

Envisioning Possible Futures for AI Research

What potential paradigm shifts could produce the next revolution in AI?

Authors

David Jensen (University of Massachusetts, Amherst), David Danks (University of California, San Diego), Sebastian Elbaum (University of Virginia), William Regli (University of Maryland), Matthew Turk (University of California, Santa Barbara), Adam Wierman (California Institute of Technology), Holly Yanco (University of Massachusetts, Lowell), Mary Lou Maher (Computing Research Association), Haley Griffin (Computing Research Association)

Introduction

The current wave of artificial intelligence (AI) innovation is being driven by a fifteen-year-old paradigm shift in AI research. That shift was just the most recent in a series of such shifts since AI got its modern start in 1956. What could be the *next* such shift in AI research? Below, we describe why this question is important to researchers, practitioners, and policymakers, and we provide six examples of paradigms that may help define the next generation of AI research.

Why Envision AI Research Futures?

Envisioning the future is challenging. It is difficult for researchers to see beyond the current scientific paradigm, for technologists to see beyond the latest technological developments, and for policymakers to see beyond the issues those new technologies raise. This is particularly true when powerful new technologies sweep rapidly into public awareness, as AI technologies have recently done.

However, it is useful to recall that the current scientific and technical moment for AI — powered by deep neural networks, large language models, and other foundation models — is just the latest in a series of such moments that the field has experienced in its comparatively brief history. Prior research paradigms for AI include symbolic processing, knowledge-based systems, and statistical machine learning. Each paradigm was hailed as ushering in a new age of AI, each produced a series of transformative applications, and each was eventually superseded by one or more new paradigms that built on those previous insights.

This suggests an obvious question: *What's next for AI research? That is, what comes after the current age of deep neural networks and foundation models?*

Answers to this question are unlikely to emerge easily because the transience of the current paradigm can be difficult to imagine. When dominant, a paradigm appears to be the end point, the apotheosis of a long line of precursors that merely paved the way for the current (and final) paradigm. One reason for this is that early partisans who successfully bring a paradigm into dominance within a field naturally focus attention on the capabilities of the new paradigm, rather than on its limitations. Typically, those limitations only become visible later. A second reason is that the current paradigm succeeds because it provides new capabilities for research and development, and researchers are often busy exploiting those capabilities rather than asking what comes next. A third reason is that successors to a dominant paradigm can be extremely challenging to identify. Such successors are often only obvious in retrospect, after the limitations of the prior paradigm (and the benefits of a successor paradigm) become evident.

That said, now may be a particularly good time to consider possible futures for AI research. The current paradigm of foundation models based on deep neural networks has been dominant for less than 15 years, but some of its limitations have already started to become apparent.¹ Some researchers, both inside and outside of the current paradigm, have begun to suggest that this paradigm alone won't be sufficient to reach many of the desired goals of AI.

Identifying and describing concrete visions of possible futures for AI research has substantial benefits. First, such visioning can demonstrate the value of broad-based research programs that include efforts to both exploit the current paradigm and to explore the value of alternative paradigms. Such a balanced approach has been responsible for all prior paradigm shifts in AI research. In particular, federal research funds have supported exploration of alternative paradigms long before it was obvious that one of them would become dominant and lead to major new technological capabilities. Second, concrete visions of alternatives to the current paradigm can help researchers themselves focus on more revolutionary advances, rather than only incremental advances that exploit the current paradigm. Finally, concrete descriptions of possible futures for AI research can help correct a potential misconception among the public and some policy makers that “AI is done, except for the scaling.”²

The CCC Task Force on AI Research Futures

The authors of this report are members of the Computing Community Consortium (CCC) Task Force on AI Research Futures. This Task Force was formed to identify, describe, and disseminate possibilities for “what comes next” — a set of potential future paradigms for AI

¹For example, researchers have identified significant limitations of LLMs in classic AI reasoning tasks involving planning (e.g., Valmeekam et al. 2025) and causal reasoning (e.g., Jin et al. 2023).

²Many of the technical performance metrics of foundation models have exhibited power-law like behavior. This has led to more general statements about the expected overall intelligence of such models in the future. For example, Sam Altman of OpenAI has said that “[t]he intelligence of an AI model roughly equals the log of the resources used to train and run it.... It appears that you can spend arbitrary amounts of money and get continuous and predictable gains; the scaling laws that predict this are accurate over many orders of magnitude” (Altman 2025).

research. To gather input, the Task Force convened sets of AI experts to propose, describe, refine, and discuss potential future directions for AI research in roundtable discussions. Initially, the roundtable discussions were held with Fellows of the Association for the Advancement of Artificial Intelligence (AAAI) during and shortly after the AAAI Conference on Artificial Intelligence in February and March 2025. The Task Force subsequently held additional roundtables with AI researchers throughout March 2025. The ideas expressed in this report, while influenced by the roundtable discussions, are those of the authors, and not the individuals with whom we spoke.

This CCC effort complements another effort — the AAAI 2025 Presidential Panel on the Future of AI Research — which released a report on their efforts in March 2025.^{3,4} That report aimed to “define the current trends and the research challenges still ahead of [the AI research community] to make AI more capable and reliable...”. The methodology for gathering input for the AAAI 2025 Presidential Report included the development of a survey on potential futures and gathering data from members of AAAI.

What are “AI Research Futures”?

Both at the outset of the Task Force’s work and through the roundtable discussions, we continually refined what we meant by a “future” for AI research. We ultimately identified four criteria that we used to elicit and refine the specific futures presented in this document.

First, research futures should be relatively *unified and bounded paradigms of research*. That is, some research projects should be consistent with a given future, while others should not. Second, research futures should be *possible, but not necessarily likely*. We are not trying to predict the future, but instead identify ways that AI research might progress, depending on various factors. Third, research futures should *describe both ends and means*. Many discussions about the future of AI research identify desirable “ends” for AI research (i.e., new capabilities that would be useful in future AI technologies). However, it is far less frequent (and more challenging) to identify the “means” for producing those ends (i.e., research approaches that might credibly produce those new capabilities). Finally, research futures should be *motivated by specific technical issues*. Examples include focusing on addressing one or more shortcomings of current AI systems (e.g., embodied AI) or building on a recently emerging technical possibility (e.g., quantum AI).

Prospective AI Research Futures

Below are six AI research futures that meet the criteria outlined above: neuro-symbolic AI, neuromorphic AI, embodied AI, multi-agent AI, human-centered AI, and quantum AI. The research futures described below are neither mutually exclusive nor comprehensive. However, they provide concrete examples of approaches that might take AI research in important new directions.

³Rossi et al. (2025).

⁴During one of the Task Force’s roundtables, the AAAI President, Francesca Rossi, noted that the work described in AAAI’s report is complementary to the work of the CCC Task Force.

Neuro-Symbolic AI

Neuro-symbolic AI merges deep neural networks, which excel at learning patterns from data, with symbolic AI, which enables reasoning based on logic and prior knowledge.⁵ This integration aims to overcome several limitations of deep neural networks alone. It could significantly enhance reliability by incorporating domain knowledge and leveraging available input structures. It could improve accountability, particularly where the opacity of neural networks limits their utility. Finally, neuro-symbolic AI could achieve greater generalization by combining logic and concepts in diverse ways, enabling systems to move beyond data distributions explicitly represented in training data.

However, the neural and symbolic paradigms are very different. While each paradigm has distinct advantages, their differences hinder seamless integration. Neural networks operate on sub-symbolic, distributed representations, where knowledge is encoded across a vast number of interconnected neurons through learned weights and biases. Their distributed nature makes it difficult to extract and manipulate explicit, human-readable symbols. Bridging this gap requires complex mapping mechanisms that attempt to translate between continuous, high-dimensional neural activations and discrete, interpretable symbols, and vice versa. Furthermore, the divergent execution modes of neural and symbolic methods pose significant challenges for creating a unified system. Neural networks excel at pattern recognition, generalization from data, and handling noisy or incomplete information. Their performance depends heavily on the quantity and quality of training data. In contrast, symbolic AI excels at tasks requiring explainability, formal verification, and adherence to predefined reasoning methods. It can provide clear, step-by-step derivations for its conclusions. A key area of research is developing hybrid architectures that facilitate dynamic switching or synergistic cooperation between these inductive and deductive approaches, demanding innovative solutions for control flow, information exchange, and conflict resolution between their respective outputs.

Research in neuro-symbolic AI focuses on identifying integration strategies for neural and symbolic modes. There are examples of neuro-symbolic approaches that provide foundations for this research, but researchers have not identified a universal approach or an understanding of which approaches are likely to lead to a generalized neuro-symbolic model. Current efforts, for example, aim to improve the accuracy of LLMs by using domain knowledge and logic to prune hallucinations or to be more capable by invoking symbolic tools. The longer-term goal is to build richer synergies that facilitate the novel application of known concepts and abstractions to understand and solve new problems. To uncover these synergies, researchers are constructing richer datasets that combine neural inputs with symbolic targets, new architectures that are both differentiable and symbolic to make the entire hybrid system trainable with gradient-based optimization, and new frameworks and tools to lower the adoption barrier for integrating neural and symbolic elements.

⁵For more information, see recent surveys on the history and scope of neuro-symbolic approaches, for example Buyan, et al. (2024). There are arguments for the potential of recent advances in LLMs to spark new interest in neuro-symbolic AI to address the Grounding Problem, see Maher, et al. (2024).

Neuromorphic AI

Neuromorphic computing uses computational hardware whose physical structure mimics some key aspects of the neural tissue of humans and other animals. Specifically, neuromorphic computing devices use hardware elements — such as transistors, memristors, spintronic memories, and threshold switches — to emulate the operation of neurons. Over the past several decades, a variety of such neural hardware has been designed, constructed, and evaluated, producing a diverse array of systems all grouped under the general category of neuromorphic computing.⁶

The current dominant paradigm of AI research focuses on deep neural networks, and thus can be said to be “brain-inspired”. However, deep neural networks can be extremely power-intensive, particularly when used to implement systems such as large language models. In contrast, neuromorphic approaches attempt to directly mimic the structure and behavior of neurons in hardware. These approaches typically excel in applications that aim to minimize size, weight, and power consumption. In addition, neuromorphic hardware tightly couples memory and computation, thus minimizing latency.

Neuromorphic approaches to AI have benefited from at least two recent trends in computer hardware and applications. First, developments in neuromorphic hardware have accelerated in recent years, with an increasing number of chips and larger systems becoming commercially available. Second, practitioners have shown increasing interest in intelligent mobile and embedded devices, where size, weight, and power considerations are particularly important. Such applications are a relatively poor match for deep neural networks and a far better match for neuromorphic systems.

Despite these advantages, research in neuromorphic computing continues to face significant challenges. First, computational results from neuromorphic hardware can vary from run-to-run and machine-to-machine, and this can decrease the accuracy and reliability of any given computation. Second, despite a long history, research in neuromorphic computing is still fairly sparse, thus limiting the availability of standard methods, systems, and benchmarks. Finally, research in neuromorphic computing draws on concepts and principles from multiple fields — including neuroscience, computer science, electrical engineering, mathematics, and physics. This multidisciplinary character can make the field less attractive to early career researchers and pose a steep learning curve for those who do choose to pursue research in this area.

Embodied AI

Embodied AI represents a paradigm that moves beyond purely computational intelligence to systems that possess a physical presence in the world. Unlike current dominant paradigms that focus on disembodied AI, embodied AI agents perceive, act, and learn within physical environments or (for computational scaling with a loss of some realism) high-fidelity simulations.⁷ Direct interaction with the physical world provides embodied AI with unique access to rich, multimodal sensorimotor data — including proprioception, touch, vision, and

⁶For a more detailed introduction to neuromorphic AI, see: Kudithipudi et al. (2025); and Muir & Sheik (2025).

⁷For more discussion of embodied AI, see: Brooks (1991); Pfeifer & Fumiya (2004); and Deitke et al. (2022).

sound — that is intrinsically grounded in reality. This access to "physical data" is fundamentally different from the symbolic and statistical patterns that large language models infer from vast training sets of human-generated information, allowing embodied AI the potential to develop a deeper understanding of causality, spatial relationships, and object properties through direct experience, mirroring the way biological intelligence develops.

The ability of embodied AI to gather and interpret data directly from the physical world, in a way that current deep learning approaches cannot, could lead to fundamental breakthroughs. For instance, the implicit understanding of physics, object permanence, and human-robot interaction that arises from direct physical engagement could enable a more grounded and robust form of intelligence that contrasts sharply with the often brittle and error-prone "understanding" of the physical world exhibited by LLMs. Building on experience rather than inference allows the accumulation of rich, real-world data and gives embodied AI the potential to achieve levels of robust common sense and dexterous manipulation that are critical for deployment in unstructured human environments.

One primary hurdle is the sheer complexity of real-world environments, which are dynamic, unpredictable, and necessitate robust perception and action capabilities under varying conditions. While LLMs enable some types of common-sense reasoning based on abstract linguistic patterns, embodied AI must contend with the physical world's tangible uncertainties and continuous data streams. The creation of hardware that is both sophisticated enough for complex tasks and resilient enough for continuous physical interaction remains a substantial engineering challenge, encompassing issues of power, durability, and miniaturization. Furthermore, developing effective learning algorithms for embodied agents requires addressing problems of sample inefficiency, safe exploration in physical spaces, and the transfer of learned skills from simulation to reality (the "sim-to-real" gap), all of which are uniquely exacerbated by the need to interact with physical dynamics.

Multi-Agent AI

Multi-agent AI moves beyond today's single, monolithic AI systems to a collaborative ecosystem of specialized AI agents. In this paradigm, multiple independent AI agents, each possessing distinct capabilities, knowledge, and objectives, interact and coordinate to achieve complex overarching goals. These agents can communicate, negotiate, and even learn from each other within a shared environment. This distributed intelligence allows for emergent behaviors and more robust, adaptive solutions to problems that would be intractable for a monolithic system. Examples range from autonomous driving systems where agents handle perception, planning, and control, to complex supply chain management where agents optimize logistics, inventory, and demand forecasting.⁸

Multi-agent AI offers a clear contrast to LLM-based approaches, which, while powerful in natural language processing and generation, typically operate as single, centralized entities. Traditional LLMs excel at tasks that require extensive knowledge recall or creative text

⁸For a more detailed introduction and discussion of multi-agent AI, see: Shoham & Leyton-Brown (2008), Zhang et al. (2021), Hadfield et al. (2025).

generation based on a single prompt. However, they are less effective in tasks requiring multi-step reasoning, dynamic problem-solving, diverse forms of expertise, or real-time interaction with complex environments. Multi-agent AI, conversely, leverages specialization and parallel processing. Instead of one large model attempting to do everything, a multi-agent system can deploy a "team" of agents, each fine-tuned for a specific role (e.g., a "planner" agent, a "coder" agent, a "critic" agent). This division of labor leads to improved accuracy, reduced hallucination, enhanced scalability, and the ability to handle more dynamic and complex scenarios by allowing agents to validate each other's work and adapt to changing conditions.

Realizing the potential of multi-agent AI requires the development of robust communication protocols and coordination mechanisms that enable effective interaction and information exchange between diverse agents. This includes establishing standard ontologies for shared understanding, designing effective message passing systems, and implementing frameworks to manage task allocation, conflict resolution, and collaborative decision-making among agents. Further, the ability for agents to learn and adapt over time, both individually and collectively, is crucial. This necessitates advancements in reinforcement learning, multi-objective optimization, and mechanisms for knowledge sharing across the agent network.

Human-Centered AI

Human-centered AI starts with the principle that the complementary nature of human and machine cognition implies that we can have significantly more powerful systems through integration of both. That is, this paradigm emphasizes collaboration and coordination between humans and AI systems to augment human capabilities, rather than attempting to have autonomous AI systems that function without human interaction to replace human intelligence.⁹ While human-AI teaming is already a major area of research focus, much of that work has emphasized user experience, interaction design, and human factors. However, effective human-centered AI may require AI systems that have "social intelligence," in the sense of being able to detect and respond to complex social cues, infer complex cognitive and affective states of human teammates, interact in temporally extended and non-manipulative ways, and reason about complicated and context-sensitive human social networks.

The goal of this approach to AI research, design, and development would be systems that contribute to the social interactions and group performance that underlie so many human successes. Such technology would stand in stark contrast with the current efforts to develop AI systems that can function autonomously, often with the goal of replacing human cognition.

Social cognition and interactions have been widely studied in the cognitive and neural sciences, but these insights have only rarely made their way into AI research. For example, multiple studies have implemented simple versions of so-called "theory of mind" inference in AI systems, but these inferential systems have not regularly been incorporated into AI systems that interact with humans in the real world. The area of social robotics has arguably done the most AI research in this space, but those insights have not been incorporated into the field of AI more generally.

⁹For more discussion of human-centered AI, see: Ozmen Garibay (2023); Riedl (2019); and Shneiderman (2022).

A central challenge is to redesign AI research so that social interactions and understanding are part of the foundation of the AI system, rather than being added later through (perhaps sophisticated) interfaces. Research comparing social and asocial animal species has consistently found very large differences in cognitive capabilities, which suggests that AI systems designed “from the ground up” to be social would potentially be quite different from our current systems. However, such a redesign would require not only research support, but also a significantly more multidisciplinary approach that incorporates insights from a range of social science and animal cognition disciplines.

Quantum AI

Quantum AI is a still somewhat speculative paradigm of AI research based on exploiting emerging principles and methods of quantum computing. Quantum computing has come to include a wide variety of theoretical ideas and, more recently, practical technologies for harnessing physics in unique ways to execute computational processes. Quantum computing exploits the properties of physics found in the various states of matter to “compute” in radically different ways. A central quantum computing concept is that of a quantum bit or “qubit”. In contrast to a traditional zero-or-one “bit”, a qubit uses the superposition property of matter to take on states in which the qubit is simultaneously both zero and one. Quantum algorithmic approaches to classical computing tasks have been of deep theoretical interest since the early 1990s and, in recent years, quantum machines have made implementation of these algorithms physically possible.¹⁰ Quantum computing is one of several areas in which advances in the broader field of computer science can help drive progress in AI.

For example, quantum computing can be applied to *optimization and search problems* in AI. In this context, quantum phenomena are employed as an analog to stochastic search or statistical optimization. A quantum approach can enable rapid convergence on the minima or maxima of the energy landscape without having to perform a discrete search over the variable assignments. Another emerging opportunity for quantum devices is their use for *physical simulation*. Quantum devices promise to be able to perform simulations using direct atomic physics—rather than as discrete approximations. As such, quantum devices could prove useful for understanding states in the physical world, creating more accurate models, and improving the fidelity of what computers are capable of. While these capabilities are not specific to AI, they could prove particularly useful for AI techniques that make explicit inferences using a world model.

The use of quantum computing for AI faces an array of challenges. Three particularly stand out. First, for a given AI problem to be aided by quantum computing, an appropriate *problem encoding* must be developed that reduces the traditionally intractable aspects of an AI problem to tasks that can be executed on a quantum device. Second, using quantum devices as a component in an AI system requires *digital-to-analog conversion* of key data required by the quantum component, as well as analog-to-digital conversion of the quantum component’s output. These conversions are time consuming and prone to errors and noise, and the

¹⁰There are many references around topics of quantum computing and AI. Specific to the use of quantum machines for search, optimization, and machine learning problems there are many dozens of papers. Three of a survey nature are: Biamonte et al. (2017); Schuld et al. (2014); and Rajak et al. (2022).

computational cost of these conversions can sometimes outweigh any benefits gained from the quantum acceleration. Third, the *scale* of quantum computing devices, usually represented as the number of qubits, is very limited in comparison to traditional CMOS devices.¹¹ Hence, state spaces requiring thousands of variables are out of direct reach for the current generation of quantum hardware.

Near-Term Advances in Deep Neural Networks

An eventual paradigm shift in AI research is virtually inevitable, but the timing and nature of that paradigm shift are far from certain. When it does occur, the shift may be toward one or more of the directions described above, or it may be toward a currently unforeseen paradigm. That said, before (or even after) such a shift occurs, continued improvements based on the current paradigm of deep networks will presumably continue, and many of those improvements will originate from the rich interactions between academic and industrial research.

Since the early 2010s, large deep neural networks (DNNs) have been the dominant paradigm in AI due to empirical evidence demonstrating performance improvements over earlier types of machine learning models across a wide range of tasks. These improvements were driven primarily by a combination of algorithmic advances, increases in the scale of computation, and the availability of large-scale data. Huge investments have been made in large foundation models using transformer architectures, pre-trained on massive amounts of data, leading to impressive capabilities in natural language understanding and generation, vision, audio, and multi-modal tasks. This has enabled wide-ranging applications from conversational AI to creative content generation and complex decision-making.

There is a growing recognition, however, that the future of scaling for DNNs is limited, with diminishing performance returns relative to the costs of training, the finite nature of high-quality data, and the various problems associated with ever-increasing energy consumption and cost. This has led to an increased focus among researchers in both academia and industry on efficient scaling techniques, sparse and modular architectures, model compression, combining specialized smaller models, employing mixtures of experts, exploring synthetic data, and other directions that aim to provide alternatives to brute-force scaling. Interestingly, many of these techniques were originally developed for AI systems from previous paradigms, but have been repurposed to improve DNN performance.

Modern AI systems are complex ecosystems integrating numerous components beyond the core foundation DNN model, including prompt management, connecting to external memory and knowledge sources, additional reinforcement learning-based training, fine-tuning modules, and more. This modularity allows for greater flexibility, efficiency, and robustness than a single DNN trying to do everything. Many of these components of the AI stack are DNNs, though generally smaller and more specialized than the core foundation model. Yet these current models still fall short of important goals such as interpretability, continual learning, and alignment with complex values and preferences, and there is currently no clear indication of

¹¹Certain commercial machines report having on the order of 1000s of physical qubits, but the number of logical qubits available for computing purposes is often considerably smaller due to the redundancies required for quantum error correction.

when these approaches will plateau and no clear roadmap to the next architectural breakthrough.

These trends point toward a new generation of heterogeneous DNN-based systems to increase both performance and capabilities. Such systems could support improvements in reasoning, causal understanding, complex planning, robustness, long-term memory, and other critical areas that the current generation of large foundation models has not yet mastered. Other existing directions of research may be promising for next-generation DNNs, including new models of attention, relational architectures, state-space sequence models, dynamic/adaptive/self-modifying networks, architectures for neural operator learning, networks explicitly structured for reasoning (e.g., causal representation learning), and networks with domain-specific information embedded directly into layers. The integration of domain knowledge or constraints into networks may significantly improve performance, interpretability, and reliability. For example, physics-informed neural networks (PINNs) directly integrate knowledge of physical laws, principles, and equations; graph neural networks represent the relational structure of the data; and causal inference neural networks aim to learn causal relationships between variables rather than just correlations.

Seeding the Next AI Revolution

One of the most visible recent developments in AI has been the rapid emergence of intense industrial competition among both startups and traditional tech companies to deploy AI technology. Virtually all these companies have staked their future on exploiting the unexpected capabilities of deep neural networks, large language models, and other foundation models. All are building on a common base of existing research results and computing technology with relatively little “moat” to separate their products from those of other companies. This has resulted in fierce competition for talent, capital, customers, and ideas.

Such fierce competition can be good for consumers, as it focuses companies on rapid development and deployment of products that satisfy customer needs. However, it further reinforces one of the most common limitations of industrial research — a focus on short-term product development rather than long-term, paradigm-shifting research. That is, industrial research tends to focus on exploiting and extending what has already been shown to work well, rather than pursuing riskier research directions with higher long-term payoff.

Such long-term, paradigm-shifting research is where academic research typically excels. Indeed, the current paradigm of AI research resulted from decades of academic research on neural networks, language models, computer vision, and other areas that were predominantly supported by government funding. However, that research was not supported because the recent extraordinary success of deep neural networks was easy to foresee. Rather, that support was part of an intentional investment strategy to support a diversified portfolio of alternative paradigms, explicitly acknowledging and managing the uncertainty about where revolutionary new technological developments would emerge.

Exploring many of the potential future paradigms for AI research requires the sort of sustained effort that only government funding typically supports. For example, successfully fostering

multidisciplinary paradigms (e.g., neuromorphic AI, quantum AI, and embodied AI) is typically accomplished by sustained federal funding of dedicated research institutes. Such institutes bring together a critical mass of research talent in the disparate specialties that make up the field and keep those researchers in close proximity for a sufficient time for collaborations to solidify and flourish. This environment can also convince new students to enter the field and can foster early-career researchers and early-stage commercialization efforts.

The United States is already benefiting tremendously because the current AI revolution has largely been centered within the borders of North America. However, what is true of the *current* revolution need not be true of the *next* revolution. Without a sustained and robust portfolio of research efforts in both industry and academia, the next AI revolution could well be centered in other nations, with those nations reaping the benefits and controlling the direction of future innovation in AI.

References

- Altman, S. (2025). Three Observations. <https://blog.samaltman.com/three-observations>.
- Biamonte, J., Wittek, P., Pancotti, N., Rebentrost, P., Wiebe, N., & Lloyd, S. (2017). Quantum machine learning. *Nature*, 549(7671), 195–202. <https://doi.org/10.1038/nature23474>.
- Bhuyan, B. P., Ramdane-Cherif, A., Tomar, R., & Singh, T. P. (2024). Neuro-symbolic artificial intelligence: a survey. *Neural Computing and Applications*, 36(21), 12809-12844.
- Brooks, R. A. (1991). Intelligence without reason. *Proceedings of 12th International Joint Conference on Artificial Intelligence (IJCAI)*, Sydney, Australia, August. 569–595.
- Deitke, M., Batra, D., Bisk, Y., Campari, T., Chang, A. X., Singh Chaitin, D., Chen, C., Pérez D'Arpino, C., Ehsani, K., Farhadi, A., Fei-Fei, L., Francis, A., Gan, C., Grauman, K., Hall, D., Han, W., Jain, U., Kembhavi, A., Krantz, J., Lee, S., Li, C., Majumder, S., Maksymets, O., Martín-Martín, R., Mottaghi, R., Raychaudhuri, S., Roberts, M., Savarese, S., Savva, M., Shridhar, M., Sünderhauf, N., Szot, A., Talbot, B., Tenenbaum, J. B., Thomason, J., Toshev, A., Truong, J., Weihs, L., & Wu, J. (2022). *Retrospectives on the Embodied AI Workshop*. arXiv e-prints, arXiv:2210.06849.
- Hadfield, J., Zhang, B., Lien, K., Scholz, F., Fox, J., & Ford, D. (2025) How we built our multi-agent research system. <https://www.anthropic.com/engineering/built-multi-agent-research-system>
- Jin, Z., Chen, Y., Leeb, F., Gresele, L., Kamal, O., Lyu, Z., Jin, Z., Blin, K., Adauto, F.G., Kleiman-Weiner, M., Sachan, M. & Schölkopf, B. (2023). CLadder: Assessing causal reasoning in language models. In *Proceedings of the 37th International Conference on Neural Information Processing Systems*, 31038-31065.
- Kudithipudi, D., Schuman, C., Vineyard, C. M., Pandit, T., Merkel, C., Kubendran, R., Aimone, J.B., Orchard, G., Mayr, C., Benosman, R., Hays, J., Young, C., Bartolozzi, C., Majumdar, A., Cardwell, S.G., Payvand, M., Buckley, S., Kulkarni, S., Gonzalez, H.A., Cauwenberghs, G., Thakur, C.S., Subramoney, A. & Furber, S. (2025). Neuromorphic computing at scale. *Nature*, 637(8047), 801-812.
- Maher, M. L., Ventura, D., & Magerko, B. (2024). The grounding problem: An approach to the integration of cognitive and generative models. *Proceedings of the AAAI Symposium Series*, 2(1), 320-325. <https://doi.org/10.1609/aaais.v2i1.27695>
- Muir, D. R., & Sheik, S. (2025). The road to commercial success for neuromorphic technologies. *Nature Communications*, 16(1), 3586.
- Ozmen Garibay, O., Winslow, B., Andolina, S., Antona, M., Bodenschatz, A., Coursaris, C., Falco, G., Fiore, S.M., Garibay, I., Grieman, K., Havens, J.C., Jirotko, M., Kacorri, H., Karwowski, W., Kider, J., Konstan, J., Koon, S., Lopez-Gonzalez, M., Maifeld-Carucci, I., McGregor, S., Salvendy, G., Shneiderman, B., Stephanidis, C., Strobel, C., Holter, C.T. & Xu, W.

(2023). Six human-centered artificial intelligence grand challenges. *International Journal of Human-Computer Interaction*, 39(3), 391-437.

Pfeifer, R. and Iida, F. (2004). *Embodied Artificial Intelligence: Trends and Challenges*. Lecture Notes in Computer Science (Vol 3139). 1-26.

Rajak, A., Suzuki, S., Dutta, A., & Chakrabarti, B. K. (2022). Quantum annealing: An overview. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 381(2241), 20210417. <https://doi.org/10.1098/rsta.2021.0417>.

Riedl, M. O. (2019). Human-centered artificial intelligence and machine learning. *Human Behavior and Emerging Technologies*, 1(1), 33-36.

Rossi, F., Bessiere, C., Biswas, J., Brooks, R., Conitzer, V., Dietterich, T. G., Dignum, V., Etzioni, O., Forbus, K. D., Freuder, E., Gil, Y., Hoos, H., Horvitz, E., Kambhampati, S., Kautz, H., Kim, J., Kitano, H., Mackworth, A., Myers, K., De Raedt, L., Russell, S., Selman, B., Stone, P., Tambe, M. & Wooldridge, M. (2025). *AAAI 2025 Presidential Panel on the Future of AI Research*. AAAI. March.

Schuld, M., Sinayskiy, I., & Petruccione, F. (2014). An introduction to quantum machine learning. *Contemporary Physics*, 56(2), 172–185.
<https://doi.org/10.1080/00107514.2014.964942>; (3)

Shneiderman, B. (2022). *Human-Centered AI*. Oxford University Press.

Shoham, Y and Leyton-Brown, K. (2008). *Multiagent Systems: Algorithmic, Game Theoretic, and Logical Foundations*. Cambridge University Press.

Valmeekam, K., Stechly, K., Gundawar, A., & Kambhampati, S. (2025). A systematic evaluation of the planning and scheduling abilities of the reasoning model o1. *Transactions on Machine Learning Research*.

Zhang, K., Yang, Z., & Başar, T. (2021). Multi-agent reinforcement learning: A selective overview of theories and algorithms. *Handbook of Reinforcement Learning and Control*. 321-384.

Suggested Citation

Jensen, D., Danks, D., Elbaum, S., Regli, W., Turk, M., Wierman, A., Yanco, H., Maher, M., Griffin, H. *Envisioning Possible Futures for AI Research*. Washington, D.C.: Computing Research Association (CRA), July 2025.
https://cra.org/ccc/wp-content/uploads/sites/2/2025/07/Envisioning_Possible_Futures_for_AI_Research.pdf

ACKNOWLEDGMENTS

Reviewers

CCC and the authors acknowledge the reviewers whose thoughtful comments improved the report:

- William Gropp, University of Illinois Urbana-Champaign
- Manish Parashar, University of Utah

U.S. National Science Foundation



The CCC AI Futures Task Force was supported by the Computing Community Consortium through the National Science Foundation under Grant No. 2300842. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of the National Science Foundation.

About the Computing Community Consortium (CCC)

A programmatic committee of the Computing Research Association (CRA), CCC enables the pursuit of innovative, high-impact computing research that aligns with pressing national and global challenges. Of, by, and for the computing research community, CCC is a responsive, respected, and visionary organization that brings together thought leaders from industry, academia, and government to articulate and advance compelling research visions and communicate them to stakeholders, policymakers, the public, and the broad computing research community.