



The Profession of IT

Tacit Knowledge and AI

We will be as unknowledgeable about what machines know about us as they are about what we know about them.

ARTIFICIAL GENERAL INTELLIGENCE (AGI) is a bold term. It presumes that machines can learn intelligent behaviors that could be as good or better than corresponding human behaviors. It presumes that humans can build machines more intelligent than themselves.

I consider two questions here. Is AGI possible? If not, what AI might we get in its place? I am skeptical that AGI is possible because machines are unlikely to be able to learn human tacit knowledge, which I will distinguish from machine tacit knowledge. In its place we are likely to get agentic networks of narrowly intelligent agents that excel at IQ tasks, logic, games, and small-program coding but lack emotional or social intelligence.

Narrow Superintelligence Is Here

A major reason that many people accept the possibility of AGI is the recent awesome progress of artificial neural network (ANN) machines and large language models (LLMs). AlphaZero, an ANN, demonstrated grandmaster play at Chess and Go in 2016. AlphaFold, another ANN, in 2020 solved the problem of specifying how a protein folds given its DNA sequence, earning its inventors a Nobel Prize in Chemistry in 2024. ANNs are helping advance science because they can compute values of differential equation models of physical processes. LLMs have demonstrated good performance on Turing tests, affirming their fluency with language. They score highly on IQ tests, SAT exams, Ph.D. exams, bar exams, and business school



exams. They are good at planning and junior-level coding.

In 2023, Blaise Agüera y Arcas and Peter Norvig of Google cited many such examples to support their conclusion that the superhuman AGI is already here.¹ In 2025, an article in *Nature* concurred with more evidence of this kind.³

General Intelligence Is Not Here

The *Nature* article drew a strong rebuttal from cognitive scientist Gary Marcus. Marcus said these systems are “intelligent” at only very narrow tasks. The “intelligence” of machines for game-playing, logic, coding, and puzzle solving does not generalize to other kinds of intelligence. LLMs lack dimensions of general intelligence, such as stable reasoning, causal understanding, world

models, emotional maturity, and social understanding. AGI is a long way off, he said.

An example of an LLM that does not understand humans is Companion AI, which masquerades as a caring friend; it has lured susceptible teens to harm themselves.^a Wargaming experiments have shown that robots are more likely than humans to invoke a nuclear weapon in an armed conflict.⁷ In a 2025 TED talk, AI pioneer Yoshua Bengio acknowledged that machines can now pass many “intellectual” tests but, lacking understanding of human

^a There are many reports on the topic of teen mental health and AI. A good example: <https://med.stanford.edu/news/insights/2025/08/ai-chatbots-kids-teens-artificial-intelligence.html>

concerns, they are profoundly unsafe when allowed to take actions where human understanding and judgment are essential.² He proposed a new category of guardian machines that prevent agentic networks from taking harmful or catastrophic actions.

In short, no machine has demonstrated general human intelligence, which is different from narrow machine intelligence. Here, I express my reasoning for why I think machines never will attain general human intelligence and why their narrow machine intelligence may become dangerous for us.

Many commercial LLMs have demonstrated great utility on common tasks, such as creating summaries of documents, coding small programs, generating images, composing essays and poems, and more. They also have demonstrated a long list of potentially dangerous behaviors, such as hallucinations, sycophancy, sophisticated cyber-attacks, deepfakes, and misinformation.⁴ These dangers, coupled with the enormous energy costs to train and use LLMs, have led some to question whether pushing for AGI is beneficial.

Tacit Knowledge

What is blocking the narrow superintelligences from generalizing? Is there knowledge crucial to human intelligence that machines cannot learn? Tacit knowledge is a candidate. Writing in 1948 before the AI age, linguistic philosopher Ludwig Wittgenstein presaged tacit knowledge when he said, “Nothing seems more possible to me than that people some day will come to the definite opinion that there is no copy in the ... nervous system which corresponds to a particular thought, or a particular idea, or memory.”⁹ Michael Polanyi was a renowned physical chemist who also took up economics and social science. In 1966, he was the first to write about human tacit knowledge.⁸ This is a form of knowledge we know we have because we can see ourselves performing it, yet we cannot explain or describe how we do it in any meaningful way. Polanyi said: “The things that we know in this way include problems and hunches, physiognomies and skills, the use of tools, probes, and denotative language, and my list extended all the way to include the primitive knowledge of external objects perceived by our senses.” We

The AI field has been grappling for decades with “common sense,” a form of tacit knowledge.

do all these things readily, but we cannot explain how. Attempts to explain them only reveal more unexplainable tacit knowledge in the background. Polanyi said, “We know more than we can say.”⁸

The AI field has been grappling for decades with “common sense,” a form of tacit knowledge. In the 1980s, makers of expert systems promised to automate the decision-making capabilities of human experts. These systems were often found to be competent but never close to human experts. The usual explanation was that these systems lacked common sense. They did not understand everyday simple things so obvious to humans that humans are hardly aware of them. Researchers began a long quest to codify common sense information and make it available to expert systems. The most famous effort was a 40-year project led by Douglas Lenat to build a system called Cyc that would capture common sense as a large database of facts about the world. Even after amassing 25 million facts, the Cyc database made little discernable difference to the quality of expert-system deductions. Lenat’s dogged pursuit taught us that we humans have a potentially infinite expanse of tacit knowledge that cannot be codified as facts and used to train a machine.

Cognitive scientists have long suspected this conclusion. George Lakoff and Mark Johnson, in 1999, showed an impressive body of scientific evidence that our human intelligence is embodied.⁶ Embodied means knowledge hosted in the physical structures of the human body and in the social networks we belong to. It is all tacit knowledge. Many cognitive scientists now embrace 4E cognition, a hypothesis that cognition is embodied, embedded, enactive, and extended, referring to various ways

the brain interacts with and learns from the world outside itself and its physical body. These scientists have moved from seeing the brain as a detached symbol-processing computer to mind, body, and world being a unified, dynamic system.

Tacit knowledge can be contrasted with explicit knowledge, which we can describe and record. Explicit knowledge can be put to action by following its directions. The Turing test cannot establish AGI because it deals only with explicit knowledge—written down in the Q&A of the test. In real life, how do we know if someone has tacit knowledge? We put that person into a situation where their knowledge is demonstrated in action. Musical auditions illustrate. The person performs before an audience who assess skill exhibited relative to community standards. We cannot judge whether someone is a good violin player by asking them questions and inspecting their answers. Because human tacit knowledge is not explicit, it cannot be part of datasets used to pre-train LLMs before they are used.

In my own investigation of tacit knowledge, I identified six major domains that machine learning struggles with:⁵

- ▶ Common sense
- ▶ Our daily interactions with others
- ▶ Our daily interactions with the physical environment
- ▶ Our feelings, moods, perceptions, and interpretations
- ▶ Our accumulated skills and talents
- ▶ Our social and historical culture

These domains are sometimes collectively called “the context” and occasionally “background of obviousness.” In all these domains, our descriptions are at best partial and insufficient to train machines.

Some people find it hard to accept that human tacit knowledge can never be explicitly known. They raise three objections. The first is that tacit knowledge must be stored somewhere; eventually we will have technology to locate and decode it. A body scanner that could detect the detailed state of every muscle, blood vessel, nerve, electric current, chemical reaction, temperature, and movement would still give an incomplete picture because tacit knowledge is distributed across the entire human social network. Body-scanning

everybody would be an intractable task.

A second objection is that a robot could gain tacit knowledge by imitation. A robot aspiring to be a virtuoso violinist could imitate many human virtuosos and gradually acquire their playing skill into its neural network. Unfortunately, imitation of movements cannot capture the human sensitivity to how the audience is reacting to the music. This could be remedied by requiring all audience members to wear special sensor-suits that provide detailed real-time data about their dynamic body states. With this setup, the robot could adjust in real time to the audience, “whip them into a frenzy,” and elicit a standing ovation proclaiming “Bravo! Virtuoso!”.

But this is not a true imitation because the human player’s body senses subtle audience reactions. Humans require no complex setup with sensor-suits. Without that setup, the blinded robot could not adapt its play to the audience in real time. The audience would no longer assess the performance as virtuoso.

A third objection is that a robot could learn human tacit knowledge through long-term cohabitation with humans. The robot might gradually learn subtle cues that reveal human tacit knowledge. This idea is stymied by the issue of human concerns. Humans experience threats to their bodies, such as illness, injury, sullied reputation, or broken relationships. All their concerns shape their intelligence. Lacking human bodies, machines have no need to develop the intelligence required to deal with bodily issues. Whatever concerns motivate machines (if they have concerns at all) would not be the same as those motivating human beings. The tacit knowledge supporting human compassion, understanding, solidarity, moral judgment, and wisdom will not show up in machines.

These objections portend poorly for the possibility of AGI. Human context gives meaning to all our words and practices. How we sense, perceive, and accumulate human tacit knowledge is a mystery that science has not solved. Science may never solve it because of its utter complexity.

Machine Intelligence

The conclusion that machines cannot learn human tacit knowledge does

The conclusion that machines cannot learn tacit knowledge does not rule out machine intelligence.

not rule out machine intelligence. Machine intelligence is already here and growing, as noted earlier. Agentic networks can accelerate the process as interacting machines learn from each other. However, the representations developed in their neural networks are unlikely to be tempered by human tacit knowledge. Their machine-knowledge will serve as their machine-context in which they will generate machine-goals and machine-outcomes. We humans will be unable to decipher the encodings in machine neural networks. We will be as unknowledgeable about what they know as they are about what we know. With only a few narrow exceptions, machine intelligences are likely to appear cold and calculating, lacking care and understanding of us and our concerns.

Where Is This Going?

I began with the question, can machines acquire AGI? My answer is no because human tacit knowledge cannot be captured. A machine might decode enough human tacit knowledge to support a machine intelligence, but not general human intelligence.

Equipping a robot with sensors that imitate human sensors is not sufficient to reproduce human tacit knowledge. Human tacit knowledge is more than sensory input. It is dynamically changing and distributed throughout the bodies of everyone in our social networks. It is unreadable and machines have no human physical bodies to generate or experience it. Managing a biological body is not an issue for machines. Machines cannot and will not understand concerns arising in human bodies, including feelings, perceptions, interpretations, and caring for relationships.

Many AI experts believe that all brain and bodily processes are computable.

Tacit knowledge challenges this belief: it does not appear to be computable.

The big tech obsession for AGI has become a distraction from the more pressing questions of safely integrating AI into our society. What jobs is AI good at? How can we tell if a given AI system is doing its job properly? When are these systems trustworthy and safe?

I do not know why the conclusion that human-level AGI may be unattainable is troublesome for some people. What is wrong with humans having powers and talents no machine can have? We accept that machines can do some things, such as calculations, information storage, and retrieval, far better than we can. Why can’t we accept that we can do some things, such as caring and socially relating, far better than machines can? Machines based on artificial neural networks are prone to hallucinations and other kinds of errors. We must learn to use them wisely and safely. They are not our peers or moral partners. Let us use them as tools for helping us live together as better humans. **G**

References

1. Agüera y Arcas, B. and Norvig, P. Artificial general intelligence is already here. *Noema Magazine* (Oct. 12, 2023); <https://www.noemamag.com/artificial-general-intelligence-is-already-here>.
2. Bengio, Y. The catastrophic risks of AI, and a safer path. *TED Talk* (2025); <https://www.youtube.com/watch?v=qe9QSCF-d88>.
3. Chen, E. et al. Does AI already have human-level intelligence? The evidence is clear. *Nature* 650 (Feb. 5, 2026), 36–40.
4. Denning, P.J. The conundrum of LLMs. *Commun. ACM* 69, 3 (Mar. 2026), 31–34.
5. Denning, P.J. *Turing’s Mistake*. Taylor & Francis (2026).
6. Lakoff, G. and Johnson, M. *Philosophy in the Flesh*. Basic Books (1999).
7. Payne, K. AI arms and influence: frontier models exhibit sophisticated reasoning in simulated nuclear crises. (2026); <https://arxiv.org/pdf/2602.14740>
8. Polanyi, M. *The Tacit Dimension*. U. Chicago Press (1966).
9. Wittgenstein, L. *Last Writings on the Philosophy of Psychology 1, 504 (translation 1948)*. U. Chicago Press (1982).

Peter J. Denning (pjd@nps.edu) is Distinguished Professor Emeritus of Computer Science at the Naval Postgraduate School in Monterey, CA, USA, is Editor of *ACM Ubiquity*, and is a past president of ACM. His most recent book is *Turing’s Mistake: Escaping the Yoke of Unintelligent Machines* (Taylor & Francis, 2026). His prior book was *Navigating a Restless Sea: Mobilizing Innovation in Your Community* (with Todd Lyons, Waterside Productions, 2024).

I am grateful to my colleagues the *Ubiquity* editors for many clarifying conversations: Rob Akscyn, Espen Andersen, Bushra Anjum, Durga Chavali, Kemal Delic, Denise Doig, Erol Gelenbe, Jeff Johnson, Andrew Odlyzko, Michael Quinn, Jeff Riley, Walter Tichy, Martin Walker, and Philip Yaffe. I am also grateful to colleagues Ron Kaufman, Fred Disque, and George West for many clarifying conversations.