



Rethinking Intelligence: From Human Cognition to Artificial Futures

Habib Hamam^{1-4*}

¹ Faculty of Engineering, Uni de Moncton, Moncton, NB E1A3E9, Canada, Habib.Hamam@umoncton.ca

² School of Electrical Engineering, University of Johannesburg, Johannesburg 2006, South Africa

³ International Institute of Technology and Management (IITG), Av. Grandes Ecoles, Li-breville BP 1989, Gabon

⁴ College of Computer Science and Eng. (Invited Prof), University of Ha'il, Ha'il 55476, Saudi Arabia

Article Info

Article history:

Received July 16, 2025

Revised August 01, 2025

Accepted August 15, 2025

Keywords:

Intelligent Systems

Human-AI Collaboration

Ethical Frameworks

Interdisciplinary Approach

Adaptive Intelligence

ABSTRACT

The rapid advancement of AI technologies raises pressing questions about the nature and future direction of intelligence. A key challenge is to understand how human and artificial intelligences differ, not just in form but in function, and how they should be evaluated in a shared context. This paper proposes a structured framework based on 15 measurable conditions of intelligence, such as memory, adaptability, specialization, and ethical alignment. Our main contribution lies in connecting these conditions to nine key directions of AI development—such as responsible AI, human-machine collaboration, and quantum AI—to outline how intelligence can be evaluated and guided across both natural and synthetic domains. Methodologically, we cross-analyze these dimensions using a 15×9 matrix, providing both a diagnostic tool and a conceptual roadmap for future AI development. This approach blends insights from cognitive science, applied AI, ethics, and philosophy. Our findings show that intelligence must be judged not just by computational capability but by interpretability, ethical grounding, and social utility. Contextual and hybrid systems—those that adapt to environments and align with human values—emerge as the most promising. We conclude by calling for an interdisciplinary approach to build intelligence systems that are not only powerful but also trustworthy and socially meaningful.

This is an open access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.



1. INTRODUCTION

Over the last decade or so, Artificial Intelligence (AI) has undergone an incredible renaissance, evolving from narrow rule-based/representational systems to highly complex, adaptive models that generate language, synthesize images, make decisions, and drive cars. This technological revolution has not only transformed industries and economies but also provoked profound philosophical and societal questions about what intelligence truly is and how we might create genuine intelligence in a machine. The existence of machine intelligence is no longer limited to research labs; instead, it is woven throughout modern life, from virtual assistants built into everyday gadgets to AI copilots in fields such as medicine, law, and education.

Yet, as AI systems grow in sophistication and autonomy, an enduring question persists: “Who is more intelligent — humans or machines?” Seemingly reasonable as it might be, this question is essentially meaningless. It confuses origins and outcomes of intelligence, and it neglects the context, constraints, and capacity. Intelligence isn't a quality, immutable; it is a theatrical matter, constituted by action, adaptation and consequence. Instead of contrasting human and machine intelligence absolutely, we should ask a more helpful question:

“Under what conditions does intelligence — human or artificial — produce meaningful, effective, and responsible outcomes?”

This paper suggests that the discussion of AI can be refocused by presenting 15 conditions that specify the behavior of any intelligent system, whether biological or artificial. These conditions cover structural

*Corresponding Author

Email: Habib.Hamam@umoncton.ca

abilities such as: (memory) and (multitasking), (speed), cognitive features such as (specialization), (data analysis), and ethical aspects such as (transparency) and (responsibility).

Together, they provide a neutral yet powerful lens through which to analyze both natural and artificial intelligence. Building on this framework, we identify nine primary technological directions in current and emerging AI development, including explainable AI, specialized AI, responsible AI, quantum AI, and hybrid human-machine systems, and examine how they align with the proposed conditions. The purpose is not to outdo winners in an intelligence arms race, but to establish the basis for understanding how intelligence works, how it should develop, and how to ensure that its power is employed responsibly.

Our method reconciles philosophy's insights with the practical wisdom of science. It acknowledges that the question of intelligence is not simply a technical one, but is also ethical, epistemological and human to the core. By connecting these terrains, this contribution makes a case for the pressing need to redefine intelligence for the machine age and direct its course toward socially inclusive and ecologically viable futures.

2. METHODOLOGY: Rethinking the Nature of Intelligence

2.1. Conceptual Foundations of Intelligence - Biological and Artificial

Intelligence is commonly defined in terms of human-like abstract reasoning, creativity, empathy, and language use. Yet, this anthropocentric framing restricts our comprehension of intelligence as an all-existent phenomenon. These even risks overshadowing what the biological evolution, the environmental context and the substrate (biological or artefactual) has contributed to the content and structure of intelligent behaviour.

Korteling et al.[1] that intelligence is not to be measured with respect to its similarity to human cognitive behavior. Instead, they offer a non-anthropocentric explanation of intelligence in terms of "the ability to accomplish complex goals autonomously and efficiently." This more expansive view allows for a wide array of intelligences to exist, including other forms of artificial intelligence, animal minds, or potential extraterrestrial intellects. It focuses more on cognitive tasks driven by goals, rather than being tied explicitly to human or other biological structures.

This perspective is compatible with the growing understanding in cognitive science that natural intelligence is context-sensitive. Based on decades of ecological psychology [2], Hartley [3] insists that human cognition is neither fixed nor universal. Instead, the brain is moulded throughout life by ongoing interactions with diverse and multidimensional environments. Individuals who were raised in a city with a more intricate spatial structure have enhanced spatial navigation skills [4]. Individuals who are exposed to visual tasks with a fast pace, such as video gaming, exhibit better selective attention [5]. These results underscore the fact that the environment plays a role as a scaffolding for the development of intelligent behavior, and that intelligence is best understood within the framework of adaptation and situated interaction [6]-[8].

In this light, human intelligence can be seen as a biologically evolved solution to specific survival and social problems. As Slijepcevic [9] argues, cognition is a biological universal — even bacteria exhibit forms of semiotic agency. Intelligence, therefore, encompasses multiple evolutionary forms and substrates, not just the neural type typically found in humans. Sternberg [10] elaborates on this insight by differentiating between general intelligence and adaptive intelligence, where the latter has a closer connection with evolutionary fitness and the survival of a species by acting in a context-dependent manner.

In the meantime, AI, which is built on digital computation, has proven to be particularly strong in such areas as logical analysis, statistical pattern recognition and large-scale data processing — areas where human cognition is weakest. It was evident that the human brain cannot do anything (Working memory [1], Biased [1], Multitasking [1]) at this scale [1], can operate 24/7/365, totally fatigueless [1], leave alone Consistency, that can survive for ages. AI systems can pull and process terabytes of data in parallel, so you will never get better/faster results from your calculations, and AI systems are deterministic which means that it is deterministic [1] if you monitor the system on the same point you will get a repetitive response comforting that the system is Consistent!

It is in the nature of architecture and task alignment that the real difference lies, rather than simply the nature (biological vs engineered). Biological intelligence is the result of evolutionary pressures that drive survival, sensory fusion, and social bonding. Computational intelligence is described as the set of optimization algorithms that can solve mathematical problems or identify patterns. This leads to two different and possibly compatible intelligence profiles.

Instead of comparing AI to human intelligence and judging its success by the fidelity with which it reproduces human-like reasoning, we argue that intelligence in humans — and in machines — is best evaluated in terms of the degree to which it is capable of pursuing its goals in the context of its constraints, which are often different from our own. That is, intelligence is best assessed as adaptive performance, not as structural similarity.

This ecological and functional view opens the door to more meaningful comparisons between natural and artificial systems. It also calls for a multi-dimensional framework of intelligence performance — one that

accounts for memory, learning, context-sensitivity, interaction, and ethical alignment — which we introduce in the next section through our proposal of 15 foundational conditions.

2.2. The Creator–Creature Paradox

It is so ingrained that the creator must be more intelligent than the creation. Well, how can anything devised by humans ever outdo their creators? This notion, however reassuring, is a human-centric perspective that intelligence is inherently contingent upon anthropomorphic or human-like attributes and origins. However, if we reframe the question from where intelligence comes from to what it can do, that logic begins to fall apart.

As Korteling et al. [1], note, machines created by humans surpass us in specific ways — for instance, in crunching numbers, accessing vast pools of data, and detecting patterns in enormous amounts of information. That doesn't make them any better at “understanding” than we are. Instead, they are optimized to succeed in exploiting our cognitive weaknesses. No one argues that humans are less intelligent simply because AI is now superior in certain domains; instead, it's led us to redefine what “intelligence” is.

Psychologist Robert Sternberg [10] invites us to think of intelligence not as raw IQ, but as adaptive intelligence: the ability to act in ways that support survival and success in a particular context. By this definition, an AI system that helps a pilot avoid a crash or assists in early cancer detection is undeniably intelligent—even if it can't hold a conversation or pass a Turing test.

Slijepcevic [9] goes even further and explains how intelligence arises in nature, even in bacteria. Intelligence, he contends, is not just for brains or neurons. It's a more general concept: the ability of a system to engage in meaningful interaction with its environment.

This broader understanding has support from cognitive science. As researchers like Hartley [3] have noted, the capacity of human intelligence itself isn't constant. It grows over time, influenced by its environment and interactions [2][4][6]. AI systems can, in the meantime, quickly reconfigure themselves based on enormous flows of data, adapting and learning on a scale and at a rate that dwarfs that of human beings.

This is all a rebuke to old “creator vs. creature” logic. It is time, perhaps, to stop asking whether machines can be intelligent and start asking what kind of intelligence do we need? The objective may not be to create machines in our image so much as to design systems that expand our reach, compensate for our limitations and help us navigate a more and more bewildering world.

2.3. Intelligence as Potential vs. Intelligence as Action

The classical sense of intelligence as an interior, relatively fixed essence — usually measured via IQ tests or abstract reasoning — is now seen as grossly inaccurate. Korteling et al. [1] claim human intelligence to be no less specialized, inconsistent, or effective than presumed. Bricks and mortar of cognitive performance are moulded and restricted by general characteristics, for example:

- Limited working memory capacity (10–50 bits/sec),
- Poor cognitive multitasking, often confused with rapid task-switching,
- Rapid decay of declarative memory over time,
- Numerous cognitive biases (e.g., anchoring, hindsight, confirmation) that distort reasoning [1].

This limitation makes it clear that intelligence is not a substance, but a performance (under constraints) — a performance that is a product of the architecture, environment, and task-specific operations of an intelligent system.

From this perspective, biological intelligence is tuned, not tuned optimally. As Hartley [3] and others demonstrate, cognition is significantly influenced by real-world environments, which, over time, shape the development of attention, memory, and learning [4][6][8]. Intelligence is not merely inherited; it emerges through ecological interaction [2], and may differ substantially across individuals and contexts due to developmental or environmental exposure [5][11].

Viewed this way, artificial intelligence is not less intelligent simply because it lacks emotions or consciousness — it is differently intelligent. AI systems frequently outperform humans in domains where biological cognition struggles: large-scale pattern recognition, high-volume memory retention, and rapid, bias-free computation [1]. Their performance stems not from deeper “understanding,” but from structural advantages in speed, scale, and optimization. Narrow in scope, yet highly efficient, AI reveals how alternative architectures can produce powerful — albeit different — forms of intelligence.

And this split is indeed starting to become more apparent in the now-abundant dialog between neuroscience and deep learning. As Saxe et al. [12] have also motivated neuroscientists to challenge their understanding of perception, attention and executive functions. Although artificial as well as biological systems differ in architecture and purpose, comparing their dynamics of learning and internal representations is in fact contributing significant understanding in both domains. 'Once again, this would seem to indicate that intelligence is better considered as a functional and adaptive process and not as some kind of fixed or unique thing.

Similarly, Angermueller [13] et al. show just how foundational deep learning is increasingly becoming in computational biology, e.g., in the domains of regulatory genomics and cellular imaging. These are domains where traditional human analysis simply cannot keep pace with the scale and complexity of modern data. That said, these models are not trying to replicate human thought. Instead, they redefine what effective performance looks like, discovering hidden patterns and making confident predictions in ways that go far beyond conventional cognitive strategies.

Together, these perspectives reinforce the idea that intelligence is best judged by what a system can accomplish within the boundaries of its design and environment — not by how closely it imitates human reasoning. Artificial systems and biological minds each operate within their own unique constraints and strengths. Indeed, a fair assessment of intelligence must consider the ability to act effectively under those constraints, whether the system runs on neurons or transistors.

This understanding sets the stage for a more structured exploration: a framework of conditions that define what makes intelligence — human or machine — truly performative. That framework is the focus of the next section.

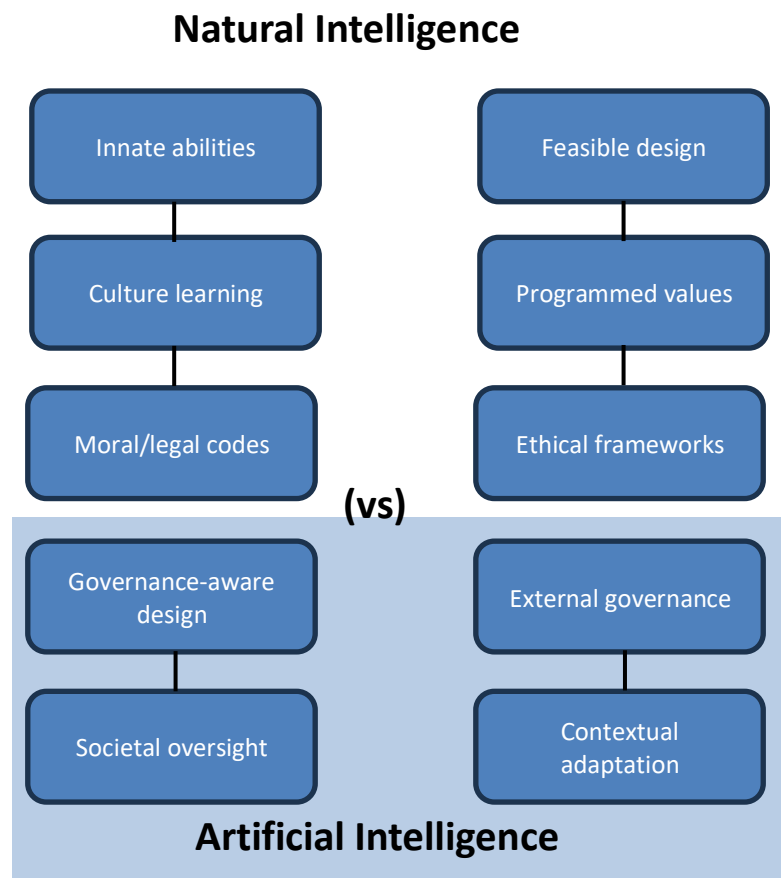


Figure 1. Natural vs Artificial Intelligence Performance Dimensions

Figure 1 illustrates a conceptual framework that maps 15 key conditions of intelligence performance across nine emerging directions in AI development. These conditions are classified as either structural (e.g., working memory, computation speed) or functional (e.g., data access, interpretability, ethical framing), and their influence is evaluated in both natural and artificial systems.

The diagram emphasizes intersections — such as how domain specialization, pattern correlation, and computational speed are core to Specialized AI, or how interpretability and ethical framing underpin Responsible AI. It also highlights societal and technical enablers, such as Contextual Adaptation—the ability of AI systems to adjust their behavior based on situational inputs and evolving norms—and Programmed Values, which reflect human ethical input at the code level.

Think of it like designing a smart home. Specialized AI is like your thermostat: it knows exactly how to regulate temperature efficiently, but nothing else. Responsible AI, on the other hand, is akin to an intelligent security system that not only detects intrusions but also adheres to household rules about when and how to notify authorities — ensuring both privacy and fairness. Contextual Adaptation is your home's ability to adjust

lighting or music based on who's in the room and what time it is. At the same time, Programmed Values are the family guidelines encoded into the system; for instance, never unlock the door remotely for unknown visitors, regardless of the situation.

This visual synthesis serves as a navigational aid for aligning AI development with human values, revealing how different performance conditions relate to distinct AI directions, and ultimately proposing a holistic benchmark for evaluating intelligence—natural or artificial—not by its origin but by its societal impact and adaptability under constraints.

2.4. From “Who is More Intelligent?” to “How Can Intelligence Be Used Effectively?”

“Who’s smarter — people or machines?” is sort of the wrong question. As Korteling and his group [1] say, it’s not so much that one is better than the other overall, but that it’s becoming clear how humans and AI can complement each other. The real question is: how can AI help us do things better, especially where we are not very good?

This perspective has the benefit of allowing us to stop comparing machines to humans all the time. Our brains evolved to help us survive, not to be the most rational thinkers around. We misremember, we’re easily distracted, and often make decisions based on intuition rather than logic. Machines, on the other hand, don’t get tired or emotional — they’re quick, precise, and excellent multitaskers. It’s like using a power drill instead of twisting a screwdriver by hand — the tool doesn’t replace the person; it just speeds up and eases the task.

This concept also appears in brain science. Macpherson and colleagues [14] discuss how AI and neuroscience are learning from each other. Early AI systems tried to mimic how brain cells work. Today, AI is actually helping scientists understand how our brains function — how we see, or what happens when mental health conditions affect the mind. But AI doesn’t need to think like a brain to be useful. Think of it like using GPS: the machine doesn’t think like you do, but it helps you get where you’re going faster and without getting lost. And that’s what makes AI so helpful — it supports us even if it “thinks” in an entirely different way.

So really, trying to build AI that thinks exactly like a human might not be the most important goal. What matters is building systems that work with us, fill in our gaps, and do their jobs safely and responsibly. The near future lies in diverse, specialized, and interconnected narrow AI systems, built to tackle the weaknesses of biological cognition — an approach captured in our proposed 15 conditions for intelligence performance and 8 key AI directions.

Ultimately, intelligence is not about who wins a cognitive contest. It’s about designing systems — human, machine, or hybrid — that perform effectively, ethically, and adaptively within a given context.

3. Analysis: The 15 Conditions of Intelligence Performance

The effectiveness of any intelligent system — whether human or artificial — depends not merely on its origin or substrate, but on how well it performs under real-world constraints. Based on an integrative analysis of cognitive science, AI systems design, and philosophical reasoning, we identify 15 essential conditions that govern the performance of intelligence.

3.1. Structural vs. Functional Conditions”

We propose a distinction between:

- Structural conditions: Intrinsic attributes embedded in the architecture of the system (biological or artificial).
- Functional conditions: Operational capacities that emerge through training, adaptation, interaction, and contextual deployment.

This categorization enables us to compare intelligences not by their nature, but by the performance enablers they possess.

3.2. Structural Conditions

Structural properties refer to the specific characteristics that define the architecture and the set point of an intelligent, biological, or artificial system. These can be limitations or enablers such as memory, design lineage, fatigue resistance, and operation speed. They tend to be more stable and basic than function conditions, and can affect how intelligence can be made, maintained, and even born out over the course of your life. Table 1 summarizes these structural aspects, highlighting how each condition differs between human intelligence and artificial systems. It emphasizes that while machines benefit from unlimited uptime, high-speed processing, and vast memory, they are entirely dependent on human design, collaboration, and programming intent.

Table 1. Structural Conditions of Intelligence Performance

#	Condition	Definition & Rationale	Implications (Human vs AI)
1	Designer Category	The intelligence level of those who design or train the system.	Humans build AI; AI reflects the intelligence of its creators.
2	Team-Based Design	Collective intelligence from multidisciplinary teams.	Most AI is a product of collaboration, not a lone genius.
3	Domain Specialization	Intelligence performs better when fine-tuned for specific tasks.	AI excels when purpose-built; humans benefit from expertise.
4	Time Availability	Time invested in learning, processing, or building intelligence.	Machines have infinite uptime; humans are time-limited.
5	Fatigue Resistance	Ability to function without cognitive or physical exhaustion.	AI systems don't tire; humans do.
8	Working Memory	Ability to retain and use information in short-term.	AI systems have large RAM; humans are limited to ~10–50 bits/s.
10	Speed of Computation	Latency and throughput of cognitive operations.	Machines far exceed humans in processing speed.
11	Parallel Processing / Multitasking	Ability to run multiple processes simultaneously.	Humans are poor multitaskers; AI thrives in parallelization.

3.3. Functional Conditions

Functional requirements specify the behavioral capabilities that explain how intelligence is utilized at runtime. Among these are data access, the ability to learn, detection of correlations, and interpretability. Unlike structural conditions, functional conditions can change with training, adaptation, or external perturbations. They determine how well an intelligent system can operate in typical everyday situations, when interacting with people, or when required to navigate complex ethical and social situations. These functional conditions and their implications in human and artificial intelligence are summarized in Table 2. This table emphasizes AI systems as being well-suited to data heavy tasks, correlation analysis, and continuous retraining, but also shows the importance of human-machine fluency and ethical governance frameworks.

Table 2. Functional Conditions of Intelligence Performance

#	Condition	Definition & Rationale	Implications (Human vs AI)
6	Access to Large Datasets	Quantity and quality of data available for learning or decision-making.	AI can access and process massive datasets; humans cannot.
7	Data Mining Capacity	Ability to extract patterns, correlations, or insights from data.	AI models can identify deep correlations; humans rely on heuristics.
9	Pattern Correlation Abilities	Capacity to identify hidden or non-obvious relations across data points.	AI excels in statistical patterning; humans in intuitive linking.
12	Continuous Learning / Evolution	Ability to improve over time with new data or experiences.	Humans evolve biologically and socially; AI retrains rapidly.
13	Human-Machine Interaction Fluency	Efficiency and intuitiveness of communication between humans and AI.	A critical design challenge for intelligent interfaces.
14	Interpretability & Transparency	Capacity to explain decisions and ensure traceability of reasoning.	AI often lacks this clarity; human reasoning, though biased, is narratively coherent.
15	Ethical and Regulatory Framing	Integration of moral, legal, and societal safeguards into decision-making.	Humans are held accountable; AI must be governed externally.

3.4. Summary

These 15 conditions provide a unifying framework for assessing, contrasting, and tuning natural and artificial intelligence. Whereas structural conditions determine what a system can do, functional conditions indicate how well it can be expected to perform in the real world. Interpreted in this way, intelligence makes sense as a capacity insofar as it proves itself in action under constrained circumstances, and this view focuses our thinking on a more productive appreciation of AI-human complementarity.

The quantity 15 is not a fixed constant – it is selected to bring out the dominant and common elements of intelligence performance. Some of those conditions may fall under more general classifiers, while others may be categorized by specific domain or application context. This framework is therefore flexible and extensible, intended to support more profound conceptual clarity and practical guidance.

4. Results: Nine Technological Directions of AI

Building upon the 15 conditions of intelligence performance, we identify nine emerging directions shaping the future of artificial intelligence. Each direction emphasizes a distinct pathway in which AI systems are being developed or applied, reflecting different constellations of structural and functional conditions. For each direction, we describe its scope, connect it with the relevant conditions, and illustrate a representative

application. This framework also opens a space for anticipating key challenges and opportunities for future research and regulation.

4.1. Specialized AI

Specialized AI refers to systems designed and optimized to perform narrowly defined tasks with high precision and efficiency. Unlike Artificial General Intelligence (AGI), which aspires to broad, human-like cognitive capabilities, specialized AI thrives in constrained domains where the problem is well-defined, the data is structured, and the objectives are specific. These systems do not attempt to mimic the full range of human reasoning but instead excel at focused functions such as detecting patterns, making predictions, or classifying data within a specific context [15] [16].

The strength of specialized AI lies in its narrow scope. By focusing on clearly bounded tasks—such as diagnosing pneumonia from X-rays, translating between languages, or flagging fraudulent transactions—it can achieve levels of performance that often exceed human capabilities [17] [18]. For instance, in healthcare, specialized AI tools are revolutionizing medical imaging, pathology, and personalized treatment planning. These systems can analyze thousands of scans or patient records in seconds, aiding clinicians with faster and more accurate decision-making [17] [18].

Beyond healthcare, specialized AI plays a vital role in domains such as education and manufacturing. In classrooms, AI-powered applications personalize learning experiences and support students with disabilities through adaptive learning systems and communication aids [18] [19]. In industrial settings, AI is used for real-time monitoring, quality control, predictive maintenance, and supply chain optimization—leading to gains in productivity and efficiency [15].

Ultimately, specialized AI is not a step toward general intelligence, but rather a powerful and practical solution for real-world challenges in domains where deep data and task-specific expertise are available. Its success depends not on versatility, but on precision, scalability, and contextual understanding [16].

Related Conditions:

This direction aligns with several core conditions from our framework (Tables 1 and 2):

- [Condition (Cond) 3] Domain Specialization – Purpose-driven optimization increases relevance and performance.
- [Cond 6] Access to Large Datasets – Training specialized systems requires extensive and high-quality data.
- [Cond 9] Pattern Correlation Abilities – Statistical modeling uncovers non-obvious patterns, enhancing diagnosis or classification.
- [Cond 10] Speed of Computation – Enables real-time responsiveness critical in high-stakes environments.

Application Example:

A notable example is the application of convolutional neural networks (CNNs) in medical imaging diagnostics. CNNs can detect diabetic retinopathy, lung nodules, or tumors with precision that rivals or surpasses trained radiologists. According to Khalifa et al. [20], AI-enabled clinical decision support systems have demonstrated excellence across six domains—including diagnostic assistance, EHR analysis, and risk prediction—making them powerful tools in personalized healthcare.

In diagnostic modeling for medical imaging, AI has enabled the early detection of diseases such as breast cancer and neurological diseases through the detection of imaging markers that go unnoticed by humans. Likewise, forward looking analytics within specialized AI models can predict patient decompensation or surgical complications [20].

Future Implications and Open Challenges:

The increasing effectiveness of task-specific AI raises new concerns regarding generalizability, or the lack thereof. These systems are rarely very flexible and tend to generalize poorly across domains or in cases of contextual ambiguity. In addition, the potential of missing social and ethical implications cannot be excluded when specially trained models are sent to work on their own.

As Khalifa et al. [20], achieving success in niche applications - like clinical diagnostics or personalized therapy – requires the combination of openness in design, ethical oversight, and the collaboration of technologists with clinicians. The way forward is to combine interpretability, regulation, and adaptive learning to obtain AI that allows for meaningful contributing to secure and effective decision-making in the highest impact domains.

4.2. Personalized AI

Personalized AI systems are tailored to individual users by adapting to their behaviors, preferences, contexts, and needs. These systems aim to enhance user experience, decision-making, and relevance through

dynamic customization. Rather than a one-size-fits-all model, personalized AI is context-sensitive and user-centric.

Related Conditions:

Personalized AI aligns with several conditions in our framework, especially (Tables 1 and 2):

- [Cond 1] Designer Category: Designers must integrate user diversity into system behavior.
- [Cond 2] Team-Based Design: Interdisciplinary approaches are required to ensure personalization is ethical and inclusive.
- [Cond 4] Time Availability: Continuous user interaction enables long-term adaptation.
- [Cond 8] Working Memory: Local or cloud-based memory mechanisms help retain context over sessions.

Typical Application Example:

Recommendation engines in platforms like Netflix or Spotify use real-time learning to suggest content based on viewing/listening history, temporal patterns, and even mood signals.

Future Implications / Open Challenges:

One unsolved problem is that of user autonomy and privacy versus personalization. Personalized AI systems might even accidentally reinforce biases or filter bubbles. Transparent, ad hoc adaptation mechanisms and user-controlled customization interfaces are increasingly being recognized as key elements.

Recent studies recognize the growing importance of machine learning tools in educational platforms, which infer learning styles on the fly to adapt content delivery to a student's learning preferences. Overviews of the work discussed in Reference [21], along with other relevant research, emphasize the shift from static profiling methods (e.g., questionnaires) toward behavior-based modeling using deep neural networks and hybrid algorithms. This shift aims to make learner profiles — and the models built on them — even more adaptive and relevant.

4.3. Human–Machine Hybrid Intelligence

Hybrid intelligence refers to the collaborative potential between human and machine intelligence, where each compensates for the other's limitations to achieve superior outcomes. Rather than competing, humans and AI systems work side by side—each doing what they do best. The idea is simple but powerful: combine human intuition, empathy, and contextual judgment with machine-level speed, logic, and pattern recognition to create systems capable of achieving what neither could accomplish alone [22].

This collaboration takes many forms, from human-led processes where AI offers decision support, to machine-led systems with human oversight, and even fully integrated partnerships. The allocation of tasks typically depends on the strengths required—humans bring insight and creativity, while machines provide computational efficiency and data-driven precision [23]. For instance, in medical image classification or tumor detection, hybrid models blend deep learning with clinical oversight to ensure both accuracy and ethical responsibility [24] [25].

A key model in this paradigm is the "human-in-the-loop" approach, where humans remain involved in critical steps of the process. This ensures not only the reliability of AI outputs but also enhances user trust and accountability. This is especially important in domains where high-stakes decisions—such as healthcare, autonomous driving, or customer service—require nuance, ethics, or empathy [22].

Ultimately, hybrid intelligence is not about replacing human cognition but augmenting it. As emphasized in recent studies, mutual augmentation—where both humans and machines evolve through collaboration—represents a promising direction in the development of AI systems that are not only powerful but also socially and contextually aware [22] [26].

Related Conditions:

This direction resonates with several key conditions, including:

- [Cond 4] Time Availability: AI extends cognitive support beyond human time limits.
- [Cond 5] Fatigue Resistance: Machines assist humans in sustained or repetitive tasks.
- [Cond 7] Data Mining Capacity: AI extracts patterns; humans interpret them meaningfully.
- [Cond 11] Parallel Processing / Multitasking: Division of labor allows hybrid teams to handle complex tasks efficiently.
- [Cond 13] Human–Machine Interaction Fluency: Effective collaboration depends on intuitive, adaptive interfaces.

Typical Application Example:

In clinical settings, hybrid diagnostic systems support doctors by flagging potential anomalies in medical images or suggesting personalized treatment plans—combining data-driven analysis with human expertise.

Future Implications / Open Challenges:

We are only at the beginning of building trustworthy and seamless interfaces between humans and machines. When perceptions of smart machines diverge from their actual capabilities — or when we come to rely too heavily on them — problems arise. Effective hybrid team design requires shared context, explainability, and dynamic task allocation.

4.4. Explainable AI

Explainable Artificial Intelligence (XAI) is an emerging area of research focused on making AI systems and their decision-making processes transparent, interpretable, and understandable to humans. Unlike traditional “black box” models, XAI techniques enable stakeholders to trace how an AI system arrives at its conclusions, thereby fostering trust, accountability, and informed decision-making [27].

Recent advancements highlight the practical value of XAI in real-world applications. For instance, in the domain of clean energy policy, Khan et al. [28] proposed ExplainableClassifier, a machine learning model that classifies the eligibility of Clean Alternative Fuel Vehicles (CAFEVs). The model integrates XAI techniques like SHAP and LIME to provide transparency into feature contributions and ensure interpretability of each decision. Beyond model performance, the study emphasizes data de-duplication and ethical reclassification of ambiguous cases, illustrating how XAI can serve as a bridge between technical accuracy and regulatory fairness.

Similarly, Ahmed et al. [29] applied XAI to the bioinformatics domain by leveraging SHAP visualizations and TabNet interpretability tools to decode the regulatory role of pseudogenes. Their hybrid deep learning framework, which combines autoencoders, cGAN-based data augmentation, and TabNet classification, demonstrates how explainability can support large-scale genome annotation while remaining accessible to biologists via interactive tools like Gradio.

These examples underscore XAI’s growing relevance across disciplines—from electric vehicle policy to genomics—where clarity, traceability, and ethical grounding are essential. As AI systems continue to permeate high-stakes environments, explainability will not only help refine algorithms but also enhance societal acceptance and regulatory integration.

Related Conditions:

This direction aligns closely with:

- [Cond 1] Designer Category: The quality of explainability depends on the intelligence and ethical awareness of system architects.
- [Cond 3] Domain Specialization: Interpretation must be tailored to the field of application (e.g., medicine, law, finance).
- [Cond 12] Continuous Learning / Evolution: Models should update explanations as they evolve.
- [Cond 14] Interpretability & Transparency: The core condition driving this direction.
- [Cond 15] Ethical and Regulatory Framing: Regulatory compliance increasingly demands explainability (e.g., GDPR).
-

Typical Application Example:

In credit scoring, an explainable AI system clarifies why a loan application was approved or rejected, identifying factors like income, repayment history, or debt ratio—ensuring transparency and avoiding hidden biases.

Future Implications / Open Challenges:

Balancing accuracy with interpretability remains difficult, especially for deep learning models. There is also a growing need for domain-specific explainability metrics and human-centered design of explanations to accommodate users’ cognitive and emotional contexts.

4.5. Frugal / Edge AI

Frugal AI—often implemented through Edge AI—focuses on building models that are resource-efficient in terms of computation, energy, memory, and bandwidth. These systems are designed to operate locally (on the “edge” of the network), without relying heavily on cloud infrastructure, making them suitable for real-time, low-power, or remote applications.

Related Conditions:

This direction mainly involves:

- [Cond 4] Time Availability: Edge AI reduces latency by processing data in real time at the source.
- [Cond 10] Speed of Computation: Efficient processing is essential for low-resource settings.
- (Optional tie-in with [Cond 5] Fatigue Resistance): Devices run 24/7 without needing human supervision, like AI's fatigue immunity.

Typical Application Example:

A wearable health monitor utilizes an edge-based AI model to detect anomalies in heart rate or oxygen levels in real-time, eliminating the need to send data to the cloud, thereby ensuring privacy and responsiveness.

Future Implications / Open Challenges:

As discussed in Reference [30], Edge AI is full of promise—think autonomous cars, factory robots, smart homes that actually feel smart, and health monitors that don't nag—but, of course, the road ahead isn't all smooth. One of the biggest hurdles? Squeezing large, complex models onto devices with limited power and memory. It's like trying to run a marathon in flip-flops—technically possible, but far from ideal.

There's also the matter of energy efficiency. These edge devices need to perform real-time AI tasks without draining their batteries faster than you can say "optimization." Add to that the need to function reliably in ever-changing network environments—sometimes with perfect 5G, sometimes with... well, less than ideal connectivity.

Another layer of complexity? The edge ecosystem itself. We've got MEC, fog computing, cloudlets—you name it. But unless these approaches start speaking the same language, scalable deployment remains a dream rather than a reality. As Reference [30] points out, the future of Edge AI will depend heavily on progress in standardization, security, and edge-aware training techniques.

And, of course, no frugal AI conversation is complete without discussing energy-aware design. A recent study from the VUBETA Journal [31] demonstrates how energy-aware node selection using hierarchical clustering in multilevel IoT-based wireless sensor networks (WSNs) significantly improves energy consumption without compromising functionality—precisely the kind of optimization that next-generation edge applications will need to thrive.

4.6. Specialized Generative AI

Specialized Generative AI refers to models like GPT, Stable Diffusion, or any of those clever domain-specific LLMs trained to generate meaningful content—but not in a "write me a poem about cheese" kind of way. We're talking serious stuff here: legal contracts, architectural blueprints, molecular structures. The idea is that instead of trying to know a bit about everything (hello, generalist models), these systems are finely tuned for high-stakes, high-context applications—like drafting legal arguments or generating drug candidates [32].

This specialization isn't just for show. It helps reduce hallucinations, makes the outputs more relevant, and generally helps these models feel less like overconfident interns and more like reliable collaborators. Of course, there's still a long way to go—interpretability, domain transferability, and regulatory integration are just a few of the open tabs on the to-do list. But if the current momentum holds, specialized generative models could become indispensable partners in fields where accuracy, consistency, and context aren't optional—they're mission-critical.

Related Conditions:

- [Cond 3] Domain Specialization: Training on a targeted corpus enhances accuracy and relevance.
- [Cond 6] Access to Large Datasets: Generative models require abundant, high-quality domain-specific data.
- [Cond 9] Pattern Correlation Abilities: Generative AI excels at leveraging learned associations for synthesis.
- [Cond 10] Speed of Computation: High-speed inference enables real-time creativity and interaction.

Typical Application Example:

A generative model trained on biomedical literature proposes novel protein sequences for vaccine candidates, supporting researchers with hypotheses that would take years to derive manually.

Future Implications / Open Challenges:

Despite its promise, specialized generative AI still encounters challenges such as hallucination, data contamination, and explainability. Ensuring that outputs are safe, accurate, and ethically aligned continues to be an open area of interdisciplinary research.

Specialized AI vs. Specialized Generative AI

While both directions focus on narrow-domain optimization, their core purpose differs. Details can be found in [Table 3](#).

Table 3: Comparison - Specialized AI vs. Specialized Generative AI

Criteria	Specialized AI (4.1)	Specialized Generative AI (4.6)
Goal	Decision-making, classification, optimization	Content generation (text, image, code)
Output Type	Discrete actions or labels	New content or artifacts
Key Use Case	Fraud detection, quality control, diagnostics	Text generation, image synthesis, creative design
Architectural Focus	Often rule-based or task-specific neural models	Transformer or diffusion models adapted to domain
Main Challenge	Accuracy and generalization	Coherence, factual consistency, safety

Thus, while both directions align with condition [2] Specialization, the generative aspect introduces new concerns, especially related to creativity, unpredictability, and ethical risk.

4.7. Autonomous Intelligent Systems

When you apply this idea to machinery and code, an autonomous intelligent system is like “set it and let it go,” the equivalent of fire-and-forget AI. These are machines or software agents that perform tasks, make decisions, and improve as they go — and they do it with little to no human input. We’re talking about systems that are a blend of sensing (what’s going on?), rational thinking (what is that anyway?), strategy (what may I do?), and action (go do it!).

They’re designed to travel through complex, unpredictable conditions — think driverless cars dodging potholes, drones delivering parcels through wind and rain, or factory robots making on-the-fly adjustments to a production line. It’s akin to the internet having its own driver’s license — and the keys to the car. However, as with humans, the brighter the agent, the more we must consider responsibility, safety, and trust.

Related conditions from the framework:

- [Cond 4] Time Availability — Autonomy often implies continuous operation beyond human working limits.
- Cond 5] Fatigue Resistance — These systems operate under harsh or long-duration conditions without degradation.
- [Cond 7] Data Mining Capacity — Real-time environment scanning and decision-making depend on constant data extraction and interpretation.
- [Cond 9] Pattern Correlation Abilities — Crucial for real-time prediction, navigation, and anomaly detection.
- [Cond 10] Speed of Computation — Necessary for fast decision loops in high-stakes contexts.
- [Cond 15] Ethical and Regulatory Framing — As these systems gain autonomy, accountability and compliance become critical.

Typical application example:

Autonomous vehicles exemplify a significant application: AI agents must continuously process data from cameras, LIDAR, radar, and GPS, perform localization and trajectory planning, and make ethically charged decisions under uncertainty. These tasks demand a tight integration of structural and functional intelligence performance conditions.

Future implications and open challenges:

- Balancing autonomy with human oversight and ensuring fail-safe behaviors.
- Designing multi-agent coordination protocols in swarms or fleets.
- Embedding moral reasoning and explainability in high-risk systems (e.g., medical robotics, defense AI).
- Establishing legal responsibility and liability frameworks for decisions made by non-human agents.

4.8. Responsible AI

Responsible AI is what happens when we stop asking “Can we do this?” and start seriously asking, “Should we?” It’s not just about building smart systems— it’s about creating systems that behave. In other words, AI needs to work effectively and align with human values. That means being fair, accountable,

transparent, and respectful of privacy and basic rights—especially in sensitive areas like healthcare, education, finance, and criminal justice, where mistakes can carry serious consequences [33] [34] [35].

At its core, Responsible AI is guided by foundational principles often captured under the SHIFT framework: Sustainability, Human-centeredness, Inclusiveness, Fairness, and Transparency [33]. These values ensure that as AI technologies evolve, they do so with people—and the planet—in mind. It's not just about accuracy or performance; it's about ensuring systems don't accidentally (or worse, intentionally) discriminate, violate privacy, or cause harm.

Making Responsible AI a reality isn't just about good intentions. It requires clear governance structures, transparent design processes, and continuous evaluation [34]-[36]. Organizations must operationalize these ideals through both internal policies and external regulatory compliance. While some companies have attempted “ethics washing” in the past, many now recognize that a combination of self-regulation and legally enforced accountability is necessary for truly trustworthy systems [37] [38].

Sector-specific considerations are also key. In healthcare, for example, Responsible AI principles are critical—because algorithmic bias or opacity can literally be a matter of life or death [4]. Similarly, sustainability is gaining traction as a major concern, not just in terms of energy use but in the broader social and environmental impacts of AI systems [39] [40] [41] [42] [43].

So, if AI is the rocket ship, Responsible AI is the navigation system—and the seatbelt. It's not about slowing innovation down. It's about making sure we arrive safely and with integrity, dignity, and fairness intact.

Related conditions from the framework:

- [Cond 1] Designer Category — The values and intent of creators shape system behavior.
- [Cond 12] Continuous Learning / Evolution — Responsible AI must adapt while avoiding harmful feedback loops.
- [Cond 14] Interpretability & Transparency — Trust and governance depend on the system's ability to explain its actions.
- [Cond 15] Ethical and Regulatory Framing — Responsible AI is fundamentally about embedding external constraints and moral reasoning into system design and operation.

Typical application example:

In the healthcare sector, AI tools used for diagnostics or treatment recommendations must ensure not only high accuracy, but also fairness across demographic groups, explainability of results, and compliance with privacy laws (e.g., HIPAA, GDPR). A Responsible AI approach would include audit trails, human-in-the-loop decision-making, and bias mitigation strategies.

Future implications and open challenges:

- Designing auditable models that maintain privacy and robustness.
- Integrating ethics-by-design in development pipelines rather than post-hoc corrections.
- Harmonizing international regulatory frameworks for cross-border AI applications.
- Creating interdisciplinary teams that bring together technologists, ethicists, legal experts, and affected communities.

4.9. Quantum AI

Quantum AI is where things start getting wild — but in a good way. It blends two of the most exciting fields in tech: quantum computing and artificial intelligence. While classical bits can only be 0 or 1, quantum bits (qubits) can exist in both states at once (thanks to superposition), and they can be mysteriously linked across space (thanks to entanglement). This gives quantum systems the potential to solve problems that would take traditional computers practically forever. But Quantum AI isn't just about doing the same things faster — it's about doing some things differently altogether.

Take, for example, recent work by Shahwar et al. [44], where a hybrid ZFNet–Quantum Neural Network was developed to detect pneumonia from chest X-ray images. In their method, features extracted by the classical ZFNet architecture were passed into a parameterized quantum circuit. This circuit, running on actual quantum devices via platforms like PennyLane, transformed thousands of features into a compressed set of highly informative quantum-enhanced signals. The result? A classification accuracy of 96.5%, outperforming the CNN baseline.

This study demonstrates that quantum computing isn't just a theoretical exercise—it can already be applied in high-stakes domains, such as healthcare. Quantum models can enhance feature selection, optimize

learning dynamics, and deliver higher accuracy with fewer resources, all while working in tandem with classical deep learning models.

So, what does this mean for AI? Quantum AI represents a paradigm shift. It provides us with new ways to think about computation, learning, and pattern recognition — not just faster, but fundamentally different. As research like [44] shows, we’re not just imagining science fiction anymore. We’re building a future where intelligence operates in both classical and quantum dimensions.

Related conditions from the framework:

- [Cond 4] Time Availability — Quantum algorithms could drastically reduce training and inference times.
- [Cond 10] Speed of Computation — Quantum computing can exponentially speed up specific AI processes (e.g., search, optimization).
- [Cond 11] Parallel Processing / Multitasking — Quantum systems naturally encode and process multiple states simultaneously.
- [Cond 12] Continuous Learning / Evolution — Quantum-enhanced learning models could achieve new levels of adaptability.

We may consider a separate condition: Quantum coherence management could be framed as a novel constraint for intelligence.

Typical application example:

Quantum machine learning (QML) algorithms, such as the Variational Quantum Classifier (VQC) or Quantum Support Vector Machine (QSVM), are already being explored for use in drug discovery, material science, cryptography, and climate modeling — tasks involving vast state spaces or combinatorial complexity.

Future implications and open challenges:

- Hardware limitations: Scalability, error correction, and decoherence are major challenges.
- Algorithm development: Need for AI models natively designed for quantum architectures.
- Talent and tools: Bridging the gap between AI and quantum computing communities.
- Ethical foresight: Quantum AI could outpace current regulatory and ethical frameworks due to its sheer power and opacity.

4.10. Discussion

4.10.1. Recapitulation

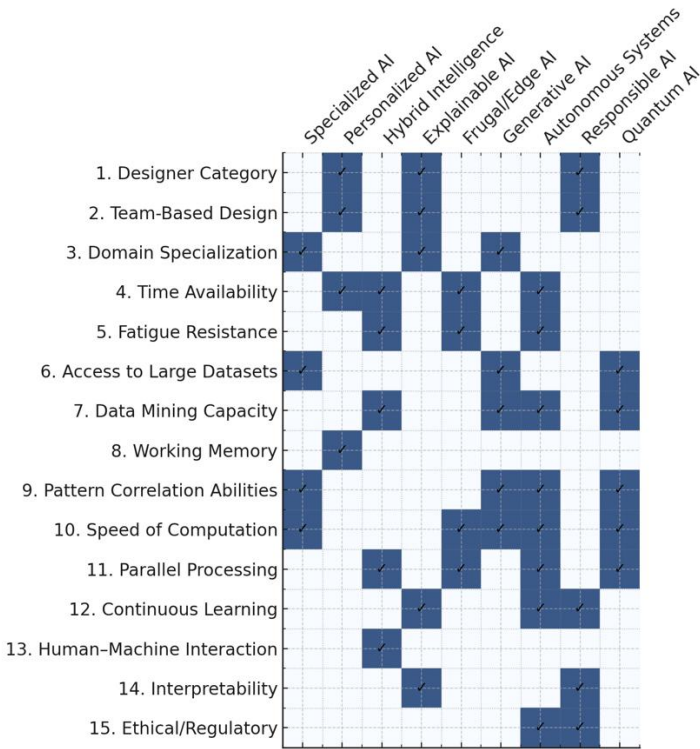


Figure 2: Mapping of 15 Intelligence conditions to 9 AI development directions

Figure 2 provides a visual representation of how everything fits together. It maps the 15 conditions of intelligence performance — divided into structural and functional categories — against the nine technological directions we’ve explored so far. Each cell shows whether a condition is critical to a given direction.

Think of it as a kind of “matchmaking chart” for intelligence: which capabilities are needed where, how they reinforce one another, and where the most significant gaps (and opportunities) might be. It’s both a diagnostic lens and a planning tool — one we hope will spark new conversations and clearer pathways forward.

4.10.2. Ethics and Regulation as Prerequisites for Intelligence

Let’s be honest: intelligence without ethics can get dangerous — fast. In humans, we’ve got a pretty robust safety net for this. Our moral compass doesn’t come from one place; layers shape it:

- Our gut instincts (innate moral sense),
- The stuff we learn growing up — family, school, culture,
- The law, with its rules, rights, and consequences,
- And society itself, which quietly (and sometimes not so quietly) nudges us to do the right thing.

Together, these form a kind of moral GPS. It’s not perfect, but it helps keep our intelligence from veering into harmful territory.

Now, AI doesn’t have feelings or free will. But it does make decisions that affect people’s lives. So it too needs a moral framework — something that functions like the human version, even if it’s not built the same way.

That means ethics and regulation aren’t just afterthoughts. They need to be baked into every step of the AI process — from the initial code to how the system behaves once it’s out in the world. Developers should follow a digital version of the Hippocratic oath: write code that helps, not harms. And those values must be reflected in how the AI behaves, particularly when it is used by humans who may not always have the best intentions.

In short: whether intelligence is made of neurons or algorithms, it only works for the good if it’s guided by values.

Two core conditions in our framework reinforce this principle:

- Condition 14: Interpretability and Transparency ensures that AI decisions can be understood and traced, enabling oversight and correction.
- Condition 15: Ethical and Regulatory Framing calls for an explicit embedding of moral values, legal constraints, and social responsibility into AI systems.

These are not ancillary features but prerequisites for trustworthy and sustainable intelligence. Just as no human is considered a responsible citizen without moral and legal alignment, no AI system should be deployed without a parallel ethical infrastructure.

Ultimately, intelligence must not be judged solely by its ability to solve problems, but by the values it respects while doing so. Ethics and regulation are not limitations—they are the essential operating system of meaningful intelligence, natural or artificial.

4.10.3. Cross-analysis: Mapping Conditions to Directions

When we look closely at how the 15 conditions of intelligence performance intersect with the primary AI development directions, some fascinating patterns emerge. It’s not just about technical progress anymore — it’s about strategic synergy.

Consider Responsible AI or Autonomous Intelligent Systems. These are not built only in fancy code; they rest on further layers: ethics, explainability, and regulatory accountability (Conditions 14, 15). Others are more like Frugal AI or Specialized AI, and they have come from different soil: structural efficiency and custom optimization.

Have you noticed how modern AI is shifting away from isolated tools and toward ecosystems? Increasingly, we’re seeing systems that adapt based on who’s using them, what they’re doing, and what’s happening around them. This is the age of hybrid, adaptive, contextual intelligence — and it’s changing the game.

But let’s not kid ourselves. Can purely technological solutions ever be enough? Without strong design safeguards and a human-in-the-loop philosophy, even our most advanced systems can produce outcomes that no one wants. That’s why academia, industry, and public governance must step up — not just to build better tech, but to shape how it impacts human life.

4.10.4. Trends: Hybrid, Adaptive, Contextual Intelligence

What's next for AI? It's not super-general or narrowly specialized — it's contextual and hybrid. These are systems that evolve with the situation, collaborate with humans, and specialize in ways that fit real environments.

Think about it: Wouldn't it be ideal to have AI that not only learns fast, but knows when and how to act in the moment? This is where we're headed — toward human-AI collaboration platforms, cognition at the edge, and even quantum-enhanced decision-making.

And here's another question: Is building more intelligent machines only about better engineering? Or does it also require perspectives from neuroscience, philosophy, law, and behavioral science? If we want truly effective intelligence, it must be interdisciplinary. That's the only way to shape systems that understand both data and human context.

4.10.5. Limitations of Purely Technological Approaches

Let's face it: technology doesn't of itself provide us with meaning, responsibilities or goals. It's a powerful engine — but unmoored from a steering wheel, it can go off the rails.

Have you ever wondered why, in AI, some of our most significant problems — such as bias, black-box decisions, or value misalignment — are not simply technical bugs? That's for the simple reason that they are not engineering problems. They're conceptual, moral, and philosophical ones.

If we allow ourselves to deal with these by proxy, he argued, are we moving artificial intelligence forward — or just enabling its dangers?

4.10.6. Strategic Role of Academia, Industry, and Governance

So, who is responsible for getting this right?

The answer: everyone.

Academia must ask the big questions, design ethical frameworks, and train the next generation (why not start with "Intelligence Awareness" curricula?). The industry needs to take those insights and build scalable, responsible technology. And governments? They're the ones who can create fair and enforceable rules — rules that protect innovation and people at the same time.

But here's a challenge: Can these three players align on shared values? If they can — values like transparency, accountability, and inclusiveness — and translate them into code, policy, and protocol, we're on the path to something much more meaningful.

4.10.7. The Rise of Quantum AI

Now let's talk about the new kid on the block: Quantum AI.

It's not just "faster AI." It's an entirely new way of thinking about intelligence. With quantum bits (qubits) operating in superposition, these systems can explore massive decision spaces in parallel. That means we're unlocking serious potential for optimization, speed (Condition 10), and parallelism (Condition 11).

But here's the twist: Are we ready for the ethical questions quantum AI brings with it? If we can't trace how decisions are made, what does that mean for accountability (Condition 15)? We're entering uncharted waters.

Quantum AI may even alter the very standards we use to measure intelligence — whether human or artificial. That's why we argue for a unified framework now — to stay ahead of disruption before it arrives [45]- [49]

5. Conclusion

And as the pace of innovation accelerates, the nature of intelligence — human or artificial — will increasingly be defined not by a fixed set of competencies, and certainly not by any one nation's talents, but by a dynamic ability to learn and adapt. It's no longer about whether machines become more like humans or whether humans can outperform algorithms. The real question is: what can each system do, within its constraints, to enable outcomes that are safe, adaptive, and meaningful?

To explore this, we developed a conceptual framework of 15 conditions of intelligence performance and nine developmental directions in AI. This 15×9 grid provides a systematic approach to assessing the limitations and potential of natural and synthetic agents. It emphasizes the importance of context sensitivity, ethical framing, interpretability, and participatory design — key process elements that are just as important as the intelligent results of intelligent processes.

Above all, this work invites a shift in mindset. Rather than pursuing AGI in the image of human cognition, we should focus on building intelligences that extend and complement human potential. Doing so requires not

only engineering prowess but also interdisciplinary thinking — drawing from philosophy, neuroscience, ethics, computer science, and law to align the growing power of intelligent systems with societal purpose.

REFERENCES

- [1] J. E. Korteling, G. C. van de Boer-Visschedijk, R. A. M. Blankendaal, R. C. Boonekamp, and A. R. Eikelboom, "Human- versus Artificial Intelligence," *Frontiers in Artificial Intelligence*, vol. 4, 2021. <https://doi.org/10.3389/frai.2021.622364>.
- [2] M. H. Bornstein, "The Ecological Approach to Visual Perception," 1980, *JSTOR*.
- [3] C. A. Hartley, "How do Natural Environments Shape Adaptive Cognition Across The Lifespan?," *The Journal of Aesthetics and Art Criticism*, vol. 26, no. 12, pp. 1029–1030, 2022. <https://doi.org/10.2307/429816>
- [4] A. Coutrot *et al.*, "Entropy of City Street Networks Linked to Future Spatial Navigation Ability," *Nature*, vol. 604, no. 7904, pp. 104–110, 2022. <https://doi.org/10.1038/s41586-022-04486-7>
- [5] C. S. Green and D. Bavelier, "Action video Game Modifies Visual Selective Attention," *Nature*, vol. 423, no. 6939, pp. 534–537, 2003. <https://doi.org/10.1038/nature01647>
- [6] D. W. Massaro and D. Friedman, "The Adaptive Character of Thought," 1991, *JSTOR* <https://doi.org/10.2307/1423253>.
- [7] S. A. Nastase, A. Goldstein, and U. Hasson, "Keep it Real: Rethinking The Primacy of Experimental Control in Cognitive Neuroscience," *Neuroimage*, vol. 222, p. 117254, 2020. <https://doi.org/10.1016/j.neuroimage.2020.117254>
- [8] K. E. Adolph, J. E. Hoch, and W. G. Cole, "Development (of walking): 15 Suggestions," *Trends in Cognitive Sciences*, vol. 22, no. 8, pp. 699–711, 2018. <https://doi.org/10.1016/j.tics.2018.05.010>
- [9] P. Slijepcevic, "Natural Intelligence and Anthropropic Reasoning," *Biosemitotics*, vol. 13, no. 2, pp. 285–307, 2020. <https://doi.org/10.1007/s12304-020-09388-7>
- [10] R. J. Sternberg, "A Theory of Adaptive Intelligence and Its Relation to General Intelligence," *Journal of Intelligence*, vol. 7, no. 4, p. 23, 2019. <https://doi.org/10.3390/jintelligence7040023>
- [11] L. H. Howard, C. Carrazza, and A. L. Woodward, "Neighborhood Linguistic Diversity Predicts Infants' Social Learning," *Cognition*, vol. 133, no. 2, pp. 474–479, 2014. <https://doi.org/10.1016/j.cognition.2014.08.002>
- [12] A. Saxe, S. Nelli, and C. Summerfield, "If Deep Learning is The Answer, What is The Question?," *nature reviews neuroscience*, vol. 22, no. 1, pp. 55–67, 2021. <https://doi.org/10.1038/s41583-020-00395-8>
- [13] C. Angermueller, T. Pärnamaa, L. Parts, and O. Stegle, "Deep learning for computational biology," *Molecular Systems Biology*, vol. 12, no. 7, p. 878, 2016. <https://doi.org/10.15252/msb.20156651>
- [14] T. Macpherson *et al.*, "Natural and Artificial Intelligence: A brief introduction to the interplay between AI and neuroscience research," *Neural Networks*, vol. 144, pp. 603–613, 2021. <https://doi.org/10.1016/j.neunet.2021.09.018>.
- [15] Y. Wang, "Innovative Applications of Artificial Intelligence in Specific Domains," *Innovation in Science and Technology*, vol. 3, no. 5, pp. 40–46, 2024. <https://doi.org/10.56397/IST.2024.09.05>
- [16] Y. Chun, J. Hur, and J. Hwang, "AI Technology Specialization and National Competitiveness," *PLoS One*, vol. 19, no. 4, p. e0301091, 2024. <https://doi.org/10.1371/journal.pone.0301091>
- [17] D. R. Serrano *et al.*, "Artificial Intelligence (AI) Applications In Drug Discovery and Drug Delivery: Revolutionizing Personalized Medicine," *Pharmaceutics*, vol. 16, no. 10, p. 1328, 2024. <https://doi.org/10.3390/pharmaceutics16101328>
- [18] P. Rajpurkar, E. Chen, O. Banerjee, and E. J. Topol, "AI in Health and Medicine," *Nature Medicine*, vol. 28, no. 1, pp. 31–38, 2022. <https://doi.org/10.1038/s41591-021-01614-0>
- [19] S. Hopcan, E. Polat, M. E. Ozturk, and L. Ozturk, "Artificial Intelligence in Special Education: A Systematic Review," *Interact. Learn. Environ.*, vol. 31, no. 10, pp. 7335–7353, 2023. <https://doi.org/10.1080/10494820.2022.2067186>
- [20] M. Khalifa, M. Albadawy, and U. Iqbal, "Advancing clinical decision support: The role of artificial intelligence across six domains," 2024, *Elsevier*. <https://doi.org/10.1016/j.cmpbup.2024.100142>
- [21] S. G. Essa, T. Celik, and N. E. Human-Hendricks, "Personalized Adaptive Learning Technologies Based on Machine Learning Techniques to Identify Learning Styles: A Systematic Literature Review," *IEEE Access*, vol. 11, pp. 48392–48409, 2023. <https://doi.org/10.1109/ACCESS.2023.3276439>
- [22] M. H. Jarrahi, C. Lutz, and G. Newlands, "Artificial Intelligence, Human Intelligence and Hybrid Intelligence Based On Mutual Augmentation," *Big Data and Society*, vol. 9, no. 2, p. 20539517221142824, 2022. <https://doi.org/10.1177/20539517221142824>
- [23] R. Filieri, Z. Lin, Y. Li, X. Lu, and X. Yang, "Customer Emotions in Service Robot Encounters: A Hybrid Machine-Human Intelligence Approach," *Journal of Service Research*, vol. 25, no. 4, pp. 614–629, 2022. <https://doi.org/10.1177/10946705221103937>
- [24] A. Raza *et al.*, "A Hybrid Deep Learning-Based Approach for Brain Tumor Classification," *Electronics*, vol. 11, no. 7, p. 1146, 2022. <https://doi.org/10.3390/electronics11071146>
- [25] S. M. Saqib, O. Muhammad, T. Mazhar, M. Iqbal, S. Guizani, and H. Hamam, "Medicine Image Classification Using Deep Learning: Highlighting the Mednet-Mobil Hybrid Model," *Discover Computing*, vol. 28, no. 1, p. 103, 2025. <https://doi.org/10.1007/s10791-025-09619-w>
- [26] A. N. Raikov and M. Pirani, "Human-machine duality: What's Next in Cognitive Aspects of Artificial Intelligence?," *IEEE Access*, vol. 10, pp. 56296–56315, 2022. <https://doi.org/10.1109/ACCESS.2022.3177657>
- [27] Y. Y. Ghadi *et al.*, "Explainable AI Analysis for Smog Rating Prediction," *Scientific Reports*, vol. 15, no. 1, p. 8070, 2025. <https://doi.org/10.1038/s41598-025-92788-x>
- [28] M. A. Khan *et al.*, "Decoding'Eligibility Unknown': transparent classification and feature-based reclassification in

- CAFV analysis,” *International Journal of Electrical Power & Energy Systems*, vol. 170, p. 110894, 2025. <https://doi.org/10.1016/j.ijepes.2025.110894>
- [29] Z. Ahmed, K. Munir, M. U. Tanveer, S. R. Hassan, A. U. Rehman, and H. Hamam, “Deep Pseudogene Categorization and Genome-Wide Transcription Prediction Using GANP-Based Feature Selection and TabNet Interpretability,” *IEEE Access*, 2025. <https://doi.org/10.1109/ACCESS.2025.3585601>
- [30] R. Singh and S. S. Gill, “Edge AI: a survey,” *Internet of Things and Cyber-Physical Systems*, vol. 3, pp. 71–92, 2023. <https://doi.org/10.1016/j.iotcps.2023.02.004>
- [31] M. Iyobhebhe *et al.*, “A review on Energy Consumption Model on Hierarchical clustering techniques for IoT-based multilevel heterogeneous WSNs using Energy Aware Node Selection,” *Vokasi Unesa Bulletin of Engineering Technology and Applied Science*, vol. 2, no. 2, pp. 100–111, 2025. <https://doi.org/10.26740/vubeta.v2i2.34882>
- [32] T. Mazhar *et al.*, “Generative AI, IoT, and Blockchain in Healthcare: Application, Issues, and Solutions,” *Discover Internet of Things*, vol. 5, no. 1, p. 5, 2025. <https://doi.org/10.1007/s43926-025-00095-8>
- [33] H. Siala and Y. Wang, “SHIFtIng Artificial Intelligence to be Responsible in Healthcare: A Systematic Review,” *Social Science & Medicine*, vol. 296, p. 114782, 2022. <https://doi.org/10.1016/j.socscimed.2022.114782>
- [34] N. Díaz-Rodríguez, J. Del Ser, M. Coeckelbergh, M. L. De Prado, E. Herrera-Viedma, and F. Herrera, “Connecting the Dots in Trustworthy Artificial Intelligence: from AI Principles, Ethics, and Key Requirements to Responsible AI Systems and Regulation,” *Information Fusion*, vol. 99, p. 101896, 2023. <https://doi.org/10.1016/j.inffus.2023.101896>
- [35] E. Papagiannidis, P. Mikalef, and K. Conboy, “Responsible Artificial Intelligence Governance: A Review and Research Framework,” *The Journal of Strategic Information Systems*, vol. 34, no. 2, p. 101885, 2025. <https://doi.org/10.1016/j.jsis.2024.101885>
- [36] P. Mikalef, K. Conboy, J. E. Lundström, and A. Popović, “Thinking Responsibly about Responsible AI and ‘the dark side’ of AI,” 2022, *Taylor & Francis*. <https://doi.org/10.1080/0960085X.2022.2026621>
- [37] R. Baeza-Yates, “Introduction to Responsible AI,” *Proceedings of the 17th ACM International Conference on Web Search and Data Mining*, 2024, pp. 1114–1117. <https://doi.org/10.1145/3616855.3636455>
- [38] H. Herrmann, “What’s next for responsible artificial intelligence: a way forward through responsible innovation,” *Heliyon*, vol. 9, no. 3, 2023. <https://doi.org/10.1016/j.heliyon.2023.e14379>
- [39] Y.-L. Chang and J. Ke, “Socially Responsible Artificial Intelligence Empowered People Analytics: A Novel Framework Towards Sustainability,” *Human Resource Development Review*, vol. 23, no. 1, pp. 88–120, 2024. <https://doi.org/10.1177/15344843231200930>
- [40] J. B. Lyons, K. Hobbs, S. Rogers, and S. H. Clouse, “Responsible (use of) AI,” *Front. Neuroergonomics*, vol. 4, p. 1201777, 2023. <https://doi.org/10.3389/fnrgo.2023.1201777>
- [41] M. I. Merhi, “An Assessment of the Barriers Impacting Responsible Artificial Intelligence,” *Information Systems Frontiers*, vol. 25, no. 3, pp. 1147–1160, 2023. <https://doi.org/10.1007/s10796-022-10276-3>
- [42] A. K. Shukla, V. Terziyan, and T. Tiihonen, “AI as a user of AI: Towards Responsible Autonomy,” *Heliyon*, vol. 10, no. 11, 2024. <https://doi.org/10.1016/j.heliyon.2024.e31397>
- [43] B. C. Stahl, “Embedding Responsibility in Intelligent Systems: From AI Ethics to Responsible AI Ecosystems,” *Scientific Reports*, vol. 13, no. 1, p. 7586, 2023. <https://doi.org/10.1038/s41598-023-34622-w>
- [44] T. Shahwar, F. Mallek, A. U. Rehman, M. T. Sadiq, and H. Hamam, “Classification of Pneumonia Via a Hybrid Zfnct-Quantum Neural Network using a Chest X-Ray Dataset,” *Current Medical Imaging*, vol. 20, no. 1, p. e15734056317489, 2024. <https://doi.org/10.2174/0115734056317489240808094924>
- [45] S. Khan *et al.*, “Optimizing load Demand Forecasting in Educational Buildings using Quantum-Inspired Particle Swarm Optimization (QPSO) with Recurrent Neural Networks (RNNs): A Seasonal Approach,” *Scientific Reports*, vol. 15, no. 1, p. 19349, 2025. <https://doi.org/10.1038/s41598-025-04301-z>
- [46] T. Shahwar *et al.*, “Automated detection of Alzheimer’s via hybrid classical quantum neural networks,” *Electronics*, vol. 11, no. 5, p. 721, 2022. [https://doi.org/10.47363/JAICC/2025\(4\)125](https://doi.org/10.47363/JAICC/2025(4)125)
- [47] H. Hamam, “The Convergence of AI, Cloud and Quantum Computing: Preparing for the Next Leap,” *Journal of Artificial Intelligence & Cloud Computing*, vol. 4, no. 2, pp. 1–3, 2025. [https://doi.org/10.47363/JAICC/2025\(4\)125](https://doi.org/10.47363/JAICC/2025(4)125)
- [48] M. M. Saeed *et al.*, “A comprehensive survey on 6G-security: Physical connection and service layers,” *Discover Internet of Things*, vol. 5, no. 1, p. 28, 2025. <https://doi.org/10.1007/s43926-025-00123-7>
- [49] J. Ktari, T. frikha, M. Hamdi, N. Affes, and H. Hamam, “Enhancing Blockchain Security and Efficiency through FPGA-based Consensus Mechanisms and Post-quantum Cryptography,” *Recent Advances in Electrical & Electronic Engineering*, vol. 18, no. 7, pp. 946–958, 2025. <https://doi.org/10.2174/0123520965288815240424054237>

BIOGRAPHY OF THE AUTHOR

Habib Hamam     obtained the B.Eng. and M.Sc. degrees in information processing from the Technical University of Munich, Germany 1988 and 1992, and the PhD degree in Physics and applications in telecommunications from Université de Rennes I conjointly with France Telecom Graduate School, France 1995. He also obtained a postdoctoral diploma, “Accreditation to Supervise Research in Signal Processing and Telecommunications”, from Université de Rennes I in 2004. He was a Canada Research Chair holder in “Optics in Information and Communication Technologies”, the most prestigious research position in Canada – which he held for a decade (2006-2016). The title is awarded by the Head of the Government of Canada after a selection by an international scientific jury in the related field. He is currently a full Professor in the Department of Electrical Engineering at Université de Moncton. He is OSA senior member, IEEE senior member and a registered professional engineer in New-Brunswick. He obtained several pedagogical and scientific awards. He is among others editor in chief and founder of CIT-Review, academic editor in Applied Sciences and associate editor of the IEEE Canadian Review. He also served as Guest editor in several journals. His research interests are in optical telecommunications, Wireless Communications, diffraction, fiber components, signal and image processing, IoT, data protection, AI and Big Data.