstrated that an attacker who is hired through deception will outflank all technical defenses from the inside. At the same time, defenders will have to consider any security documentation they create as a known deception surface, as auditors, regulators and boards will now need to check operating controls, instead of the pretty paperwork, and assurance practices will change accordingly. These adaptations are not only beneficial to prevent fraud, but they also benefit society. In a significant way, a world that values ability over credentials and polish is a more just world than the one it replaces, it gives opportunities to those who are self-taught, to those who are changing careers, to those who are honest learners, from all backgrounds, and it denies opportunities to those who inherited prestige and those who manufactured confidence. What future leaders must not do is what past leaders did at scale: they thought that if people made impressive things, they were the people they needed; they thought that if they rewarded people on the presentation timeline, consequences would be worked out quietly on their own, they thought that if they weren't found out until the outage, the breach, or the audit, that was fine. This lesson has been bought at a high price now. Should not require second purchase.

17. A Ninety-Day Organizational Action Plan

Awareness without implementation will not create any changes, then this article concludes its practical analysis with a phased roadmap that a typical enterprise can take into one quarter, which is divided into three thirty-day movements discover, validate, and institutionalize.

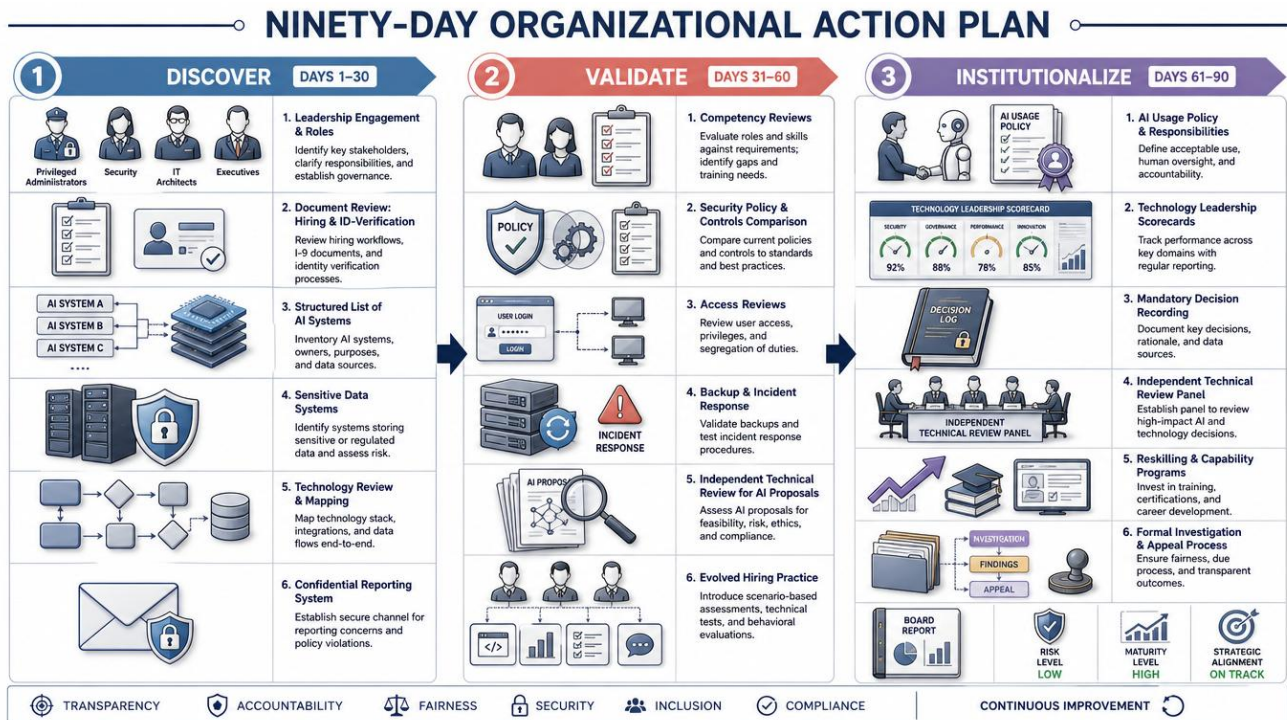


Fig -15: Ninety-Day Organizational Action Plan

The initial 30 days are spent on discovery, because an organization can't manage what it hasn't mapped. Leadership should determine which roles are most important to have artificial credibility privileged administrators, security, architects, and executives who make decisions that involve large amounts of resources. It should take a look at existing hiring and ID-verification procedures, create a list of AI systems that have been deployed and under which guidelines, identify the systems that contain the most sensitive data, review key technology decisions made over the past few years to determine ownership and evidence, and establish a confidential reporting system so that employees who have been quietly working around issues can be heard.

The next 30 days are used to validate the map and make it knowledge. Practical competency reviews should be implemented for critical positions, based on notice, fairness and should be based on reasoning and adaptation, not trivia. Security policies should be compared to the operating controls to identify the differences between maturity in the security policies and the maturity of the operating controls. There should be privileged and remote access reviews; backups and incident-response procedures should not be read but tested major AI-generated proposals that are awaiting funding should be independently technically reviewed before the money moves and hiring should start immediately to include work samples and scenario-based interviews, as new hires are the lowest cost point of control.

The last 30 days are a thirty day version of the first 60 days. The organization should have a policy on the use of AI with the premise that AI can help in the work, but a named human is responsible for the outcome. Technology leadership scorecards should be put in place that are linked to operational results. Significant decisions should be recorded and in the use of decision records and authorship logs should be made mandatory. Access for an independent technical-review panel with access to the board. Reskilling programs should provide an honest way for workers with gaps to be reskilled, and formal investigation and appeal processes should be established to ensure that future concerns are addressed through a proper process and not ad hoc. The quarter should conclude with a report to the board not a reassurance report, but a report on the facts of what was identified, what was addressed, and what will be assessed next quarter. But organisations that go through this cycle will not have eradicated deception, rather they will have made it difficult, detectable, and unprofitable which is what eradication is realistically.

18. RESEARCH GAP

While the literature to support this analysis is real, it is still young and there are several empirical questions that are still open that are important for sound policy. It is important to be aware of these gaps for two reasons for intellectual honesty and to guide future research.

First, there is no longitudinal information on the prevalence and trajectory of competence fraud with the help of AI. Government advice and prosecutions reveal that identity-fraud schemes do exist and can be of significant proportions but only shed light on schemes that are detected. There is no base measurement of the percentage of undetected artificial credibility, and no idea of the number of hires, promotions, and funded proposals that rely on it, and whether the percentage is increasing.

Second, there is no common and validated model available to verify competences in the generative-AI age. The methods presented in this article are based not on instruments validated against impersonation with AI but on professional experience and general assessment science. They are unstudied in terms of their false-positive and false-negative rates, and fairness across languages, cultures, and neurodiverse communication styles.

Third, little empirical research has been done on the misattribution of AI-generated ideas to executives. While credit allocation and impression management have been studied in organizational research for some time, the dynamic of boards ascribing strategy to machine and human strategic genius, and the impact of this dynamic on promotion, compensation, and capital allocation, has yet to be systematically studied.

Fourth, there are immature models for the due process of allegations of AI-assisted deception. There are no reliable, tested procedures yet to determine the difference between fraud and honest use of AI, to prevent an employee from being suspected of discrimination, and to hold up to legal scrutiny. The above sequence of procedures is a rational combination and is not a validated procedure.

Fifth, the economic size of the problem is unknown. Artificial credibility costs failed programs, security incidents, talent attrition, capital misallocation are spread out among budget lines that no accounting system can consolidate, while existing productivity research, while useful, only looks at the effect of tools, not deception effects. The business case for the controls recommended in this article will be based on documented risk until researchers develop ways to isolate and measure these costs. All these gaps



represent their own area of research and advancement in any of these areas would have a tangible impact on fairness and effectiveness of organisational practice.

19. CONCLUSION

Dishonesty is not a creation of Artificial Intelligence. All the above credential fraud, corporate politics, impression management, and incompetent leadership have been around for generations before generative models. What's different is not human nature, it's the economics of deception how fast, how big, how shiny it can be. What used to take years of expertise to simulate can be simulated in minutes and institutions that were founded on the premise that great things mean great people haven't caught on yet.

The modification that this article suggests is not panicking or prohibition. Those that react by curbing the use of AI will be doing a disservice to their honest majority while doing little harm to their dishonest minority, who were never inhibited by policy to begin with. Those that are suspicious in general will lose the flexible workforce that the coming decades will demand remote workers, immigrants, career changers and self-taught professionals, and sophisticated fraud, designed to look conventional, will slip past their suspicions. The right answer is structural, assess actual identity, not claimed identity, assess actual reasoning, not slick recitations, assess actual outcomes, not promised outcomes, assess understanding of trade-offs, not mastery of jargon, assess transparency of the use of AI, not the extent to which it is hidden, assess willingness to accept responsibility, not the ability to evade it. All these attributes are not something that can be created at will which is why they have to be the coin of professional assessment.

Those who turn their backs on artificial intelligence and those who rely on it without thinking will not have a future in the technology professions. It will be the domain of those who have real knowledge, who are continually learning, who are responsible for their tools, who exercise human judgment and who are disciplined enough to know them. In short, what is being said for the last time is that there is no intention to eradicate flawed people from society, all people are flawed, and honest students don't deserve suspicion they deserve help. The goal is to create a system where fraud is impractical, where the developing professional can develop openly, where the honest contributor is rewarded for his efforts, where no executive or engineer can continue to get away with a facade of competence forever, created by an AI. This can be achieved through available tools, available law, and available management practice. What it needs is the courage to ignore the shiny surface of the object and pose the question which is the one that has always distinguished substance from performance, Can you explain it. Can you defend it, Will you answer it.

REFERENCES

- [1] Cernicova-Buca, M., & Palea, A. (2026). Leading through transformation: University responses to generative AI in romanian higher education. *Professional Communication and Translation Studies*, 19, 44-55. <https://doi.org/10.59168/ddjk2862>
- [2] Demirer, M., Musolff, L., & Yang, L. (2026). Writing code vs. shipping code: Productivity effects across generations of AI coding tools. <https://doi.org/10.3386/w35275>
- [3] Egunjobi, J. P. (2025). The misuse of ai-generated content in academic and religious settings. *International Journal of Research and Scientific Innovation*, XII(XV), 871-879. <https://doi.org/10.51244/ijrsi.2025.121500077p>

- [4] Evangeline, S. I. (2025). Policy implications of AI in combatting cyber fraud and identity theft. *Advances in Computational Intelligence and Robotics*. <https://doi.org/10.4018/979-8-3373-4149-1.ch007>
- [5] Johnson, V. (2022). Malleability of source credibility and its impact on knowledge revision (poster 20). *AERA* 2022. <https://doi.org/10.3102/ip.22.1892118>
- [6] K Mogal, A., M Pagar, T., & K Mahale, A. (2026). Ai-driven transformation of labor markets: Skill shifts, hybrid employment, and governance. *International Journal of Science and Research (IJSR)*, 277-282. <https://doi.org/10.21275/sc26211090305>
- [7] Karhu, K. (2002). Expertise cycle – an advanced method for sharing expertise. *Journal of Intellectual Capital*, 3(4), 430-446. <https://doi.org/10.1108/14691930210448332>
- [8] Malleswari, B., & Indira Priyadarshini, Y. (2026). AI based web platform for cyber-fraud prevention in job portals using machine learning. 2026 IEEE 15th International Conference on Communication Systems and Network Technologies (CSNT). <https://doi.org/10.1109/csnt69054.2026.11502345>
- [9] (2020). Limitations of AI. *Artificial Intelligence and the Legal Profession*. <https://doi.org/10.5040/9781509931842.ch-010>
- [10] Henle, C. A., Dineen, B. R., & Duffy, M. K. (2019). Assessing intentional resume deception: Development and nomological network of a resume fraud measure. *Journal of Business and Psychology*, 34(1), 87-106. <https://doi.org/10.1007/s10869-017-9527-4>
- [11] Jain, V., & Mitra, A. (2025). Ensuring compliance and ethical standards with generative AI in fintech. *Generative Artificial Intelligence in Finance*, 253-264. <https://doi.org/10.1002/9781394271078.ch13>
- [12] Maire, Q. (2021). Credential markets and credential theory. *International Study of City Youth Education*. https://doi.org/10.1007/978-3-030-80169-4_13
- [13] Smith, P. H. B. (2014). Document creation, image acquisition and document quality. *Handbook of Document Image Processing and Recognition*. https://doi.org/10.1007/978-0-85729-859-1_3
- [14] (2009). Abbreviations to be used without prior definition in the professional animal scientist. *The Professional Animal Scientist*, 25(6), iv. [https://doi.org/10.15232/s1080-7446\(15\)30797-x](https://doi.org/10.15232/s1080-7446(15)30797-x)
- [15] (2015). Security expertise: An introduction. *Security Expertise*. <https://doi.org/10.4324/9781315744797-9>
- [16] (2016). Law enforcement approaches to the investigation of fraud. *Fraud*. <https://doi.org/10.4324/9781315583075-19>
- [17] (2026). Why artificial intelligence generated clinical content still needs a clinician's eye. *Default Digital Object Group*. <https://doi.org/10.1200/adn.26.300481>
- [18] Burzminski, N. (2003). DEMONSTRATED LEADERSHIP BEHAVIORS AND LEADERSHIP STYLES OF ENTRY-LEVEL DIETITIANS. *Journal of the American Dietetic Association*, 103, 126. [https://doi.org/10.1016/s0002-8223\(08\)70201-3](https://doi.org/10.1016/s0002-8223(08)70201-3)
- [19] Corbin, M. A. (1994). Approaching environmental cleanup costs liability through insurance principles. <https://doi.org/10.21236/ada456706>
- [20] George, D. (2026f). The company Brain-Building institutional memory that outlives employee turnover. *Zenodo (CERN European Organization for Nuclear Research)*. <https://doi.org/10.5281/zenodo.20415478>
- [21] Empson, L. (2017). *Leadership evolution*. Oxford University Press. <https://doi.org/10.1093/oso/9780198744788.003.0008>
- [22] Gottschalk, P. (2007). CIO leadership roles. *CIO and Corporate Strategic Management*. <https://doi.org/10.4018/978-1-59904-423-1.ch003>
- [23] George, D. (2026g). Employee retention in the modern workplace: the role of organizational culture, leadership, and career growth beyond compensation. *Zenodo (CERN European Organization for Nuclear Research)*. <https://doi.org/10.5281/zenodo.20819050>
- [24] Greve, G. (2015). Organizational burnout?. *Organizational Burnout*. https://doi.org/10.1007/978-3-8349-4764-2_1
- [25] Howell, J. M., & Boies, K. (2004). Champions of technological innovation: The influence of contextual knowledge, role orientation, idea generation, and idea promotion on champion emergence. *The Leadership Quarterly*, 15(1), 123-143. <https://doi.org/10.1016/j.leaqua.2003.12.008>
- [26] George, D. (2026d). Security Service Edge (SSE) and SASE: A complete guide to Cloud-Native Zero Trust architecture for enterprise security. *Zenodo (CERN European Organization for Nuclear Research)*. <https://doi.org/10.5281/zenodo.19974566>
- [27] Iskanderova, T. (2024). Analysis of fake news narratives: Exercises. *Unveiling Semiotic Codes of Fake News and Misinformation*. https://doi.org/10.1007/978-3-031-53751-6_5



- [28] George, D. (2026c). Architectural Convergence in Security Operations: a technical framework for AI-Augmented Threat Detection, Automated response, and Organizational cyber resilience. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.19986642>
- [29] Lemieux, P. (2014). Artificial jobs. Who Needs Jobs?. https://doi.org/10.1057/9781137353511_12
- [30] Leška, R. (2026). Paternity right and false attribution (right to attribution). Elgar Encyclopedia of Intellectual Property Law. <https://doi.org/10.4337/9781800886926.paternity.right.and>
- [31] George, D. (2026b). Self-Driving Networks: AI automation for Enterprise IT. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.19335608>
- [32] Lykke, S., Leif, E., & Anders, M. (2008). Measuring leadership talent and training leadership ability. PsycEXTRA Dataset. <https://doi.org/10.1037/e428942008-001>
- [33] Madsen, J. K. (2019). Source credibility. The Psychology of Micro-Targeted Election Campaigns. https://doi.org/10.1007/978-3-030-22145-4_4
- [34] George, D. (2026a). Is India running out of water – An analysis of water bankruptcy, AI data centers, and big finance. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.19259949>
- [35] Olson, J. M., & Ross, M. (1988). False feedback about placebo effectiveness: Consequences for the misattribution of speech anxiety. *Journal of Experimental Social Psychology*, 24(4), 275–291. [https://doi.org/10.1016/0022-1031\(88\)90021-2](https://doi.org/10.1016/0022-1031(88)90021-2)
- [36] George, D. (2026e). Data centers: critical infrastructure, global risk & geopolitical power. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.20443237>
- [37] George, D. (2026h). When Technology works but adoption fails: Leadership capability gaps in enterprise Artificial intelligence transformation. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.20759385>
- [38] PHEYSEY, D. (1973). INDIVIDUAL SKILLS AND ORGANIZATIONAL REQUIREMENTS. *Personnel Review*, 2(1), 14–24. <https://doi.org/10.1108/eb055221>
- [39] George, D., & Siranchuk, D. (2026). The hidden economics of artificial intelligence: why compute costs, token accounting, and human intuition are rewriting the AI value equation. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.20354347>
- [40] Ricchiute, D. N. (1997). Effects of judgment on memory: Experiments in recognition bias and process dissociation in a professional judgment task. *Organizational Behavior and Human Decision Processes*, 70(1), 27–39. <https://doi.org/10.1006/obhd.1997.2691>
- [41] George, D., Dr.T.Baskar, & Srikanth, P. B. (2026). Beyond the perimeter Rethinking enterprise network security in an age of distributed threats. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.20329717>
- [42] Shuju, F., & Aifen, X. (2023). International soft law governance of artificial intelligence: Advantages, approaches, and credibility. *Journal of Business Theory and Practice*, 11(3), p35. <https://doi.org/10.22158/jbtp.v11n3p35>
- [43] Wang, D. (2025). Labor market distortion, agricultural modernization and wage inequality. *Wage Inequality from Agricultural Modernization*. https://doi.org/10.1007/978-981-96-6851-9_5
- [44] George, D., & Dr.T.Baskar. (2026). Degrees, skills, and the Architecture of Human Capital: Reconceptualising India's path to inclusive Economic development. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.20843907>
- [45] Young, E. (2012). Overly optimistic. *New Scientist*, 214(2864), 30. [https://doi.org/10.1016/s0262-4079\(12\)61222-6](https://doi.org/10.1016/s0262-4079(12)61222-6)
- [46] (1995). Team compensation. *Personnel Economics*. <https://doi.org/10.7551/mitpress/5311.003.0007>
- [47] (2007). Theory-based business problem-solving. *Problem Solving in Organizations*. <https://doi.org/10.1017/cbo9780511618413.006>
- [48] (2007). Mentoring. *Encyclopedia of Industrial and Organizational Psychology*. <https://doi.org/10.4135/9781412952651.n182>
- [49] (2013). Practical idea generation. *Basics Graphic Design 03: Idea Generation*, 106–135. <https://doi.org/10.5040/9781350088689.ch.004>
- [50] (2013). Credit card theft. *Encyclopedia of Street Crime in America*. <https://doi.org/10.4135/9781452274461.n43>
- [51] (2019). A false sense of security. Farewell, My Beautiful Homeland. <https://doi.org/10.2307/j.ctvq4c0vm.34>

- [52] (2020). Hafezparast, nasrin, (born 17 sept. 1982), chief clinical data officer, outcome based healthcare, since 2020 (co-founder and chief technology officer, 2013–20). Who's Who. <https://doi.org/10.1093/ww/9780199540884.013.u293066>
- [53] (2024). Encouraging young people to consider a career in law enforcement: Validating ai-generated strategies and ideas. *Archives of Law and Economics*, 1(1). <https://doi.org/10.31038/ale.2024113>
- [54] Adobor, H. (2026). Navigating the AI adoption trap: A framework for organizational practice. *Organizational Dynamics*, 55(2), 101232. <https://doi.org/10.1016/j.orgdyn.2026.101232>
- [55] Anand, B., & Yuxin, H. (2024). Play testing and reflective learning AI tool for creative media courses. *Proceedings of the 16th International Conference on Computer Supported Education*. <https://doi.org/10.5220/0012633600003693>
- [56] Aninakwah, K. A. (2026). Digital transformation and AI adoption in government: Evaluating the productivity gains, implementation barriers, and governance risks. *Journal of Information and Technology*, 10(1), 1–17. <https://doi.org/10.53819/81018102t4368>
- [57] Baer, M., & Zhang, J. H. (2024). Discerning creativity: A group process perspective on idea selection. *Innovation*, 26(3), 387–400. <https://doi.org/10.1080/14479338.2024.2334212>
- [58] Dunn, J. (2005). Why good employees make unethical decisions: The role of reward systems, organizational culture, and managerial oversight. *Managing Organizational Deviance*. <https://doi.org/10.4135/9781452231105.n2>
- [59] Gelles, M. G. (2021). Insider threat prevention, detection, and mitigation. *International Handbook of Threat Assessment*. <https://doi.org/10.1093/med-psych/9780190940164.003.0037>
- [60] Gmyrek, P., Bescond, D., & Berg, J. (2023). Generative AI and jobs. *ILO*. <https://doi.org/10.54394/fvnq9406>
- [61] Hornburg, H. M. (2002). Strategic attack. <https://doi.org/10.21236/ada482915>
- [62] Kim, S., & Kim, H. (2019). An empirical study on investment decision-making factors of personal investment associations: Focused on investment decision-making factors of venture firms. *The Korean Academic Association of Business Administration*, 32(11), 2051–2084. <https://doi.org/10.18032/kaaba.2019.32.11.2051>
- [63] Pandey, A. (2026). Why 60% of transformation programs fail (and no one talks about the real reason). *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.6436800>
- [64] (2022). Government agencies issue stern security warnings. *Computer Fraud & Security*, 2022(2). [https://doi.org/10.12968/s1361-3723\(22\)70019-x](https://doi.org/10.12968/s1361-3723(22)70019-x)
- [65] Abubakari, M. S. (2024). Overviewing the maze of research integrity and false positives within ai-enabled detectors. *Advances in Educational Technologies and Instructional Design*. <https://doi.org/10.4018/979-8-3693-0884-4.ch013>
- [66] Dempsey, K., Ross, R., & Stine, K. (2014). Supplemental guidance on ongoing authorization: Transitioning to near real-time risk management. <https://doi.org/10.6028/nist.cswp.06032014>
- [67] Díez-de-Castro, E., Peris-Ortiz, M., & Díez-Martín, F. (2018). Criteria for evaluating the organizational legitimacy: A typology for legitimacy jungle. *Organizational Legitimacy*. https://doi.org/10.1007/978-3-319-75990-6_1
- [68] Shoir, A., Shakire, U., Sayyora, K., & Timur, S. (2023). Examining the communication competence of psychologists in professional settings. *Proceedings of the 2nd Pamir Transboundary Conference for Sustainable Societies*. <https://doi.org/10.5220/0012964600003882>
- [69] Ugli, D. S. B. (2025). The concept of prosecutorial verification of law enforcement and its regulation in legal acts. *International Journal of Law And Criminology*, 5(6), 40–43. <https://doi.org/10.37547/ijlc/volume05issue06-09>
- [70] (2001). Epistemologically authentic scientific reasoning. *Designing for Science*. <https://doi.org/10.4324/9781410600318-21>
- [71] (2016). The prosecution route. *Fraud*. <https://doi.org/10.4324/9781315583075-46>
- [72] Agndal, H., Arvidsson, A., & Nilsson, U. (2023). Managing appropriation concerns and coordination costs in complex vendor relationships: Integration and isolation as governance strategies. *Industrial Marketing Management*, 113, 116–127. <https://doi.org/10.1016/j.indmarman.2023.05.023>
- [73] Beck, S. (2026). The future of recruitment transformation in the USA: Smarter hiring, faster careers. <https://doi.org/10.55277/researchhub.rfq9zsza.1>
- [74] Gosangi, S. R. (2022). SECURITY BY DESIGN: BUILDING a COMPLIANCE-READY ORACLE EBS IDENTITY ECOSYSTEM WITH FEDERATED ACCESS AND ROLE-BASED CONTROLS. *JOURNAL OF ADVANCED RESEARCH ENGINEERING AND TECHNOLOGY*, 1(2), 1–14. https://doi.org/10.34218/jaret_01_02_001



- [75] Ji, J., Jun, J., Wu, M., & Gelles, R. (2024). Cybersecurity risks of ai-generated code. <https://doi.org/10.51593/2023ca010>
- [76] Lal, S. (2026). AI for cloud cost optimization. Revolutionizing the Cloud. https://doi.org/10.1007/978-3-032-07479-9_13
- [77] Mak, A. K. Y., Ho, S. S., & Kim, H. J. (2016). Hiring cancer survivors outcome expectation scale. PsycTESTS Dataset. <https://doi.org/10.1037/t55536-000>
- [78] Martin, P. (2017). Recruiter incentives and alternatives. Oxford University Press. <https://doi.org/10.1093/oso/9780198808022.003.0008>
- [79] Pieters, W. (2010). Revealing the risks. *Techné: Research in Philosophy and Technology*, 14(3), 194-206. <https://doi.org/10.5840/techné201014321>
- [80] Pinto, L. H., Portugal, R., & Viana, P. (2024). What is in your résumé? the effects of multiple social categories in résumé screening. *Personnel Review*, 53(5), 1331-1358. <https://doi.org/10.1108/pr-04-2023-0290>
- [81] Sravanthi, P. N., & Kumar, M. K. (2025). Verifying food influencers on blockchain for transparency in endorsements and combating fake reviews. *Advances in Marketing, Customer Relationship Management, and E-Services*. <https://doi.org/10.4018/979-8-3693-8542-5.ch004>
- [82] Tiri, K. (2007). Side-channel attack pitfalls. 2007 44th ACM/IEEE Design Automation Conference. <https://doi.org/10.1109/dac.2007.375044>
- [83] (2012). Innovation and the global financial crisis: Systemic consequences of incompetence. Challenging the Innovation Paradigm. <https://doi.org/10.4324/9780203120972-15>
- [84] (2022). Counterfeiting and anti-counterfeiting costs: An application of cost of quality concepts. Determining the value of brand protection programs: identifying and assessing performance metrics in brand protection. <https://doi.org/10.4337/9781839105821.00022>
- [85] (2025). Customer portfolio lifetime value. Customer Portfolio Management. <https://doi.org/10.7551/mitpress/15411.003.0006>
- [86] (2026). Science-based national water assessment: A blueprint for sustainable water management, demonstration case from the republic of korea. <https://doi.org/10.54677/iwez4468>
- [87] Boosey, L., Brookins, P., & Rvkin, D. (2019). Contests between groups of unknown size. *Games and Economic Behavior*, 113, 756-769. <https://doi.org/10.1016/j.geb.2018.09.001>
- [88] Eaton, C. (2026). Chapter 8. ai-generated reading summaries are enough ✨ AI tools can rhetorically support but not replace reading. *Bad Ideas about AI and Writing*; Generative Practices for Teaching, Learning, and Communication. <https://doi.org/10.37514/per-b.2026.2777.2.08>
- [89] Harris, P. L. (2018). Revisiting privileged access. Oxford University Press. <https://doi.org/10.1093/oso/9780198789710.003.0005>
- [90] Hogue, C. (2011). Resources: Boosting efficiency while curbing environmental harm. *Chemical & Engineering News*, 051911142334. <https://doi.org/10.1021/cen051911142334>
- [91] Krishnan, S. (2026). Longitudinal effects of ethical AI pedagogy on AI literacy, moral competence, and educational equity in indian higher education: A three-year mixed-methods study with implications for academic libraries (NEP 2020 alignment). *Information Services and Use*. <https://doi.org/10.1177/18758789261452185>
- [92] L. Canalli, R. (2024). Interpretable AI models for judicial decisionmaking: Beyond explicability towards legal due process. *e-Publica*, 11(1). <https://doi.org/10.47345/v11n1art6>
- [93] López Montiel, Á. G. (2025). Using artificial intelligence to develop research competencies: A tool to avoid academic dishonesty. 2025 Institute for the Future of Education Conference (IFE). <https://doi.org/10.1109/ife63672.2025.11024603>
- [94] Moses, S., Rowe, D. C., & Cunha, S. A. (2015). Addressing the inadequacies of role based access control (RBAC) models for highly privileged administrators: Introducing the SNAP principle for mitigating privileged account breaches. *International Journal of Intelligent Computing Research*, 6(3), 583-591. <https://doi.org/10.20533/ijicr.2042.4655.2015.0073>
- [95] Person, R. (2014). Publishing balanced scorecards and dashboards. *Balanced Scorecards & Operational Dashboards with Microsoft® Excel®*. <https://doi.org/10.1002/9781118984000.ch30>
- [96] Schoenfeldt, M. (2018). Impractical criticism: Close reading and the contingencies of history. *Texts and readers in the Age of Marvell*. <https://doi.org/10.7228/manchester/9781526113894.003.0002>
- [97] Šat, R., Launer, M. A., & Rosak-Szyrocka, J. (2025). Designing effective reskilling programs. *Reskilling and Upskilling in the Age of AI*. <https://doi.org/10.1201/9781003611677-4>