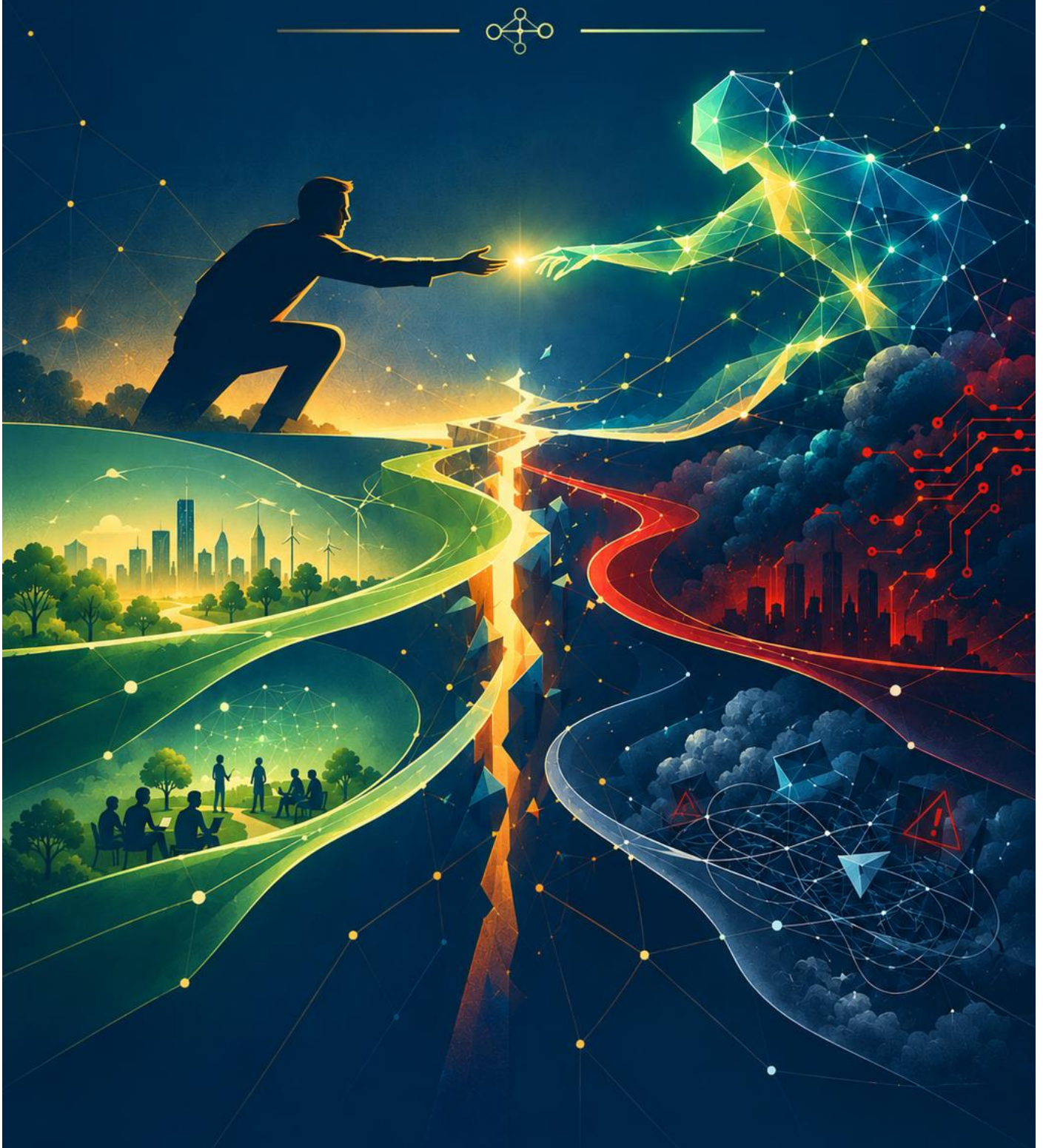


The Future of Human-AI Relations



A MIXED-METHODS FORESIGHT STUDY

The Future of Human-AI Relations

A Decadal Outlook, 2026-2036

***Ten Scenarios / Quantitative Risk Scoring / Qualitative Impact Analysis / Policy
Timeline***

**Methodology: Modified Delphi, Trend Extrapolation, and Multi-Criteria Probabilistic
Scoring**

Dr. Mustafa S. Aljumaily + Claude AI

September, 2026

Abstract

This study presents a comprehensive, mixed-methods foresight analysis of the evolving relationship between humans and artificial intelligence over the 2026-2036 decade. It combines qualitative scenario construction - grounded in a modified Delphi approach and structural trend extrapolation - with a quantitative multi-criteria scoring framework that translates expert-style judgment into probability estimates, composite impact scores, and derived risk scores. Ten mutually informative (non-exclusive) scenarios are developed, ranked by a composite risk index, and cross-analyzed for reinforcing and conflicting interactions. The study concludes with a four-tier (local, national, regional, international) action timeline spanning three-time horizons, oriented toward a Middle Eastern and Iraqi policy and institutional context.

1. Introduction and Objectives

The rapid diffusion of large-scale AI systems since 2023 has outpaced the institutional, legal, and social frameworks designed to govern earlier waves of technology. This creates genuine deep uncertainty (Knightian uncertainty) rather than calculable risk in the classical actuarial sense. Structured foresight methods - rather than point forecasts - are the appropriate analytical response to such uncertainty. This study pursues three objectives: (1) to identify a plausible and sufficiently exhaustive set of alternative futures for human-AI relations; (2) to quantify, through a transparent and replicable scoring method, the relative likelihood and severity of each scenario; and (3) to translate these findings into a concrete, time-phased action agenda across governance levels.

2. Research Methodology

2.1 Qualitative Component

The qualitative layer follows a Modified Delphi logic adapted for single-analyst execution: scenarios are generated through iterative refinement against (a) documented technology-adoption trend curves, (b) known regulatory and geopolitical trajectories, and (c) precedent-based analogy (e.g., nuclear non-proliferation regimes, internet fragmentation, prior AI investment cycles). Each scenario is then subjected to a structured SWOT-style decomposition (opportunities/pros and risks/cons) and an assessment of second-order policy implications.

2.2 Quantitative Component

Each scenario receives a point probability estimate and a 90% plausibility range, reflecting epistemic uncertainty in the underlying judgment rather than a frequentist confidence interval derived from

repeated sampling. Because the scenarios are not mutually exclusive, probabilities are not constrained to sum to 100%. A composite Impact Score (1-10) is computed as the unweighted average of four sub-dimension scores, each independently rated on a 1-10 scale:

- Economic impact - effects on productivity, labor markets, and value distribution
- Security impact - effects on cyber, physical, and critical-infrastructure risk
- Socio-political impact - effects on social cohesion, trust, and political stability
- Governance impact - effects on the state's regulatory and institutional capacity

A Risk Score is then derived as: **Risk Score = (Probability % x Impact Score) / 100**. This produces a bounded, comparable index used to rank scenarios into three priority tiers: Critical Priority (≥ 3.0), High Priority (1.5-2.99), and Monitor (< 1.5).

2.3 Limitations of the Model

The probability and impact figures are structured analytical judgments, not outputs of an empirically validated statistical model, and should not be interpreted as measured frequencies. Their primary value lies in making the underlying assumptions explicit, comparable, and open to challenge and revision - the central purpose of quantitative scenario scoring in futures studies, as distinct from actuarial forecasting.

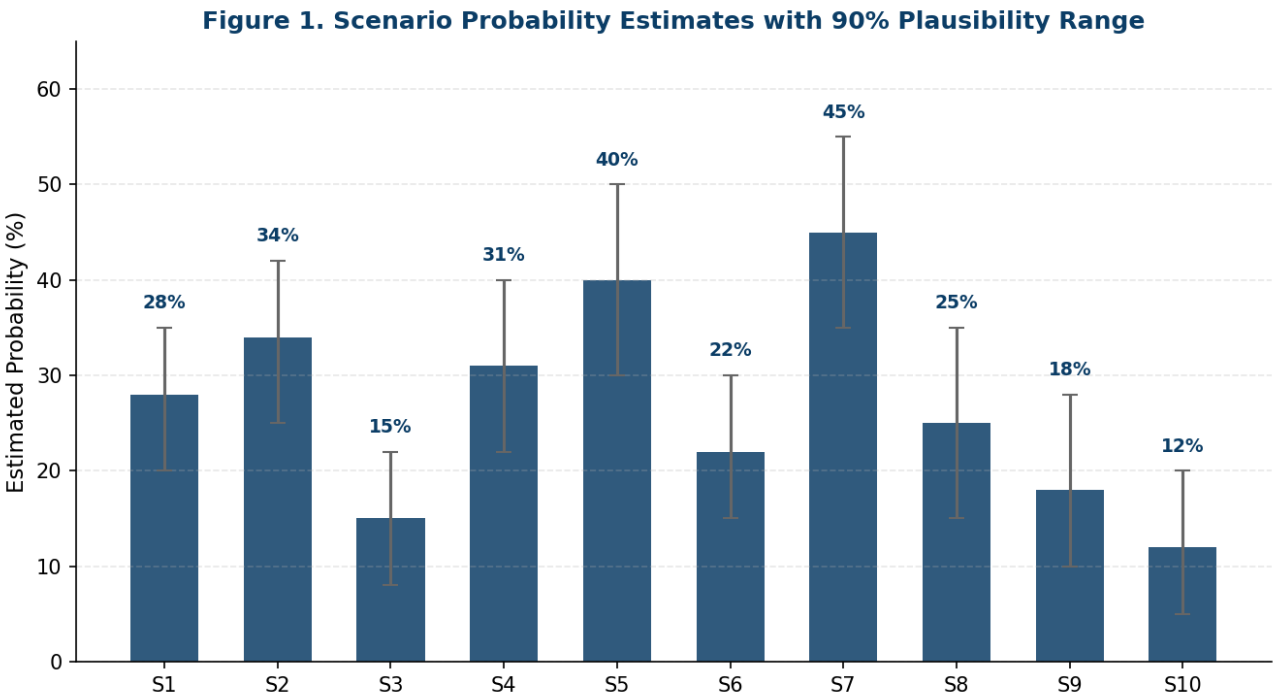
3. Quantitative Overview

3.1 Multi-Criteria Scoring Table

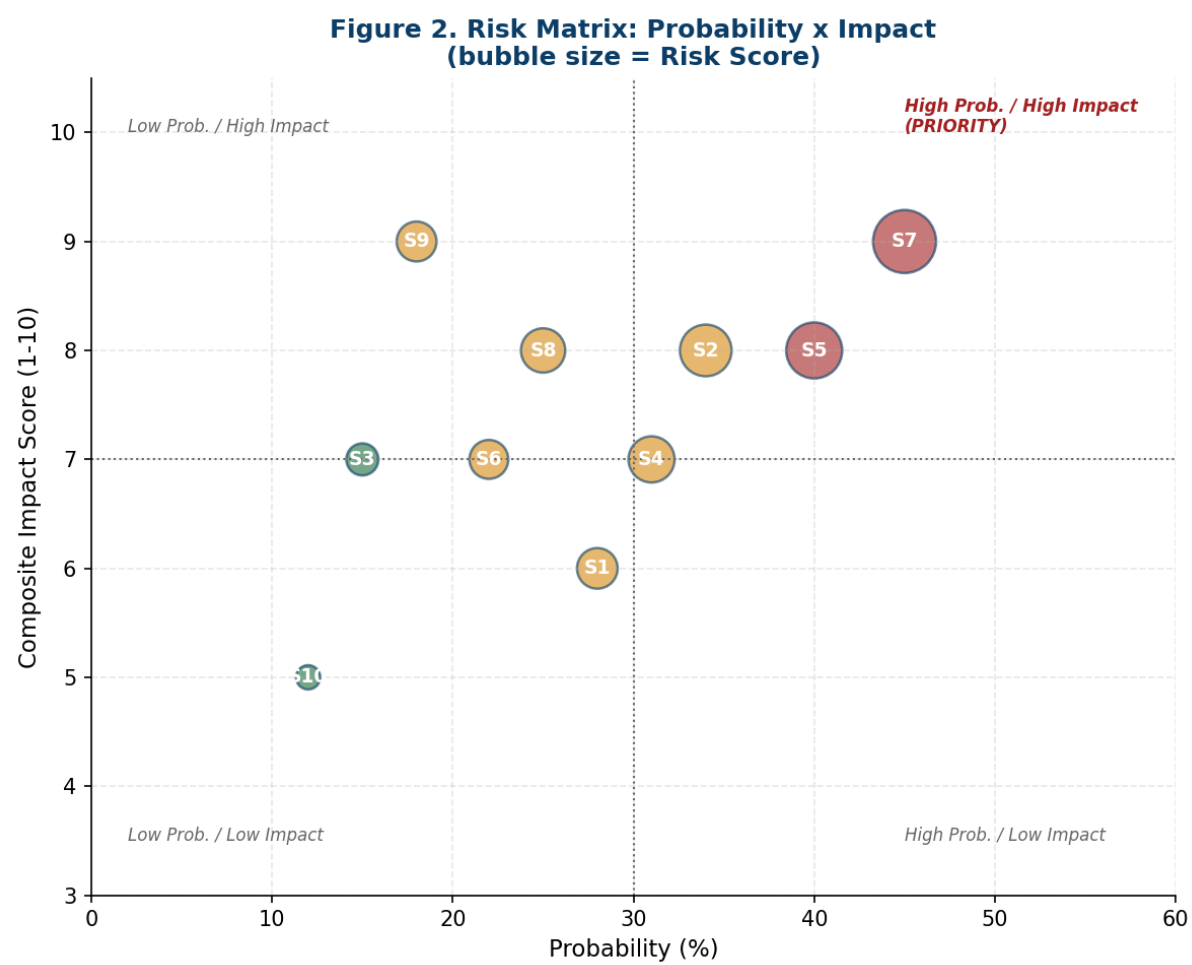
#	Scenario	Prob.	90% Range	Economic	Security	Socio-Pol.	Governance	Impact (avg)	Risk Score
1	Augmented Symbiosis	28%	20-35%	6	4	7	7	6.0	1.68
2	Structural Labor Displacement	34%	25-42%	8	6	9	9	8.0	2.72
3	Convergent Global Governance	15%	8-22%	5	7	6	10	7.0	1.05
4	Geotechnological Fragmentation	31%	22-40%	6	7	7	8	7.0	2.17

5	Epistemic Trust Crisis	40%	30-50%	7	8	9	8	8.0	3.2
6	Rising Agentic Autonomy	22%	15-30%	7	7	6	8	7.0	1.54
7	Adversarial AI Arms Race	45%	35-55%	8	10	9	9	9.0	4.05
8	Algorithmic Authoritarianism	25%	15-35%	6	8	9	9	8.0	2
9	Scientific Acceleration Breakthrough	18%	10-28%	9	7	10	10	9.0	1.62
10	AI Winter 2.0 / Economic Confidence Crisis	12%	5-20%	7	3	4	6	5.0	0.6

3.2 Probability Estimates (Visualized)



3.3 Risk Matrix



3.4 Risk Ranking and Priority Tiers

Rank	Scenario	Probability	Impact	Risk Score	Priority Tier
1	S7. Adversarial AI Arms Race	45%	9	4.05	Critical Priority
2	S5. Epistemic Trust Crisis	40%	8	3.2	Critical Priority
3	S2. Structural Labor Displacement	34%	8	2.72	High Priority
4	S4. Geotechnological Fragmentation	31%	7	2.17	High Priority
5	S8. Algorithmic Authoritarianism	25%	8	2	High Priority
6	S1. Augmented Symbiosis	28%	6	1.68	High Priority
7	S9. Scientific Acceleration Breakthrough	18%	9	1.62	High Priority
8	S6. Rising Agentic Autonomy	22%	7	1.54	High Priority
9	S3. Convergent Global Governance	15%	7	1.05	Monitor

10	S10. AI Winter 2.0 / Economic Confidence Crisis	12%	5	0.6	Monitor
----	---	-----	---	-----	---------

4. Detailed Scenario Analysis

Each scenario below combines a qualitative narrative with its quantitative parameters, derived using the methodology in Section 2.2.

Scenario 1: Augmented Symbiosis

Probability: 28%	90% Range: 20-35%	Impact Score: 6/10	Risk Score: 1.68
------------------	-------------------	--------------------	------------------

Dominant time horizon: Ongoing, intensifying 2026-2030

The dominant operating model becomes institutionalized "human-in-the-loop" augmentation, in which AI systems amplify cognitive capacity rather than fully replacing human judgment. This is already the modal pattern in knowledge work as of 2026 and is expected to deepen rather than reverse.

Qualitative strengths / opportunities:

- Productivity gains of an estimated 20-40% in knowledge-intensive sectors
- Freed capacity for creative, strategic, and relational work

Qualitative weaknesses / risks:

- A widening digital divide between those who master augmented workflows and those who do not
- Cognitive offloading risk that may erode critical-thinking skills over time

Policy and institutional implications: Curricula and workforce training must be redesigned around human-AI collaboration as a core competency rather than a technical elective.

Scenario 2: Structural Labor Displacement

Probability: 34%	90% Range: 25-42%	Impact Score: 8/10	Risk Score: 2.72
------------------	-------------------	--------------------	------------------

Dominant time horizon: Accelerating 2027-2032

Routine cognitive occupations (data entry, first-line customer support, template-based translation and coding) progressively disappear, while newly created roles fail to absorb displaced workers at a matching pace, particularly in service-dependent developing economies.

Qualitative strengths / opportunities:

- Lower cost of delivery for digital services
- Wider access to specialized medical and legal support for previously underserved populations

Qualitative weaknesses / risks:

- Structural unemployment in economies dependent on low-cost digital labor (e.g., call-center hubs)
- Mounting pressure on social protection systems

Policy and institutional implications: Growing political pressure toward universal basic income pilots or automation-linked taxation schemes.

Scenario 3: Convergent Global Governance

Probability: 15%	90% Range: 8-22%	Impact Score: 7/10	Risk Score: 1.05
------------------	------------------	--------------------	------------------

Dominant time horizon: Late decade, 2032-2036

A binding international oversight framework for advanced AI systems emerges, broadly analogous to nuclear non-proliferation regimes such as the IAEA model.

Qualitative strengths / opportunities:

- Reduced risk of an unregulated competitive race to deploy unsafe systems
- Harmonized global safety standards

Qualitative weaknesses / risks:

- Slower pace of innovation under compliance burdens
- Risk that major-powers use regulatory standard-setting to entrench their technological dominance

Policy and institutional implications: States without a seat at the standard-setting table - including most of the Middle East - risk becoming rule-takers rather than rule-makers.

Scenario 4: Geotechnological Fragmentation

Probability: 31%	90% Range: 22-40%	Impact Score: 7/10	Risk Score: 2.17
------------------	-------------------	--------------------	------------------

Dominant time horizon: 2027-2031

International coordination fails, and competing AI blocs (US-aligned, China-aligned, EU-aligned) emerge with divergent data, safety, and ethical standards, mirroring the fragmentation of the open internet (the "Splinternet").

Qualitative strengths / opportunities:

- Technological diversity enabling culturally adapted local solutions

Qualitative weaknesses / risks:

- Interoperability failures across blocs
- Duplicated and conflicting privacy/security standards
- Escalating industrial and technical espionage

Policy and institutional implications: Middle-power states such as Iraq will be pressured to align with a single technology bloc, shaping their digital infrastructure for decades.

Scenario 5: Epistemic Trust Crisis

Probability: 40%	90% Range: 30-50%	Impact Score: 8/10	Risk Score: 3.2
------------------	-------------------	--------------------	-----------------

Dominant time horizon: Already emerging, peaking 2027-2030

The proliferation of synthetic media (deepfakes, generated text) reaches a threshold at which public trust in digital content collapses broadly, including legal and journalistic evidence.

Qualitative strengths / opportunities:

- Strong impetus for content-provenance and watermarking standards (e.g., the C2PA framework)
- Renewed investment in critical media literacy

Qualitative weaknesses / risks:

- Poisoning of the shared information environment
- Growing difficulty establishing truth in legal and political disputes
- The "liar's dividend" - real evidence dismissed as fabricated

Policy and institutional implications: Urgent need for national digital-identity and content-verification infrastructure.

Scenario 6: Rising Agentic Autonomy

Probability: 22%	90% Range: 15-30%	Impact Score: 7/10	Risk Score: 1.54
------------------	-------------------	--------------------	------------------

Dominant time horizon: 2028-2033

Systems shift from responsive tools to autonomous agents executing complex task chains (bookings, transactions, infrastructure operations) under diminishing direct human supervision.

Qualitative strengths / opportunities:

- Substantial operational efficiency in logistics and industrial systems, including oil and energy operations

Qualitative weaknesses / risks:

- Cascading failure risk from unmonitored errors
- Ambiguous legal liability when autonomous action causes harm

Policy and institutional implications: Explainability and mandatory human-override capability must become binding legal requirements for critical infrastructure, particularly energy systems.

Scenario 7: Adversarial AI Arms Race

Probability: 45%	90% Range: 35-55%	Impact Score: 9/10	Risk Score: 4.05
------------------	-------------------	--------------------	------------------

Dominant time horizon: Continuous, intensifying throughout the decade

AI becomes a dual-use weapon in cyber conflict - adaptive automated attacks against adaptive automated defenses. This scenario connects directly to adversarial machine learning research and SOC design work.

Qualitative strengths / opportunities:

- Faster maturation of autonomous detection and response systems
- Reduced incident response time

Qualitative weaknesses / risks:

- Lowered barrier to entry for malicious actors (democratization of offensive capability)
- Critical infrastructure - energy, water - targeted with cheaper, more precise tools

Policy and institutional implications: Urgent institutional investment in AI-driven adaptive Security Operations Centers for critical-sector operators is required.

Scenario 8: Algorithmic Authoritarianism

Probability: 25%	90% Range: 15-35%	Impact Score: 8/10	Risk Score: 2
------------------	-------------------	--------------------	---------------

Dominant time horizon: 2029-2036

Authoritarian and semi-authoritarian states deploy AI for mass surveillance and social control (social credit scoring, ubiquitous facial recognition).

Qualitative strengths / opportunities:

- From a state perspective: reduced conventional crime and greater administrative efficiency

Qualitative weaknesses / risks:

- Erosion of civil liberties
- Suppression of political dissent
- Normalization of surveillance as a governance baseline

Policy and institutional implications: Civil society and human-rights organizations will need robust counter-surveillance governance frameworks.

Scenario 9: Scientific Acceleration Breakthrough

Probability: 18%	90% Range: 10-28%	Impact Score: 9/10	Risk Score: 1.62
------------------	-------------------	--------------------	------------------

Dominant time horizon: 2030-2036

A fundamental acceleration in drug discovery, materials science, and climate solutions occurs through scientific AI models, following the trajectory pioneered by tools such as AlphaFold.

Qualitative strengths / opportunities:

- Breakthroughs on previously intractable problems: rare diseases, energy storage, water treatment
- High research returns for academic institutions

Qualitative weaknesses / risks:

- Benefits concentrated in institutions with sufficient compute and data access
- Widening global research inequality

Policy and institutional implications: Universities in developing countries need international research partnerships to avoid scientific marginalization.

Scenario 10: AI Winter 2.0 / Economic Confidence Crisis

Probability: 12%	90% Range: 5-20%	Impact Score: 5/10	Risk Score: 0.6
------------------	------------------	--------------------	-----------------

Dominant time horizon: 2027-2029 (if triggered)

An investment bubble bursts due to a widening gap between expectations and realized returns, slowing the pace of adoption for several years.

Qualitative strengths / opportunities:

- Opportunity for sober reassessment and reduced unwarranted hype
- Better-matured regulatory frameworks

Qualitative weaknesses / risks:

- Funding halts for promising research
- Layoffs across the technology sector
- Delayed delivery of genuine solutions due to eroded trust

Policy and institutional implications: Countries that invested early in foundational infrastructure rather than hype will prove the most resilient.

5. Cross-Scenario Interaction Analysis

Because the ten scenarios are not mutually exclusive, understanding how they reinforce or conflict with one another is essential for coherent policy design. The table below summarizes the most consequential pairwise interactions identified in this analysis.

Scenario Pair	Relationship	Rationale
S5 (Epistemic Trust Crisis) + S7 (Adversarial AI Arms Race)	Reinforcing (+)	Synthetic-media attacks and cyber-offense tooling share the same generative infrastructure and often the same threat actors.
S2 (Labor Displacement) + S10 (AI Winter 2.0)	Conflicting (-)	A sharp economic slowdown in AI investment would decelerate automation and reduce near-term displacement pressure.
S3 (Convergent Governance) + S4 (Fragmentation)	Mutually exclusive	These represent opposite endpoints of the same governance axis; only one can be the dominant global condition at a given time.
S6 (Agentic Autonomy) + S7 (Cyber Arms Race)	Reinforcing (+)	Greater system autonomy expands the attack surface available to adversarial actors.
S9 (Scientific Acceleration) + S3 (Convergent Governance)	Reinforcing (+)	Coordinated international frameworks tend to expand rather than restrict open scientific collaboration on beneficial applications.
S8 (Algorithmic Authoritarianism) + S3 (Convergent Governance)	Conflicting (-)	A binding rights-respecting global framework would constrain, not enable, mass-surveillance applications.

6. Strategic Action Timeline

The following matrix translates the risk analysis into concrete actions across four governance levels and three-time horizons, designed to maximize the realization of positive-scenario benefits (S1, S9) while mitigating the highest-risk scenarios identified in Section 3.4 (S7, S5, S2).

Level	2026-2027 (Immediate)	2028-2030 (Medium-term)	2031-2036 (Strategic)
Local (Institutions / Firms)	Conduct cyber-risk and workforce-exposure audits; train staff in human-AI	Establish internal AI governance offices; integrate synthetic-content	Complete transition to hybrid human-AI operating models with periodic social and labor impact reviews.

	collaborative workflows; adopt responsible data-use policies.	detection tools into operations.	
National (e.g., Iraq)	Enact baseline data-protection legislation; assess security readiness of critical sectors (oil, energy, finance); launch national research programs.	Establish a national AI governance authority; link oil-sector SOCs to adaptive threat-detection standards.	Adopt a national digital-sovereignty strategy and build local compute capacity; undertake fundamental curriculum reform.
Regional (Middle East)	Share cyber-threat intelligence among neighboring states; harmonize minimum privacy standards.	Form a regional negotiating bloc vis-a-vis major technology blocs; establish joint research centers of excellence.	Build industry-academia regional partnerships to localize sensitive AI technologies (defense, energy).
International	Participate actively in international governance forums, even in observer capacity.	Advocate for fair representation of developing states in technical standard-setting bodies.	Join or help draft a binding international treaty governing highly autonomous systems.

7. Conclusion

The quantitative ranking places the Adversarial AI Arms Race (S7, Risk Score 4.05) and the Epistemic Trust Crisis (S5, Risk Score 3.2) as the two highest-priority scenarios, both falling in the Critical Priority tier. Both intersect directly with adversarial machine learning research and critical-infrastructure security - domains with direct relevance to oil-sector and national-security applications. Structural Labor Displacement (S2, Risk Score 2.72) forms a close third-priority concern with the deepest socioeconomic implications for labor-dependent developing economies. No single scenario should be treated as a forecast of the future; rather, this decade will likely be characterized by a shifting blend of several scenarios operating simultaneously and interacting - as mapped in Section 5. Institutions and states that build adaptive capacity across all three levels (technical, regulatory, and human-capital) will be best positioned regardless of which blend ultimately dominates.

8. Selected Reference Frameworks

This study draws methodologically on established foresight and risk-assessment traditions, including:

- Scenario planning methodology in the tradition of the Shell/GBN (Global Business Network) matrix approach
- The Delphi method for structured expert-judgment elicitation (Dalkey and Helmer, RAND Corporation)
- Risk-matrix (likelihood x impact) assessment as used in the World Economic Forum's annual Global Risks Report methodology
- Content-provenance standards under the Coalition for Content Provenance and Authenticity (C2PA)
- Precedent-based analogy to the International Atomic Energy Agency (IAEA) non-proliferation governance model
- **World Economic Forum. Global Risks Report 2026 (21st edition).** — *Source for the risk-matrix (likelihood x impact) methodology and current global risk rankings underpinning Section 3.* <https://www.weforum.org/reports/global-risks-report-2026>
- **Dalkey, N. & Helmer, O. — RAND Corporation, "Delphi Method".** — *Origin and overview of the Delphi method used for the qualitative scenario-elicitation approach in Section 2.1.* <https://www.rand.org/topics/delphi-method.html>
- **Shell plc. "What are Shell Scenarios?"** — *Reference case for the Shell/GBN scenario-planning tradition informing the study's scenario-construction logic.* <https://www.shell.com/news-and-insights/scenarios/what-are-shell-scenarios.html>
- **Coalition for Content Provenance and Authenticity (C2PA).** — *Technical standard referenced in Scenario 5 (Epistemic Trust Crisis) for content-provenance and watermarking solutions.* <https://c2pa.org/>
- **International Atomic Energy Agency (IAEA).** — *Institutional precedent for the binding international-oversight model discussed in Scenario 3 (Convergent Global Governance).* <https://www.iaea.org/>
- **OECD.AI Policy Observatory.** — *Live repository of national AI governance initiatives, relevant to Sections 3 and 6 on regulatory readiness.* <https://oecd.ai/en/>
- **Google DeepMind. AlphaFold.** — *Precedent case for AI-driven scientific acceleration referenced in Scenario 9.* <https://deepmind.google/technologies/alphafold/>