

WINNING COVID-19 SPREAD IN A ROBOTIC ARTICULATION: THE INTIMATION OF ARTIFICIAL INTELLIGENCE

Bartholomew Uchenna Arum*

Abstract

Computer science can be divided into four main fields: software development, computer architecture (hardware), human-computer interfacing (the design of the most efficient ways for humans to use computers), and *artificial intelligence* (the attempt to make computers behave intelligently). Software development is concerned with creating computer programs that perform efficiently. Computer architecture is concerned with developing optimal hardware for specific computational needs. The areas of artificial intelligence (AI) and human-computer interfacing often involve the development of both software and hardware to solve specific problems.

Keywords: Covid-19, Artificial Intelligence, Robotic Articulation

Introduction

Artificial Intelligence (AI) is the study and engineering of intelligent machines capable of performing the same kinds of functions that characterize human thought. Again, Artificial intelligence is that branch of computer science that develops programs to allow machines to perform functions normally requiring human intelligence. Hence, Computer intelligence is the ability of computers to perform functions that normally require human intelligence. The concept of Artificial Intelligence (AI) dates from ancient times, but the advent of digital computers in the 20th century brought AI into the realm of possibility. AI was conceived as a field of computer science in the mid-1950s. The term AI has been applied to computer programs and systems capable of performing tasks more complex than straightforward programming, although still far from the realm of actual thought. While the nature of intelligence remains elusive, AI capabilities currently have far-reaching applications in such areas as information processing, computer gaming, national security, electronic commerce, and diagnostic systems.

The History of AI

The field of artificial intelligence (AI) officially started in 1956, launched by a small but now-famous DARPA-sponsored summer conference at Dartmouth College, in Hanover, New Hampshire. (The 50-year celebration of this conference, AI@50, was held in July 2006 at Dartmouth, with five of the original participants making it back. What happened at this historic conference figures in the final section of this entry.) Ten thinkers attended, including John McCarthy (who was working at Dartmouth in 1956), Claude Shannon, Marvin Minsky, Arthur Samuel, Trenchard Moore (apparently the lone note-taker at the original conference), Ray Solomonoff, Oliver Selfridge, Allen Newell, and Herbert Simon. From where we stand now, into the start of the new millennium, the Dartmouth conference is memorable for many reasons, including this pair: one, the term

‘artificial intelligence’ was coined there (and has long been firmly entrenched, despite being disliked by some of the attendees, e.g., Moore); two, Newell and Simon revealed a program – Logic Theorist (LT) – agreed by the attendees (and, indeed, by nearly all those who learned of and about it soon after the conference) to be a remarkable achievement. LT was capable of proving elementary theorems in the propositional calculus.

Though the *term* ‘artificial intelligence’ made its advent at the 1956 conference, certainly the *field* of AI, operationally defined (defined, i.e., as a field constituted by practitioners who think and act in certain ways), was in operation before 1956. For example, in a famous *Mind* paper of 1950, Alan Turing argues that the question “Can a machine think?” (and here Turing is talking about standard computing machines: machines capable of computing functions from the natural numbers (or pairs, triples, ... thereof) to the natural numbers that a Turing machine or equivalent can handle) should be replaced with the question “Can a machine be linguistically indistinguishable from a human?.” Specifically, he proposes a test, the “Turing Test” (TT) as it’s now known. In the TT, a woman and a computer are sequestered in sealed rooms, and a human judge, in the dark as to which of the two rooms contains which contestant, asks questions by email (actually, by *teletype*, to use the original term) of the two. If, on the strength of returned answers, the judge can do no better than 50/50 when delivering a verdict as to which room houses which player, we say that the computer in question has passed the TT. Passing in this sense operationalizes linguistic indistinguishability. Later, we shall discuss the role that TT has played, and indeed continues to play, in attempts to define AI. At the moment, though, the point is that in his paper, Turing explicitly lays down the call for building machines that would provide an existence proof of an affirmative answer to his question. The call even includes a suggestion for how such construction should proceed. (He suggests that “child machines” be built, and that these machines could then gradually grow up on their own to learn to communicate in natural language at the level of adult humans.

This suggestion has arguably been followed by Rodney Brooks and the philosopher Daniel Dennett (1994) in the Cog Project. In addition, the Spielberg/Kubrick movie *A.I.* is at least in part a cinematic exploration of Turing’s suggestion.^[5]) The TT continues to be at the heart of AI and discussions of its foundations, as confirmed by the appearance of (Moor 2003). In fact, the TT continues to be used to *define* the field, as in Nilsson’s (1998) position, expressed in his textbook for the field, that AI simply is the field devoted to building an artifact able to negotiate this test. Energy supplied by the dream of engineering a computer that can pass TT, or by controversy surrounding claims that it has *already* been passed, is if anything stronger than ever, and the reader has only to do an internet search via the string turing test passed to find up-to-the-minute attempts at reaching this dream, and attempts (sometimes made by philosophers) to debunk claims that some such attempt has succeeded.

Returning to the issue of the historical record, even if one bolsters the claim that AI started at the 1956 conference by adding the proviso that ‘artificial intelligence’ refers to a nuts-and-bolts *engineering* pursuit (in which case Turing’s philosophical discussion, despite calls for a child machine, wouldn’t exactly count as AI per se), one must confront the fact that Turing, and indeed many predecessors, did attempt to build intelligent artifacts. In Turing’s case, such building was surprisingly well-understood before the advent of programmable computers: Turing wrote a program for playing chess before there were computers to run such programs on, by slavishly following the code himself. He did this well before 1950, and long before Newell (1973) gave thought in print to the possibility of a sustained, serious attempt at building a good chess-playing computer.

From the perspective of philosophy, which views the systematic investigation of mechanical intelligence as meaningful and productive separate from the specific logicist formalisms (e.g., first-order logic) and problems (e.g., the *Entscheidungsproblem*) that gave birth to computer science, neither the 1956 conference, nor Turing’s *Mind* paper, come close to marking the start of AI. This is easy enough to see. For example, Descartes proposed TT (not the TT by name, of course) long before Turing was born.^[7] Here’s the relevant passage:

If there were machines which bore a resemblance to our body and imitated our actions as far as it was morally possible to do so, we should always have two very certain tests by which to recognize that, for all that, they were not real men. The first is, that they could never use speech or other signs as we do when placing our thoughts on record for the benefit of others. For we can easily understand a machine’s being constituted so that it can utter words, and even emit some responses to action on it of a corporeal kind, which brings about a change in its organs; for instance, if it is touched in a particular part it may ask what we wish to say to it; if in another part it may exclaim that it is being hurt, and so on. But it never happens that it arranges its speech in various ways, in order to reply appropriately to everything that may be said in its presence, as even the lowest type of man can do. And the second difference is, that although machines can perform certain things as well as or perhaps better than any of us can do, they infallibly fall short in others, by which means we may discover that they did not act from knowledge, but only for the disposition of their organs. For while reason is a universal instrument which can serve for all contingencies, these organs have need of some special adaptation for every particular action. From this it follows that it is morally impossible that there should be sufficient diversity in any machine to allow it to act in all the events of life in the same way as our reason causes us to act. (Descartes 1637, 116)

At the moment, Descartes is certainly carrying the day.^[8] Turing predicted that his test would be passed by 2000, but the fireworks across the globe at the start of the new millennium have long since died down, and the most articulate of computers still can’t meaningfully debate a sharp toddler. Moreover, while in certain focused areas machines out-perform minds (IBM’s famous Deep Blue prevailed in chess over Gary Kasparov, e.g.; and more recently, AI systems have prevailed in other

games, e.g. *Jeopardy!* and Go, about which more will momentarily be said), minds have a (Cartesian) capacity for cultivating their expertise in virtually *any* sphere. (If it were announced to Deep Blue, or any current successor, that chess was no longer to be the game of choice, but rather a heretofore unplayed variant of chess, the machine would be trounced by human children of average intelligence having no chess expertise.) AI simply hasn't managed to create *general* intelligence; it hasn't even managed to produce an artifact indicating that *eventually* it will create such a thing.

So far we have been proceeding as if we have a firm and precise grasp of the nature of AI. But what exactly *is* AI? Philosophers arguably know better than anyone that precisely defining a particular discipline to the satisfaction of all relevant parties (including those working in the discipline itself) can be acutely challenging. Philosophers of science certainly have proposed credible accounts of what constitutes at least the general shape and texture of a given field of science and/or engineering, but what exactly is the agreed-upon definition of physics? What about biology? What, for that matter, is philosophy, exactly? These are remarkably difficult, maybe even eternally unanswerable, questions, especially if the target is a *consensus* definition. Perhaps the most prudent course we can manage here under obvious space constraints is to present in encapsulated form some *proposed* definitions of AI. We do include a glimpse of recent attempts to define AI in detailed, rigorous fashion (and we suspect that such attempts will be of interest to philosophers of science, and those interested in this sub-area of philosophy).

Russell and Norvig (1995, 2002, 2009), in their aforementioned *AIMA* text, provide a set of possible answers to the “What is AI?” question that has considerable currency in the field itself. These answers all assume that AI should be defined in terms of its goals: a candidate definition thus has the form “AI is the field that aims at building ...” The answers all fall under a quartet of types placed along two dimensions. One dimension is whether the goal is to match human performance, or, instead, ideal rationality. The other dimension is whether the goal is to build systems that reason/think, or rather systems that act. The situation is summed up in this table:

	Human-Based	Ideal Rationality
Reasoning-Based:	Systems that think like humans.	Systems that think rationally.
Behavior-Based:	Systems that act like humans.	Systems that act rationally.

Four Possible Goals for AI According to AIMA

Please note that this quartet of possibilities does reflect (at least a significant portion of) the relevant literature. For example, philosopher John Haugeland (1985) falls into the Human/Reasoning quadrant when he says that AI is “The

exciting new effort to make computers think ... *machines with minds*, in the full and literal sense.” (By far, this is the quadrant that most popular narratives affirm and explore. The recent *Westworld* TV series is a powerful case in point.) Luger and Stubblefield (1993) seem to fall into the Ideal/Act quadrant when they write: “The branch of computer science that is concerned with the automation of intelligent behavior.” The Human/Act position is occupied most prominently by Turing, whose test is passed only by those systems able to act sufficiently like a human. The “thinking rationally” position is defended (e.g.) by Winston (1992). While it might not be entirely uncontroversial to assert that the four bins given here are exhaustive, such an assertion appears to be quite plausible, even when the literature up to the present moment is canvassed.

It’s important to know that the contrast between the focus on systems that think/reason versus systems that act, while found, as we have seen, at the heart of the *AIMA* texts, and at the heart of AI itself, should not be interpreted as implying that AI researchers view their work as falling all and only within one of these two compartments. Researchers who focus more or less exclusively on knowledge representation and reasoning, are also quite prepared to acknowledge that they are working on (what they take to be) a central component or capability within any one of a family of larger systems spanning the reason/act distinction. The clearest case may come from the work on planning – an AI area traditionally making central use of representation and reasoning. For good or ill, much of this research is done in abstraction (in vitro, as opposed to in vivo), but the researchers involved certainly intend or at least hope that the results of their work can be embedded into systems that actually do things, such as, for example, execute the plans.

Development of Artificial Intelligence

In 1956 American social scientist and Nobel laureate Herbert Simon and American physicist and computer scientist Allan Newell at Carnegie Mellon University in Pennsylvania devised a program called Logic Theorist that simulated human thinking on computers. The first AI conference occurred at Dartmouth College in New Hampshire in 1956. This conference inspired researchers to undertake projects that emulated human behavior in the areas of reasoning, language comprehension, and communications. In addition to Newell and Simon, computer scientists and mathematicians Claude Shannon, Marvin Minsky, and John McCarthy laid the groundwork for creating “thinking” machines from computers.

The search for AI has taken two major directions: psychological and physiological research into the nature of human thought, and the technological development of increasingly sophisticated computing systems. Some AI developers are primarily interested in learning more about the workings of the human brain and thus attempt to mimic its methods and processes. Other developers are more interested in making computers perform a specific task, which may involve computing methods well beyond the capabilities of the human brain.

Contemporary fields of interest resulting from early AI research include expert systems, cellular automata (treating pieces of data like biological cells), and artificial life. The search for AI goes well beyond computer science and involves cross-disciplinary studies in such areas as cognitive psychology, neuroscience, linguistics, cybernetics, information theory, and mechanical engineering, among many others. The search for AI has led to advancements in those fields, as well.

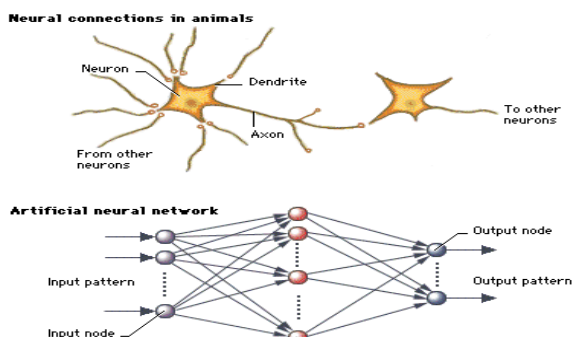
Types of Artificial Intelligence

Work in AI has primarily focused on two broad areas: developing logic-based systems that perform common-sense and expert reasoning, and using cognitive and biological models to simulate and explain the information-processing capabilities of the human brain. In general, work in AI can be categorized within three research and development types: symbolic, connectionist, and evolutionary. Each has characteristic strengths and weaknesses.

1. Symbolic AI

Symbolic AI is based in logic. It uses sequences of rules to tell the computer what to do next. Expert systems consist of many so-called IF-THEN rules: IF this is the case, THEN do that. Since both sides of the rule can be defined in complex ways, rule-based programs can be very powerful. The performance of a logic-based program need not appear “logical,” as some rules may cause it to take apparently irrational actions. “Illogical” AI programs are not used for practical problem-solving, but are useful in modeling how humans think. Symbolic programs are good at dealing with set problems, and at representing hierarchies (in grammar, for example, or planning). But they are inflexible: If part of the expected input data is missing or mistaken, they may give a bad answer or no answer at all.

2. Connectionist AI



Artificial Neural Network

Connectionism is inspired by the human brain. It is closely related to computational neuroscience, which models actual brain cells and neural circuits. Connectionist AI uses artificial neural networks made of many units working in parallel. Each unit is connected to its neighbors by links that can raise or lower the likelihood that the neighbor unit will “fire” (excitatory and inhibitory connections,

respectively). Neural networks that are able to learn do so by changing the strengths of these links, depending on past experience. These simple units are much less complex than real neurons. Each can do only one thing, such as report a tiny vertical line at a particular place in an image. What matters is not what any individual unit is doing, but the overall activity pattern of the whole network.

Consequently, connectionist systems are more flexible than symbolic AI programs. Even if the input data is faulty, the network may give the right answer. They are therefore good at pattern recognition, where the input patterns within a certain class need not be identical. But connectionism is weak at doing logic, following action sequences, or representing hierarchies of goals. What symbolic AI does well, connectionism does badly, and vice versa. Hybrid systems combine the two, switching between them as appropriate. And work on recurrent neural networks, where the output of one layer of units is fed back as input to some previous layer, aims to enable connectionist systems to deal with sequential action and hierarchy. The emerging field of *connectomics* could help researchers decode the brain's approach to information processing. See Neurophysiology; Nervous System.

3. Evolutionary AI

Evolutionary AI draws on biology. Its programs make random changes in their own rules, and select the best daughter programs to breed the next generation. This method develops problem-solving programs, and can evolve the “brains” and “eyes” of robots. A practical application of evolutionary AI would be a computer model of the long-term growth of a business in which the evolution of the business is set within a simulated marketplace. Evolutionary AI is often used in modeling artificial life (commonly known as A-Life), a spin-off from AI. One focus of study in artificial life is on self-organization, namely how order arises from something that is ordered to a lesser degree. Biological examples include the flocking patterns of birds and the development of embryos. Technological examples include the flocking algorithms used for computer animation.

Uses and Challenges of Artificial Intelligence

AI programs have a broad array of applications. They are used by financial institutions, scientists, psychologists, medical practitioners, design engineers, planning authorities, and security services, to name just a few. AI techniques are also applied in systems used to browse the Internet.

AI programs tend to be highly specialized for a specific task. They can play games, predict stock values, interpret photographs, diagnose diseases, plan travel itineraries, translate languages, take dictation, draw analogies, help design complex machinery, teach logic, make jokes, compose music, create drawings, and learn to do tasks better. AI programs perform some of these tasks well. In a famous example, a supercomputer called Deep Blue beat world chess champion Garry Kasparov in 1997. In developing its strategy, Deep Blue utilized parallel processing (interlinked and concurrent computer operations) to process 200 million

chess moves per second. AI programs are often better than people at predicting stock prices, and they can create successful long-term business plans. AI programs are used in electronic commerce to detect possible fraud, using complex learning algorithms, and are relied upon to authorize billions of financial transactions daily. AI programs can also mimic creative human behavior. For example, AI-generated music can sound like compositions by famous composers.

Some of the most widely used AI applications involve information processing and pattern recognition. For example, one AI method now widely used is “data mining,” which can find interesting patterns in extremely large databases. Data mining is an application of machine learning, in which specialized algorithms enable computers to “learn.” Other applications include information filtering systems that discover user interests in an online environment. However, it remains unknown whether computer programs could ever learn to solve problems on their own, rather than simply following what they are programmed to do.



WABOT-2 and Inventor

AI programs can make medical diagnoses as well as, or better than, most human doctors. AI programs have been developed that analyze the disease symptoms, medical history, and laboratory test results of a patient, and then suggest a diagnosis to the physician. The diagnostic program is an example of expert systems, which are programs designed to perform tasks in specialized areas as a human would. Expert systems take computers a step beyond straightforward programming, being based on a technique called rule-based inference, in which pre-established rule systems are used to process the data. Despite their sophistication, expert systems still do not approach the complexity of true intelligent thought.

Despite considerable successes AI programs still have many limitations, which are especially obvious when it comes to language and speech recognition. Their translations are imperfect, although good enough to be understood, and their dictation is reliable only if the vocabulary is predictable and the speech unusually clear. Research has shown that whereas the logic of language structure (syntax) submits to programming, the problem of meaning (semantics) lies far deeper, in the direction of true AI (or “strong” AI, in the parlance of developers). Developing natural-language capabilities in AI systems is an important focus of AI research. It involves programming computers to understand written or spoken information and

to produce summaries, answer specific questions, or redistribute information to users interested in specific areas. Essential to such programs is the ability of the system to generate grammatically correct sentences and to establish linkages between words, ideas, and associations with other ideas. “Chatterbot” programs, although far from natural conversationalists, are a step in that direction. They attempt to simulate an intelligent conversation by scanning input keywords to come up with pre-prepared responses from a database.

Much work in AI models intellectual tasks, as opposed to the sensory, motor, and adaptive abilities possessed by all mammals. However, an important branch of AI research involves the development of robots, with the goal of creating machines that can perceive and interact with their surroundings. WABOT-2, a robot developed by Waseda University in Japan in the 1980s, utilized AI programs to play a keyboard instrument, read sheet music, and converse rudimentarily with people. It was a milestone in the development of “personal” robots, which are expected to be anthropomorphous—that is, to emulate human attributes. AI robots are being developed as personal assistants for hospitalized patients and disabled persons, among other purposes. Natural-language capabilities are integral to these efforts. In addition, scientists with the National Aeronautics and Space Administration (NASA) are developing robust AI programs designed to enable the next generation of Mars rovers to make decisions for themselves, rather than relying on (and waiting for) detailed instructions from teams of human controllers on Earth.

To match everything that people can do, AI systems would need to model the richness and subtlety of human memory and common sense. Many of the mechanisms behind human intelligence are still poorly understood, and computer programs can simulate the complex processes of human thought and cognition only to a limited extent. Even so, an AI system does not necessarily need to mimic human thought to achieve an intelligent answer or result, such as a winning chess move, as it may rely on its own “superhuman” computing power.

Philosophical Debates on Artificial Intelligence

People often ask if artificial intelligence is possible, but the question is ambiguous. Certainly, AI programs can produce results that resemble human behavior. Some things that most people once assumed computers could never do are now possible due to AI research. For example, AI programs can compose aesthetically appealing music, draw attractive pictures, and even play the piano “expressively.” Other things are more elusive, such as producing perfect translations of a wide range of texts; making fundamental, yet aesthetically acceptable, transformations of musical style; or producing robots that can interact meaningfully with their surroundings. It is controversial whether these things are merely very difficult in practice, or impossible in principle.

The larger question of whether any program or robot could really be intelligent, no matter how humanlike its performance, involves highly controversial issues in the philosophy of mind, including the importance of embodiment and the nature of intentionality and consciousness. Some philosophers and AI researchers argue that intelligence can arise only in bodily creatures sensing and acting in the real world. If this is correct, then robotics is essential to the attempt to construct truly intelligent artifacts. If not, then a mere AI program might be intelligent.

British mathematician and computer scientist Alan Turing proposed what is now called the Turing Test as a way of deciding whether a machine is intelligent. He imagined a person and a computer hidden behind a screen, communicating by electronic means. If we cannot tell which one is the human, we have no reason to deny that the machine is thinking. That is, a purely behavioral test is adequate for identifying intelligence (and consciousness).

American philosopher John Searle has expressed a different view. He admits that a program might produce replies identical to those of a person, and that a programmed robot might behave exactly like a human. But he argues that a program cannot understand anything it says. It is not actually saying or asserting anything at all, but merely outputting meaningless symbols that it has manipulated according to purely formal rules—in other words, all syntax and no semantics. Searle asserts that human brains can ascribe meaning to symbols, thus deriving understanding, whereas metal and silicon cannot. No consensus exists in either AI or philosophy as to whose theory, Turing's or Searle's, is right.

Whether an AI system could be conscious is an especially controversial topic. The concept of consciousness itself is ill-understood, both scientifically and philosophically. Some would argue that any robot, no matter how superficially humanlike, would never possess the consciousness or sentience of a living being. But others would argue that a robot whose functions matched the relevant functions of the brain (whatever those may be) would inevitably be conscious. The answer has moral implications: If an AI system were conscious, it would arguably be wrong to "kill" it, or even to use it as a "slave."

Robot

Robot computer-controlled machine that is programmed to move, manipulate objects, and accomplish work while interacting with its environment. Robots are able to perform repetitive tasks more quickly, cheaply, and accurately than humans. The term *robot* originates from the Czech word *robota*, meaning "compulsory labor." It was first used in the 1921 play *R.U.R.* (Rossum's Universal Robots) by the Czech novelist and playwright Karel Capek. The word *robot* has been used since to refer to a machine that performs work to assist people or work that humans find difficult or undesirable.

Early History of Robots

The concept of automated machines dates to antiquity with myths of mechanical beings brought to life. Automata, or humanlike machines, also appeared in the clockwork figures of medieval churches, and 18th-century watchmakers were famous for their clever mechanical creatures. Feedback (self-correcting) control mechanisms were used in some of the earliest robots and are still in use today. An example of feedback control is a watering trough that uses a float to sense the water level. When the water falls past a certain level, the float drops, opens a valve, and releases more water into the trough. As the water rises, so does the float. When the float reaches a certain height, the valve is closed and the water is shut off.

The first true feedback controller was the Watt governor, invented in 1788 by the Scottish engineer James Watt. This device featured two metal balls connected to the drive shaft of a steam engine and also coupled to a valve that regulated the flow of steam. As the engine speed increased, the balls swung out due to centrifugal force, closing the valve. The flow of steam to the engine was decreased, thus regulating the speed. Feedback control, the development of specialized tools, and the division of work into smaller tasks that could be performed by either workers or machines were essential ingredients in the automation of factories in the 18th century. As technology improved, specialized machines were developed for tasks such as placing caps on bottles or pouring liquid rubber into tire molds. These machines, however, had none of the versatility of the human arm; they could not reach for objects and place them in a desired location.

The development of the multijointed artificial arm, or manipulator, led to the modern robot. A primitive arm that could be programmed to perform specific tasks was developed by the American inventor George Devol, Jr., in 1954. In 1975 the American mechanical engineer Victor Scheinman, while a graduate student at Stanford University in California, developed a truly flexible multipurpose manipulator known as the Programmable Universal Manipulation Arm (PUMA). PUMA was capable of moving an object and placing it with any orientation in a desired location within its reach. The basic multijointed concept of the PUMA is the template for most contemporary robots.

How Robots Work

The inspiration for the design of a robot manipulator is the human arm, but with some differences. For example, a robot arm can extend by telescoping—that is, by sliding cylindrical sections one over another to lengthen the arm. Robot arms also can be constructed so that they bend like an elephant trunk. Grippers, or end effectors, are designed to mimic the function and structure of the human hand. Many robots are equipped with special purpose grippers to grasp particular devices such as a rack of test tubes or an arc-welder.

The joints of a robotic arm are usually driven by electric motors. In most robots, the gripper is moved from one position to another, changing its orientation. A computer calculates the joint angles needed to move the gripper to the desired position in a process known as inverse kinematics. Some multijointed arms are equipped with servo, or feedback, controllers that receive input from a computer. Each joint in the arm has a device to measure its angle and send that value to the controller. If the actual angle of the arm does not equal the computed angle for the desired position, the servo controller moves the joint until the arm's angle matches the computed angle. Controllers and associated computers also must process sensor information collected from cameras that locate objects to be grasped, or they must touch sensors on grippers that regulate the grasping force.

Any robot designed to move in an unstructured or unknown environment will require multiple sensors and controls, such as ultrasonic or infrared sensors, to avoid obstacles. Robots, such as the National Aeronautics and Space Administration (NASA) planetary rovers, require a multitude of sensors and powerful onboard computers to process the complex information that allows them mobility. This is particularly true for robots designed to work in close proximity with human beings, such as robots that assist persons with disabilities and robots that deliver meals in a hospital. Safety must be integral to the design of human service robots.

Uses of Robots



Hospital Robot

More than 1 million robots are estimated to be in operation in the industrialized world. Many robot applications are for tasks that are either dangerous or unpleasant for human beings. In medical laboratories, robots handle potentially hazardous materials, such as blood or urine samples. In other cases, robots are used in repetitive, monotonous tasks in which human performance might degrade over time. Robots can perform these repetitive, high-precision operations 24 hours a day without fatigue. A major user of robots is the automobile industry. General Motors Corporation uses approximately 16,000 robots for tasks such as spot welding, painting, machine loading, parts transfer, and assembly. Assembly is one of the

fastest growing industrial applications of robotics. It requires higher precision than welding or painting and depends on low-cost sensor systems and powerful inexpensive computers. Robots are used in electronic assembly where they mount microchips on circuit boards.

Activities in environments that pose great danger to humans, such as locating sunken ships, cleaning up nuclear waste, prospecting for underwater mineral deposits, and exploring active volcanoes, are ideally suited to robots. Similarly, robots can explore distant planets. NASA's Galileo, an unpiloted space probe, traveled to Jupiter in 1996 and performed tasks such as determining the chemical content of the Jovian atmosphere. The robotic Mars Exploration rovers landed on Mars in 2003 and moved over the Martian surface for years, carrying out scientific examinations that they radioed back to Earth. Robots are being used to assist surgeons in installing artificial hips, and very high-precision robots can assist surgeons with delicate operations on the human eye. Research in telesurgery uses robots that may one day perform operations in distant battlefields under the remote control of expert surgeons.

Remotely controlled robots are now used by the military. These include small terrestrial robots to disable bombs and flying unmanned aerial vehicles (UAVs) equipped with powerful cameras for reconnaissance. Versions of such robots are also designed to use deadly force in military combat operations. Ground robots with cameras can carry machine guns fired remotely by an operator. UAVs equipped with bombs or missiles can strike targets from the air. Experts have raised concerns about giving future combat robots the ability to use force without direct human control. Such robots could also be used by terrorists.

Impacts of Robots

Robotic manipulators create manufactured products that are of higher quality and lower cost. But robots can cause the loss of unskilled jobs, particularly on assembly lines in factories. New jobs are created in software and sensor development, in robot installation and maintenance, and in the conversion of old factories and the design of new ones. These new jobs, however, require higher levels of skill and training. Technologically oriented societies must face the task of retraining workers who lose jobs to automation, providing them with new skills so that they can be employable in the industries of the 21st century.

Future Technologies

Automated machines will increasingly assist humans in the manufacture of new products, the maintenance of the world's infrastructure, and the care of homes and businesses. Robots will be able to make new highways, construct steel frameworks of buildings, clean underground pipelines, and mow lawns. Prototypes of systems to perform all of these tasks already exist. One important trend is the development of microelectromechanical systems, ranging in size from centimeters to millimeters. These tiny robots may be used to move through blood vessels to

deliver medicine or clean arterial blockages. They also may work inside large machines to diagnose impending mechanical problems.

Perhaps the most dramatic changes in future robots will arise from their increasing ability to reason. The field of artificial intelligence is moving rapidly from university laboratories to practical application in industry, and machines are being developed that can perform cognitive tasks, such as strategic planning and learning from experience. Increasingly, diagnosis of failures in aircraft or satellites, the management of a battlefield, or the control of a large factory will be performed by intelligent computers.

Corona Virus in a Digital Fight

As the coronavirus emergency exploded into a full-blown pandemic in early 2020, forcing countless businesses to shutter, robot-making companies found themselves in an unusual situation: Many saw a surge in orders. Robots don't need masks, can be easily disinfected, and, of course, they don't get sick.

An army of automatons has since been deployed all over the world to help with the crisis: They are monitoring patients, sanitizing hospitals, making deliveries, and helping frontline medical workers reduce their exposure to the virus. Not all robots operate autonomously—many, in fact, require direct human supervision, and most are limited to simple, repetitive tasks. But robot makers say the experience they've gained during this trial-by-fire deployment will make their future machines smarter and more capable. These photos illustrate how robots are helping us fight this pandemic—and how they might be able to assist with the next one.

Robots can act as an interface between a doctor and a patient wherein they can carry out diagnostic and treatment processes, reducing the human contact and risk of transmission of infection during the coronavirus pandemic, an expert in the field of Robotics has said. Bartłomiej Stanczyk, Robotics Engineer with ACCREA Engineering in Germany, was speaking during an e-discussion on the the topic- Using Artificial Intelligence to Tackle Epidemics: The COVID-19 Model. The event, organised by the Abu Dhabi-based TRENDS Research & Advisory, brought together leading experts from around the world who deliberated on the importance of artificial intelligence, machine learning, big data, and other technologies in the ongoing fight against the COVID-19 that has infected more than 3.8 million people and killed over 260,000 people across the world.

Stanczyk said that robots could help doctors keep a safe distance from the patient by using probes and other remote medical equipment. "We aim to build a completely autonomous diagnostician through robotics, thus enabling the transfer of the skill from the human doctor on the machine carrying out the treatment," he said. The interface between the doctor and patient means the robot can carry out all of the diagnostic and treatment processes, he said.

Explaining a wide range of uses of robots in the medical field, Stanczyk said that they can help in disinfection of inaccessible areas in hospitals. They can also be used in close proximity to humans by installing a sense of touch based on force sensors. Munier Nazzal, Professor of Surgery at the University of Toledo, in the US advocated the use of artificial intelligence (AI) in the development of a vaccine to cure COVID-19 patients. “AI can help with vaccine development by examining the virus' components. This can aid specialists gain a basic understanding and develop treatments that can be subject to pre-clinical trials,” he said.

Konrad Karcz, Professor of Medicine and Head of Minimally Invasive Surgery at the Ludwig Maximilian University Clinic in Germany, spoke about the potential for chatbots to measure body temperature and other medical indicators in patients. Sapan S Desai, Chief Executive Officer of the Surgisphere Corporation in the US, explained the transformative potential of AI illustrated by the company's collection of data on 86,000 COVID-19 cases which was used to model outcomes that suggested healthcare resources would be severely strained.

Nurses and doctors at Circolo Hospital in Varese, in northern Italy—the country's hardest-hit region—use robots as their avatars, enabling them to check on their patients around the clock while minimizing exposure and conserving protective equipment. The robots, developed by Chinese firm Sanbot, are equipped with cameras and microphones and can also access patient data like blood oxygen levels. Telepresence robots, originally designed for offices, are becoming an invaluable tool for medical workers treating highly infectious diseases like COVID-19, reducing the risk that they'll contract the pathogen they're fighting against.



Robots can't replace real human interaction, of course, but they can help people feel more connected at a time when meetings and other social activities are mostly on hold.

In Ostend, Belgium, ZoraBots brought one of its waist-high robots, equipped with cameras, microphones, and a screen, to a nursing home, allowing residents like Jozef Gouwy to virtually communicate with loved ones despite a ban on in-person visits.



And while Japan's Chiba Zoological Park was temporarily closed due to the pandemic, the zoo used an autonomous robotic vehicle called RakuRo, equipped with 360-degree cameras, to offer virtual tours to children quarantined at home.



In the Social Areas

Offices, stores, and medical centers are adopting robots as enforcers of a new coronavirus code. At Fortis Hospital in Bangalore, India, a robot called Mitra uses a thermal camera to perform a preliminary screening of patients.



In Tunisia, the police use a tanklike robot to patrol the streets of its capital city, Tunis, verifying that citizens have permission to go out during curfew hours.



Can Robot Be Moral?

The recent philosophical discussion concerning robots has been largely preoccupied with questions such as “can robots think, know, feel, or learn?” can they be conscious, technological, and self-adaptive?; can robot be in principle psychologically and intellectually isomorphic to men? Considerably less attention has been paid meanwhile to the question whether robots can be moral. Since the later problem seems to me rather intimately connected with the ones extensively discussed, I would like to raise it here in an attempt to carry the discussion to its logical conclusion.

The thesis of this paper is that there are no magic descriptive terms intelligence, consciousness, purposiveness, etc. predicable exclusively of men but not of robots, then there are no such moral terms either. If men and machines coexist in a natural continuum in which there are no gaps, quantum jumps, or insurmountable barriers preventing the assimilation of the one to the other, then they also coexist in a moral continuum in which only relative but never absolute distinctions can be made between human and machine morality.

I will argue the thesis by raising the question whether robots can be moral in two stages:

1. Can robots act morally?
2. Can we, without absurdity, treat robots agents?

The answer to these questions will be given, not in terms of a new “robot morality” but in terms of a few traditional ethical theories.

To make these questions more sophisticated than any single machine already existing. At the time, for all their complexity, they are not to have any capabilities other than the one computer scientists and cyberneticists like Turing, Wiener, Ashby, Arbib, Pask, and Uttley.

Machines cannot be assumed to be inherently capable of behaving morally. Humans must teach them what morality is, how it can be measured and optimized. For AI engineers, this may seem like a daunting task. After all, defining moral values is a challenge mankind has struggled with throughout its history. Nevertheless, the state of AI research requires engineers and ethicists to define morality and quantify it in explicit terms. Engineers cannot build a “Good Samaritan AI” as long as they lack a formula for the Good Samaritan human.

Robotic Ethics

The ethics of artificial intelligence (AI) and robotics is often focused on “concerns” of various sorts, which is a typical response to new technologies. Many

such concerns turn out to be rather quaint (trains are too fast for souls); some are predictably wrong when they suggest that the technology will fundamentally change humans (telephones will destroy personal communication, writing will destroy memory, video cassettes will make going out redundant); some are broadly correct but moderately relevant (digital technology will destroy industries that make photographic film, cassette tapes, or vinyl records); but some are broadly correct and deeply relevant (cars will kill children and fundamentally change the landscape). The task of an article such as this is to analyse the issues and to deflate the non-issues. Some technologies, like nuclear power, cars, or plastics, have caused ethical and political discussion and significant policy efforts to control the trajectory these technologies, usually only once some damage is done. In addition to such “ethical concerns”, new technologies challenge current norms and conceptual systems, which is of particular interest to philosophy. Finally, once we have understood a technology in its context, we need to shape our societal response, including regulation and law. All these features also exist in the case of new AI and Robotics technologies—plus the more fundamental fear that they may end the era of human control on Earth.

The ethics of AI and robotics has seen significant press coverage in recent years, which supports related research, but also may end up undermining it: the press often talks as if the issues under discussion were just predictions of what future technology will bring, and as though we already know what would be most ethical and how to achieve that. Press coverage thus focuses on risk, security (Brundage et al. 2018). The result is a discussion of essentially technical problems that focus on how to achieve a desired outcome. Current discussions in policy and industry are also motivated by image and public relations, where the label “ethical” is really not much more than the new “green”, perhaps used for “ethics washing”. For a problem to qualify as a problem for AI ethics would require that we do *not* readily know what the right thing to do is. In this sense, job loss, theft, or killing with AI is not a problem in ethics, but whether these are permissible under certain circumstances *is* a problem. This article focuses on the genuine problems of ethics where we do not readily know what the answers are.

A last caveat: The ethics of AI and robotics is a very young field within applied ethics, with significant dynamics, but few well-established issues and no authoritative overviews—though there is a promising outline (European Group on Ethics in Science and New Technologies 2018) and there are beginnings on societal impact (Floridi et al. 2018; Taddeo and Floridi 2018; S. Taylor et al. 2018; Walsh 2018; Bryson 2019; Gibert 2019; Whittlestone et al. 2019), and policy recommendations (AI HLEG 2019 [OIR]; IEEE 2019). So this article cannot merely reproduce what the community has achieved thus far, but must propose an ordering where little order exists.

AI & Robotics

The notion of “artificial intelligence” (AI) is understood broadly as any kind of artificial computational system that shows intelligent behaviour, i.e., complex

behaviour that is conducive to reaching goals. In particular, we do not wish to restrict “intelligence” to what would require intelligence if done by *humans*, as Minsky had suggested (1985). This means we incorporate a range of machines, including those in “technical AI”, that show only limited abilities in learning or reasoning but excel at the automation of particular tasks, as well as machines in “general AI” that aim to create a generally intelligent agent. AI somehow gets closer to our skin than other technologies—thus the field of “philosophy of AI”. Perhaps this is because the project of AI is to create machines that have a feature central to how we humans see ourselves, namely as feeling, thinking, intelligent beings. The main purposes of an artificially intelligent agent probably involve sensing, modelling, planning and action, but current AI applications also include perception, text analysis, natural language processing (NLP), logical reasoning, game-playing, decision support systems, data analytics, predictive analytics, as well as autonomous vehicles and other forms of robotics (P. Stone et al. 2016). AI may involve any number of computational techniques to achieve these aims, be that classical symbol-manipulating AI, inspired by natural cognition, or machine learning via neural networks (Goodfellow, Bengio, and Courville 2016; Silver et al. 2018).

Historically, it is worth noting that the term “AI” was used as above ca. 1950–1975, then came into disrepute during the “AI winter”, ca. 1975–1995, and narrowed. As a result, areas such as “machine learning”, “natural language processing” and “data science” were often not labelled as “AI”. Since ca. 2010, the use has broadened again, and at times almost all of computer science and even high-tech is lumped under “AI”. Now it is a name to be proud of, a booming industry with massive capital investment (Shoham et al. 2018), and on the edge of hype again. As Erik Brynjolfsson noted, it may allow us to virtually eliminate global poverty, massively reduce disease and provide better education to almost everyone on the planet. (quoted in Anderson, Rainie, and Luchsinger 2018)

While AI can be entirely software, robots are physical machines that move. Robots are subject to physical impact, typically through “sensors”, and they exert physical force onto the world, typically through “actuators”, like a gripper or a turning wheel. Accordingly, autonomous cars or planes are robots, and only a minuscule portion of robots is “humanoid” (human-shaped), like in the movies. Some robots use AI, and some do not: Typical industrial robots blindly follow completely defined scripts with minimal sensory input and no learning or reasoning (around 500,000 such new industrial robots are installed each year (IFR 2019 [OIR])). It is probably fair to say that while robotics systems cause more concerns in the general public, AI systems are more likely to have a greater impact on humanity. Also, AI or robotics systems for a narrow set of tasks are less likely to cause new issues than systems that are more flexible and autonomous.

Robotics and AI can thus be seen as covering two overlapping sets of systems: systems that are only AI, systems that are only robotics, and systems that are both.

We are interested in all three; the scope of this article is thus not only the intersection, but the union, of both sets.

The Future of Artificial Intelligence



Humanoid Robot ASIMO Walks down Stairs

Building intelligent systems—and ultimately, automating intelligence—remains a daunting task, and one that may take decades to fully realize. AI research is currently focused on addressing existing shortcomings, such as the ability of AI systems to converse in natural language and to perceive and respond to their environment. However, the search for AI has grown into a field with far-reaching applications, many of which are considered indispensable and are already taken for granted. Nearly all industrial, governmental, and consumer applications are likely to utilize AI capabilities in the future.

Conclusion

Before COVID-19, most people had some degree of apprehension about robots and artificial intelligence. Though their beliefs may have been initially shaped by dystopian depictions of the technology in science fiction, their discomfort was reinforced by legitimate concerns. Some of AI's business applications were indeed leading to the loss of jobs, the reinforcement of biases, and infringements on data privacy.

Those worries appear to have been set aside since the onset of the pandemic as AI-infused technologies have been employed to mitigate the spread of the virus. We've seen an acceleration of the use of robotics to do the jobs of humans who have been ordered to stay at home or who have been redeployed within the workplace. Labor-replacing robots, for example, are taking over floor cleaning in grocery stores and sorting at recycling centers. AI is also fostering an increased reliance on chatbots for customer service at companies such as PayPal and on machine-driven content monitoring on platforms such as YouTube. Robotic telepresence platforms are providing students in Japan with an “in-person” college

graduation experience. Robots are even serving as noisy fans in otherwise empty stadiums during baseball games in Taiwan. In terms of data, AI is already showing potential in early attempts to monitor infection rates and contact tracing.

After a vaccine for COVID-19 is developed (we hope) and the pandemic retreats, it's hard to imagine life returning to how it was at the start of 2020. Our experiences in the coming months will make it quite easy to normalize automation as a part of our daily lives. Companies that have adopted robots during the crisis might think that a significant percentage of their human employees are not needed anymore. Consumers who will have spent more time than ever interacting with robots might become accustomed to that type of interaction. When you get used to having food delivered by a robot, you eventually might not even notice the disappearance of a job that was once held by a human. In fact, some people might want to maintain social distancing even when it is not strictly needed anymore.

We, as a society, have so far not questioned what types of functions these robots will replace — because during this pandemic, the technology is serving an important role. If these machines help preserve our health and well-being, then our trust in them will increase. As the time we spend with people outside of our closest personal and work-related social networks diminishes, our bonds to our local communities might start to weaken. With that, our concerns about the consequences of robots and AI may decrease. In addition to losing sight of the scale of job loss empowered by the use of robots and AI, we may hastily overlook the forms of bias embedded within AI and the invasiveness of the technology that will be used to track the coronavirus's spread.

***Bartholomew, Uchenna Arum MSc, PDE, HND**

Email: Uchenna.bartholomew84@gmail.com

References

Allen, Colin, Iva Smit, and Wendell Wallach, 2005, "Artificial Morality: Top-down, Bottom-up, and Hybrid Approaches", *Ethics and Information Technology*, 7(3): 149–155. doi:10.1007/s10676-006-0004-4

Allen, Colin, Gary Varner, and Jason Zinser, 2000, "Prolegomena to Any Future Artificial Moral Agent", *Journal of Experimental & Theoretical Artificial Intelligence*, 12(3): 251–261. doi:10.1080/09528130050111428

Amoroso, Daniele and Guglielmo Tamburrini, 2018, "The Ethical and Legal Case Against Autonomy in Weapons Systems", *Global Jurist*, 18(1): art. 20170012. doi:10.1515/gj-2017-0012

Anderson, Janna, Lee Rainie, and Alex Luchsinger, 2018, *Artificial Intelligence and the Future of Humans*, Washington, DC: Pew Research Center.

Arnold, Thomas and Matthias Scheutz, 2017, “Beyond Moral Dilemmas: Exploring the Ethical Landscape in HRI”, in *Proceedings of the 2017 ACM/IEEE International Conference on Human-Robot Interaction—HRI '17*, Vienna, Austria: ACM Press, 445–452. doi:10.1145/2909824.3020255

Asaro, Peter M., 2019, “AI Ethics in Predictive Policing: From Models of Threat to an Ethics of Care”, *IEEE Technology and Society Magazine*, 38(2): 40–53. doi:10.1109/MTS.2019.2915154

Asimov, Isaac, 1942, “Runaround: A Short Story”, *Astounding Science Fiction*, March 1942. Reprinted in “I, Robot”, New York: Gnome Press 1950, 1940ff.

Awad, Edmond, Sohan Dsouza, Richard Kim, Jonathan Schulz, Joseph Henrich, Azim Shariff, Jean-François Bonnefon, and Iyad Rahwan, 2018, “The Moral Machine Experiment”, *Nature*, 563(7729): 59–64. doi:10.1038/s41586-018-0637-6

Baldwin, Richard, 2019, *The Globotics Upheaval: Globalisation, Robotics and the Future of Work*, New York: Oxford University Press.

Bennett, Colin J. and Charles Raab, 2006, *The Governance of Privacy: Policy Instruments in Global Perspective*, second edition, Cambridge, MA: MIT Press.

Benthall, Sebastian and Bruce D. Haynes, 2019, “Racial Categories in Machine Learning”, in *Proceedings of the Conference on Fairness, Accountability, and Transparency - FAT* '19*, Atlanta, GA, USA: ACM Press, 289–298. doi:10.1145/3287560.3287575