



Artificial Intelligence in 2027

Maria Gini (University of Minnesota; gini@umn.edu)

Noa Agmon (Bar-Ilan University; agmon@cs.biu.ac.il)

Fausto Giunchiglia (University of Trento; fausto.giunchiglia@unitn.it)

Sven Koenig (University of Southern California; skoenig@usc.edu)

Kevin Leyton-Brown (University of British Columbia; kevinlb@cs.ubc.ca)

DOI: [10.1145/3203247.3203251](https://doi.org/10.1145/3203247.3203251)

Introduction

Every day we read in the scientific and popular press about advances in AI and how AI is changing our lives. Things are moving at a fast pace, with no obvious end in sight.

What will AI be ten years from now? A technology so pervasive in our daily lives that we will no longer think about it? A dream that has failed to materialize? A mix of successes and failures still far from achieving its promises?

At the 2017 International Joint Conference on Artificial Intelligence (IJCAI), Maria Gini chaired a panel to discuss “AI in 2027.” There were four panelists: Noa Agmon (Bar-Ilan University, Israel), Fausto Giunchiglia (University of Trento, Italy), Sven Koenig (University of Southern California, US), and Kevin Leyton-Brown (University of British Columbia, Canada). Each of the panelists specializes in a different part of AI, so their visions span the field, providing an exploration of possible futures.

The panelists were asked to present their views on possible futures, specifically addressing what AI technologies they expected would be in widespread use in 2027, what they thought would still show potential but not have become widely accepted, and what they expected the AI research landscape to look like ten years from now.

This article summarizes the main points that each panelist made and their reflections on the topics. The focus in each contribution is not much on predicting the future but on bringing up specific open problems in each subarea and discuss how the current AI technologies could be steered to address them.

Noa Agmon, Bar-Ilan University¹

The discussion about the fourth industrial revolution, and the part of AI and robotics within it, is wide. In the context of this revolution, autonomous cars and other types of robots are expected to gain popularity and, among other things, to take over human labor. While surveys like “When will AI exceed human performance?” (GSD⁺17) report that some researchers expect robots to be capable of performing human tasks, such as running five kilometers, within ten years, this will probably take much longer given the current state of robotic development.

Today, the use of robots is generally limited to three categories: non-critical tasks; settings in which robots are semi-autonomous, tele-operated, or remote-controlled (namely, not fully autonomous); and highly structured settings in which uncertainties are minimal. Examples of such settings include the Amazon Robotics warehouse robots, which work autonomously in a structured environment (the warehouse), semi-autonomous drones operated in military settings (usually follow a specified route autonomously, though operative decisions are made by human operators), robots that perform cleaning tasks, which are considered non-critical, Mars rovers, which operate semi-autonomously in unstructured environments, and more. When robots are required to operate fully autonomously in unstructured settings requiring them to handle unbounded uncertainties or completely unpredictable events, they tend to fail. One of many examples is the Knightscope robot, which drove into a fountain on its first day of deployment as a security guard in Washington, D.C.

Rather than arguing about the ability of robots

¹Acknowledgments: I would like to thank Gal Kaminka and David Sarne from Bar-Ilan University for their helpful comments.

to outperform humans, and when this might happen, the following discussion examines the challenges and opportunities that will influence the development of intelligent robotics in the next ten years.

Dependence on hardware. As opposed to the progress of AI, which relies mainly on algorithmic development and benefits from processing improvements, progress in robotics is also intimately tied to the capabilities of electro-mechanics, physical sensors, and energy storage and management. Whatever apocalyptic or euphoric visions we have for working with robots, their realization is much more dependent on physical components than we, AI researchers and practitioners, tend to consider. For example, most quad-copters, which are considered to be a basis for breakthrough applications (such as home deliveries and emergency services), can only fly for 30 minutes or so. Likewise, vacuum cleaners are limited in the total area that they can cover before they have to be recharged. These energy concerns radically impact the usefulness of robots in applications which are otherwise within reach from a pure software perspective.

The good news is that the intimate connection between software and hardware works both ways. Just as modern SLAM algorithms (e.g., (DNC⁺01)) were able to overcome intrinsic sensor limitations to create reliable and accurate maps for navigation, advances in software can overcome some of the limitations posed by hardware.

AI influences on robotics. AI algorithms influence robotics not only in compensating for and improving the utilization of existing hardware capabilities, but also in enabling new tasks. Progress in natural language processing (NLP) and machine learning (used for chatbots, personal digital assistants, and surveillance, for instance) enables more natural forms of human-robot interaction with physical robots, and autonomous cars. However, such positive influences are somewhat asymmetric: AI will influence robotics more than robotics will influence AI. A personal robot benefits from NLP more than NLP can benefit from the consideration of multi-modal interactions (as in “talking with your hands.”)

Growing role for multi-robot systems (MRS). The academic research on MRS dates

back to the early 1980’s, when robots were scarce and not autonomous. Research has progressed far beyond the deployment of such systems outside of labs. Improvements in the reliability of robots will make it easier to deploy MRS in various applications, continuing and accelerating current successful trends (e.g., in warehouses and hospitals). This, in turn, will accelerate research on fully distributed, fully autonomous systems, which are beyond current capabilities. It is obvious that human-robot interactions will be a major focus of research in the next ten years, as robots enter a greater number of unstructured environments in which humans operate. However, given the foreseen growth in the role of multi-robot systems, human-MRS and multiple operator-single robot collaborations will likely see increased efforts.

Increasing ties with other disciplines. A good example of large-scale fully distributed, fully autonomous systems also raises an additional trend that of increasing ties with other disciplines. Swarms of molecular robots (nanobots), the size of which is measured in nanometers, are becoming a reality in medical applications (for example, targeted drug delivery). Trillions of such robots will be let loose in a patient’s body - the largest-scale MRS in robotics history. The computation of interactions between different types of nanobots has an immense impact on the ease and duration of development of new treatments (WKKH⁺16; KSSA⁺17). The capability to plan and reason about the interactions of these robots with each other and with the body requires deep collaboration between AI experts, biologists, and chemists.

Another example is reconfigurable robots, which can transform themselves into different shapes, depending on the environment and task. Research on such systems will benefit from close coordination with chemistry, physics, and biology, to take new findings into account.

Increasing accessibility, lower entry barrier, greater impact potential. A positive trend which I believe will continue to grow is the lowering of the entry barrier into robotics practice and research, at multiple levels. Research-grade robots for labs have seen dramatic decreases in cost, and the common availability of 3D printing, cheap embedded computers (Arduino, Raspberry Pi), as well as continued

push on STEM education will make development of robots cheaper and easier than ever. The availability of common robot software middleware, such as ROS, make it easier for researchers to focus their attention on bringing their expertise to bear on specific components.

Bottom line: More of the same (which is good!) Within the next ten years and beyond, we will not see general-purpose robots. That is, robots will still be dedicated to one task, for example delivery, cleaning, or surveillance. Progress in the development of intelligent robotic systems will continue to focus on excellence in the performance of specific tasks, and on the introduction of new tasks to new types of robots. To some extent, this will make robot use more popular.

Fausto Giunchiglia, University of Trento²

Providing machines with knowledge, e.g., common sense or domain knowledge, has always been one of the core AI issues. Two are the main approaches to this problem. The first *deductive* approach, usually categorized under the general heading of “Knowledge Representation and Reasoning (KRR)”, dates back to John McCarthy’s *advice taker* proposal (McC60), and it is based on the idea of telling machines what is the case, for instance in the form of facts codified as logical axioms. The second *inductive* approach, usually categorized under the general heading of “Machine Learning (ML)”, consists of providing machines with a set of examples from which to learn general statements, usually with a certain level of confidence.

A lot of relevant research has been done in KRR, not least the work on the Semantic Web (BLHL01), and much more will be done. However, the success of ML mainly, but not only, because of the work in Deep Learning (see, e.g., (LBH15)), has been so overwhelming that, thinking of what the research in KRR could be in the next ten years and where it could lead, a relevant question is the extent to which these two lines of work should integrate in an effort to jointly produce results that either of them alone

could not produce.

A lot of successful work in this area has been done, see for instance (GT07; RKN16). However, the extent to which the KRR research could be improved by exploiting the research developed in ML, or dually, the extent to which ML would *really* need to exploit any of the results developed in KRR is still unclear, at least for two reasons. The first is that this integration is far from being trivial. It is a fact that these two approaches start from somewhat opposite assumptions, the first assuming that knowledge consists of a set of facts which are either true or false, with nothing in between, the second having to deal with the issue that any fact learned via ML will hardly ever be guaranteed to be true or false with an infinite number of intermediate levels. Furthermore, the need for such an integration is far from being clear, at least from an ML point of view. Among other things, it is a fact that the current ML techniques have proven so powerful that, whenever applicable, they seem to be able to learn virtually unbound amounts of knowledge, far more knowledge than could be codified by any knowledge engineer.

At the same time both the KRR and ML techniques have their own weaknesses. Thus, on one side, KRR presents a main difficulty in how to express the inherent complexity and variability of the world, in particular but not only, when perception is involved. On the other side, instead, ML presents a main difficulty in making sense, in human terms, of the knowledge which is learned. In other words, the knowledge generated via ML does not often fit the people “intended semantics”, namely how they would describe what is the case, for instance as perceived or as learned from a large amount of text messages.

This difficulty of ML techniques, and data driven approaches in general, has been known for many years. An explicit reference to this problem comes from the field of Computer Vision, where it is named the *Semantic Gap Problem (SGP)*. The SGP was originally defined in (SWS⁺00) as follows: “... The semantic gap is the lack of coincidence between the information that one can extract from the visual data and the interpretation that the same data have for a user in a given situation. ...” It can be noticed how this notion is completely gen-

²Acknowledgments: I would like to thank Kobi Gal, Daniel Gatica-Perez, Loizos Michael, Daniele Miorandi, Andrea Passerini and Carles Sierra for our many useful discussions on this topic.

eral and can be taken to refer to the human-machine misalignment which may arise with any type of information that a machine can extract, i.e., learn, from any type of data. The novelty of these last years is that many more instances of the SGP are showing up and many of them are also discussed in the news, the main reason being the increased use, increased power, and increased popularity of ML systems and AI in general. Thus, we have read of cases when a system learns biased opinions, or when it learns a language that it is not human, or when an autonomous car does not track another car thus causing an accident. And this, in turn, is the cause of a lot of public discussions about the interaction between AI and humans, about AI and ethics and also of an increased fear of AI.

A convincing explanation of how to deal with the SGP seems necessary for AI to be used beyond a set of niche (possibly very large) application areas and to be adopted by the general public. I believe it will be very hard to convince people to use machines that they do not feel they fully control, in high value application domains, e.g., health, mobility, energy, retail. But a convincing explanation will not be enough. A solution of the SGP is also needed for AI to be used in practice. There are at least three mainstream application scenarios where some solution to the SGP seems crucial. The first is the *anytime anywhere delivery of personalized services*, as enabled by personal digital devices, e.g., smart phones or smart watches. But for this to happen, people will have to be able to make sense of why certain decisions have been taken by the machine, and to agree with them. The second is the *empowerment of social relations*, exactly for the same reasons mentioned above. Facebook, Whatsapp, or Snapchat are just the beginning and I foresee the rise of a new generation of social networks empowering more specialized, more personalized, more diversity-aware interactions among people. The third, and maybe the most important, again because of the pervasiveness of digital devices, is that we are more moving towards *open world* application scenarios. By this I mean application scenarios where, at design time, it is impossible to anticipate the system functional and non-functional requirements. In this type of applications the effects of the SGP can be devastating, as the diver-

gence between people and machines can only get worse in time.

In my opinion, a general solution of the SGP problem, and in particular a solution which is viable in open world application scenarios, can only be achieved via a tight integration of knowledge-based approaches and machine learning. The knocking down argument is that the only way to avoid the SGP is to make sure that machines learn representations of the world which are the same as their reference users. But, the fact that knowledge should be presented in human-like terms is exactly the assumption underlying all the work in KRR and also logic. More specifically, whatever knowledge will be learned via a data-driven approach, it will have to be compared and ultimately aligned to the human knowledge. Someone could argue that this is exactly what supervised learning does. But this is not the case, as also witnessed by the fact that even the human supervision, how it is implemented up to now, does not make the SGP disappear. The problem is far more complex and it will require major advances in both KRR and ML, and in AI in general, many of which, I believe, will be disruptive. A list of four open issues is provided below, with the understanding that this list is not meant to be complete nor correct. This list reflects only my current personal understanding of some of the problems which will have to be dealt with when trying to solve the SGP.

1. Since the early days of AI a fundamental issue has been that of building machines which would *exceed human-level intelligence*. This goal has been reached in many domains, e.g. chess or GO playing, while it is very far from being reached in other domains, e.g., robotics, as mentioned above. A solution of the SGP will require building machines which will show *human-like intelligence*, representation and reasoning, as the basis for the mutual human-machine understanding. In this context, exceeding human-level intelligence seems a desired property but not strictly necessary.
2. A fundamental property of life, and of humans in particular, is their ability to *adapt to unpredicted events and evolve*. Both the research in KRR and in ML seems very far from achieving this goal. It is however interesting to notice how a particular instance

of this inability to adapt was recognized by John McCarthy and named *the problem of lack of generality* of the current representation formalisms (McC87).

3. The net result of the ability to adapt and evolve as a function of the local context will be that the resulting knowledge will be highly diversified. In turn, the *diversity of knowledge* will generate the need for further adaptation in an infinite loop with will result in the process of *knowledge evolution*, somewhat analogously to the kind of evolution we see in life. Notice how the proposed approach is quite different from that taken by the Semantic Web for the solution of the problem of semantic heterogeneity. The focus is not on representation tools, e.g., ontologies, or formalisms, e.g., Description Logics, but on the process by which knowledge gets generated, stored, manipulated, and used. In this perspective the standard logics, e.g., monotonic non-monotonic logics, seem to solve, at most, only part of the problem, and for sure not the most important.
4. A specific, but core, subproblem of the problem of managing knowledge diversity, is the integration of the knowledge obtained via perception, e.g., via computer vision, and the knowledge obtained via reasoning or by being told. An implicit assumption which has been made so far is that the linguistic representation of an object we talk about, e.g., the word “cat”, and the representation of what we perceive as a cat is one-to-one. As discussed in detail in (GF16) this in general is not the case and there is a many-to-many mapping between linguistic representations and perceptual representations. On top of this, these mappings are highly dependent on the culture and on the single person and, even for the same person, change in time, as a function of the person current interests. A full understanding of how these mappings are built, of how linguistic representations influence the construction of perceptual representations, and vice versa, is a largely unexplored research area. Still, some form of solution to this problem will be needed in order to guarantee that the machine will describe what it will perceive coherently with what humans do.

Sven Koenig, University of Southern California³

AlphaGo (SHM⁺16) shows that it can be very difficult to judge technical progress, as also noticed by Stuart Russell in his invited IJCAI-17 talk. When it beat Lee Sedol in 2016, many experts thought that such a win was still at least a decade away. The AI techniques behind it already existed in principle. The ingenuity was in figuring out how to put them together in the right way. Progress on AI technology is often steadier than it appears, yet such engineering breakthroughs happen only from time to time, are difficult to predict, and often make AI technology visible in the public eye - creating the perception of waves of progress.

Various recent studies shed light on the expected progress of AI by 2027, such as the study on “AI and Life in 2030” as part of the One Hundred Year Study on AI (ai100.stanford.edu) and a recent survey of all ICML-15 and NIPS-15 authors (GSD⁺17). This survey, for example, predicts that AI will outperform humans around 2027 on tasks such as writing high school essays, explaining actions in games, generating top 40 pop songs, and driving trucks. Furthermore, humanoid robots will soon afterward beat humans in 5k races. Interestingly, North American researchers predicted that it will take about 74 years to reach high-level machine intelligence across human tasks, while Asian researchers thought it would take only 30 years. Indeed, there is currently lots of excitement and optimism, for example, in China about the potential of AI with large investments into application-oriented AI research by both the government and private sector.

In the following, I view AI as the study of agents to structure the discussion which kinds of research topics will be popular in 2027. I distinguish rational agents (that make good decisions), believable agents (that interact like humans), and cognitive agents (that think like humans). A large amount of AI research currently focuses on building rational agents on the task level - by studying single AI techniques in isolation and applying them to single tasks, resulting in narrowly intelligent agents. The current excitement about AI is often based on

³Acknowledgments: I would like to thank Paul Rosenbloom and Wolfgang Hönig from the University of Southern California for helpful comments.

the power of a small set of AI techniques combined with the availability of large amounts of data (due to progress on sensor technologies and the ubiquity of both smart phones and the internet) as well as progress in robotics for the embodiment of AI. For example, the term “big data” is typically used to characterize the current AI era, driven by the capability of machine learning techniques. In fact, 49 percent of submissions to the International Joint Conference on AI (IJCAI) in 2017 used as first keyword “machine learning,” which is about acquiring good models of the world. However, these models need to serve a bigger purpose, for example, to make good decisions. While machine learning can sometimes acquire evaluation functions that help with making good decisions (as AlphaZero (SHS⁺17), the successor of AlphaGo shows), it often requires lots of data, has limited capability for transfer, and has difficulty integrating prior knowledge (Mar16) and thus is of limited help for making decisions in novel or dynamic environments. Perhaps the term “big decisions” will be used to characterize the AI era around 2027, driven also by the capability of AI planning and similar AI techniques. Current faculty hiring in the US lags in research areas such as AI planning although the research community is already heading in that direction. For example, the popular textbook by Stuart Russell and Peter Norvig (RN09) views AI as the study of rational agents and thus essentially as a science of making good decisions with respect to given objectives. But many other disciplines could be characterized similarly, including operations research, decision theory, economics, and control theory (Koe12). AI researchers make use of techniques from some of these disciplines already. For example, the textbook by Stuart Russell and Peter Norvig discusses utility theory (from decision theory), game theory and auctions (from economics), and Markov decision processes (from operations research), yet research collaborations across these and other disciplines are still developing, which is why we should reach out more to researchers in other decision-making disciplines. There already exist some good but narrow interfaces, such as the Conference on the Integration of Constraint Programming, AI, and Operations Research (CPAIOR) or the ACM Conference on Economics and Computation (EC). There also exists an attempt to put a broader interface

in place, namely the International Conference on Algorithmic Decision Theory (ADT), which “seeks to bring together researchers and practitioners coming from diverse areas such as AI, Database Systems, Operations Research, Discrete Mathematics, Theoretical Computer Science, Decision Theory, Game Theory, Multi-agent Systems, Computational Social Choice, Argumentation Theory, and Multiple Criteria Decision Aiding in order to improve the theory and practice of modern decision support” (sma.uni.lu/adt2017). Such interdisciplinary integration can result in economic success. CPAIOR, for example, started in 2004 (preceded by five workshops) and still thrives. In parallel, ILOG successfully integrated software for constraint programming and linear optimization and was acquired by IBM in 2009. My hope is that we will have a thriving conference on intelligent decision making by 2027 that will be attended by researchers from all decision-making disciplines, including AI. Of course, the different decision-making techniques also need to be integrated into systems. AI can lead the way by developing agent architectures with good theoretical foundations for how different parts should interact, resulting in more broadly intelligent agents on the job level (that is, across tasks). This is no simple feat as the restricted applications of current robot architectures show. Integrating decision-making techniques from different disciplines is even more difficult, for example, because of their different assumptions (often due to different application areas studied by different disciplines) and different ideas about what constitutes a good solution (due to disciplinary training), which is why we should start to give students multi-disciplinary training in decision making.

While rational agents will continue to be important, human-aware agents will become more and more important and, with them, also believable agents that allow for interaction with gestures, speech, and other human-like modalities, understand human conventions and emotions, predict human behavior, and - in general - appear to be human-like. We already use intelligent assistants on a variety of platforms (such as Apple’s Siri or Amazon’s Echo) and will soon routinely have conversations - including negotiations - with all kinds of apparatus, perhaps including our elevators and toilets ☺.

The progress on cognitive agents, one of the

early dreams of AI, is more difficult to judge. The research community currently works on hybrid approaches that combine ideas from symbolic, statistical, and/or neural processing and on a community-wide “Common Model of Cognition” (KT16).

Finally, AI researchers and practitioners slowly gain an understanding that they should not just develop AI techniques but also have some say in how they are being used. We need to ask ourselves questions such as:

Do we need to worry about the reliability, robustness, and safety of AI systems and, if so, what to do about it? How do we guarantee that their behavior is consistent with social norms and human values? Who is liable for incorrect AI decisions? How to ensure that AI technology impacts the standard of living, distribution and quality of work, and other social and economic aspects in the best possible way? (BGK⁺17)

AAAI and ACM recently co-founded the AAAI/ACM Conference on AI, Ethics, and Society (AIES, www.aies-conference.com) to come up with answers to these questions. AIES was filled to capacity. I hope that AIES and its topics will be even more popular in 2027.

Kevin Leyton-Brown, University of British Columbia⁴

It is a daunting task to predict the direction AI research will take a decade from now, particularly given the checkered history of such prognostication in the past. In an attempt to go beyond idle speculation, I have therefore structured this reflection around three different *approaches* a forecaster might use in making such predictions. Despite recognizing the likelihood that some of what follows will appear foolish in retrospect, I strive to make bold claims about what the future will hold. I hope that AI researchers of 2027 will forgive me!

I. Forecasting via prototypes. Ten years

⁴Acknowledgments: I would like to thank the members of the AI100 2015–16 Study Panel for helping to shape my thinking about the future of AI: P. Stone, R. Brooks, E. Brynjolfsson, R. Calo, O. Etzioni, G. Hager, J. Hirschberg, S. Kalyanakrishnan, E. Kamar, S. Kraus, D. Parkes, W. Press, A. Saxenian, J. Shah, M. Tambe, A. Teller (SBB⁺16).

sounds like a long time, but in fact it takes about that long for technologies to move from the lab to widespread practice: the transformative technologies of today existed in prototype form a decade ago. One approach to AI forecasting is thus to look at today’s prototypes and to imagine their more widespread deployment.

Broadly speaking, today’s AI prototypes offer tailored solutions for specific tasks rather than general intelligence. Some AI research topics that I expect to see making considerably broader social impact by 2027 include:

- Non-text input modalities (vision; speech)
- Consumer modeling (recommendation; marketing)
- Cloud services (translation; question answering; AI-mediated outsourcing)
- Transportation (automated trucking; some self-driving cars)
- Industrial robotics (factories; some drone applications)
- AI knowledge work (logistics planning; radiology; legal research; call centers)
- Policing & security (electronic fraud; cameras; predictive policing)

By considering where today’s prototypes have achieved less traction, it is also possible to forecast sectors in which AI technologies are less likely to take off quickly. Overall, these are often areas in which major entrenched regulatory regimes need to be navigated; where there exist substantial social or cultural barriers to the adoption of new technologies; and/or where broad impact would depend on nontrivial hardware breakthroughs. Many such sectors are the focus of concerted research today and are likely to remain important in the research landscape in 2027; however, I believe that they are less poised for short-term practical impact. Some key examples are childcare, healthcare, and eldercare; education; consumer robots beyond niche applications; and semantically rich language understanding.

II. Forecasting via consumer desires. A second strategy is to assume that investment, entrepreneurial energy, and industrial R&D will focus on meeting consumer needs that are already apparent today, and hence that these areas will see future breakthroughs.

Labor automation. A fundamental consumer need is for someone else to perform unpleasant, routine tasks. The promise of automating such tasks has been part of the AI story from the beginning (e.g., Shakey the robot delivering coffee in an office setting (Nil84)) and is increasingly becoming a reality (e.g., robot vacuum cleaners in the home; ordering books via Amazon Alexa). However, there is much scope for additional innovation in this space, centering on currently unaddressed tasks to which large numbers of people currently devote considerable time. Some potential examples are household cleaning, yard work, pet care, shopping, and food preparation. Some needs may be met by directly replacing human with robotic labor; others may be met via “gig economy” platforms that use AI on the back end to more efficiently allocate human labor; and still others may be met in entirely new ways, such as by combining AI-driven logistics platforms with centralized industrial processes (e.g., replacing supermarkets with apps, warehouses, and courier services).

Social connection. We are highly social creatures, and are willing to pay handsomely for technologies that help us to make and strengthen connections with others. Current instantiations of such technologies (e.g., social networks; remote work platforms; online dating) are highly valuable, but relatively primitive from an AI perspective, relying mainly on micro-blogging, direct messaging, user modeling, and newsfeed curation. There is scope for more AI mediation of social connection, reducing the frictions that prevent people from easily finding others to interact with in the moment and making those interactions richer.

Entertainment. Our research community’s focus on solving industrially or socially important problems sometimes may cause us to pay insufficient attention to AI’s potential for transforming the entertainment industry, which addresses another fundamental consumer need. Gaming is already bigger than Hollywood (Che17), but the future of AI in entertainment will go far beyond what we now see as computer games. Future AI entertainments will increasingly be interactive and multimodal, and will intersect with sectors we now see as distinct, such as fitness, learning, performing useful tasks, and spending quality time with friends. AI will also play an increasingly criti-

cal role in the creation, delivery, and personalization of traditional, broadcast entertainment such as TV.

Education. Education is poised to grow as a consumer sector (ROO18), both as workers respond to the need to reskill and as individuals with extra leisure time follow their passions. I argued above that education will not be transformed by AI in a decade; however, particularly because of the dual role many AI researchers hold as educators, there is nevertheless considerable scope for AI technology to make incremental progress in improving the content and delivery of educational materials. Some examples include tailoring lessons to a student’s skill level, making exercises more interactive, facilitating communication between both peers and instructors outside classroom settings, and reducing the drudgery currently entailed by grading student work. Such innovations could improve student outcomes, lower the cost of education, and broaden its reach.

III. Forecasting via extrapolation. A final strategy is to ask what we will be concerned with if current progress in AI continues.

AI beyond ML. Much recent progress in AI has arisen from improved techniques for learning to make predictions: finding a model that is currently built by hand and replacing it by a model that is learned from data (LBH15). We might therefore ask what problems would remain or become important if our capacity to build black-box models from data were to become arbitrarily effective. It is clear that even in such a world, we would be far from having achieved strong AI. Some problems that would remain open are still close to machine learning: explaining why a model made the prediction it did, or certifying fairness or compliance with legal requirements. Others, such as making counterfactual predictions (how would a system perform under a perturbation of the generating distribution?), go beyond the assumptions inherent in most supervised learning methods, requiring instead new, structural assumptions about an underlying setting. Still other problems extend beyond prediction to decision making, both in single-actor settings (e.g., optimization; planning) and multi-agent domains (weighing competing objectives via preference aggregation or mechanism design).

Increasing regulation. AI is touching the

lives of individuals, the economy, and the political system in ever increasing ways. Many in society find specific instantiations of AI frightening; many special interests are threatened by new technologies. It is thus inevitable that politicians will increasingly see a need to respond, and that AI technologies will face increasing regulation. This is something we should welcome; any mature technology must be accountable to the society in which it operates. But the details will matter enormously. A major focus of AI research in 2027 will be helping to shape regulations before they become law and designing systems within the constraints implied by these regulations afterwards.

Superhuman intelligence. AI systems will increasingly become capable of reaching human-level performance in a variety of application domains. There is nothing special about this threshold, and so we should expect the advent of AI systems exhibiting superhuman intelligence in a growing set of domains. This is often cast as a frightening prospect, but I argue that we will quickly become comfortable with it. After all, superhuman intelligences are already commonplace: governments, corporations, and NGOs are all autonomous agents that exhibit behavior much more sophisticated and complex than that of any human. We are typically unconcerned that no one person can even fully understand decisions made by the French government, by General Motors, or by the Red Cross. Instead, we aim to manage and to gain high-level understanding about such actors via reporting requirements, specifications of the interests that they must act to advance, and laws that forbid bad behavior. Society encourages the creation of such superhuman intelligences today for the same reason it will welcome superhuman AI tomorrow: many important problems are beyond the reach of individual people. Some key examples are improved collective decision making; more efficient allocation and use of scarce resources; addressing under-served communities; and limiting and responding to climate change.

References

E. Burton, J. Goldsmith, S. Koenig, B. Kuipers, N. Mattei, and T. Walsh. Ethical considerations in Artificial Intelligence courses. *Artificial Intelligence Magazine*, 38(2):22–34, 2017.

Tim Berners-Lee, James Hendler, and Ora Lassila. The semantic web. *Scientific American*, 284(5):34–43, 2001.

Chet Chetterson. How games overtook the movie industry, 2017. <http://www.geekadelphia.com/2017/06/29/how-games-overtook-the-movie-industry>, Accessed March 22, 2018.

M.W.M. Gamini Dissanayake, Paul Newman, Steve Clark, Hugh F Durrant-Whyte, and Michael Csorba. A solution to the simultaneous localization and map building (SLAM) problem. *IEEE Transactions on Robotics and Automation*, 17(3):229–241, 2001.

Fausto Giunchiglia and Mattia Fumagalli. Concepts as (recognition) abilities. In *Formal Ontology in Information Systems: Proc. 9th Int. Conference (FOIS 2016)*, pages 153–166, 2016.

K. Grace, J. Salvatier, A. Dafoe, B. Zhang, and O. Evans. When will AI exceed human performance? Evidence from AI experts. ArXiv Preprint arXiv:1705.08807v2, 2017.

Lise Getoor and Ben Taskar. *Introduction to Statistical Relational Learning (Adaptive Computation and Machine Learning)*. The MIT Press, 2007.

S. Koenig. Making good decisions quickly. *The IEEE Intelligent Informatics Bulletin*, 13(1):14–20, 2012.

Gal A. Kaminka, Rachel Spokoini-Stern, Yaniv Amir, Noa Agmon, and Ido Bachelet. Molecular robots obeying Asimov’s three laws of robotics. *Artificial Life*, 23(3):343–350, 2017.

I. Kotseruba and J. Tsotsos. A review of 40 years of cognitive architecture research: Core cognitive abilities and practical applications. ArXiv Preprint arXiv:1610.08602, 2016.

Yann LeCun, Yoshua Bengio, and Geoffrey Hinton. Deep learning. *Nature*, 521(7553):436, 2015.

G. Marcus. Deep learning: A critical appraisal. ArXiv Preprint arXiv:1801.00631, 2016.

John McCarthy. *Programs with common sense*. RLE and MIT Computation Center, 1960.

John McCarthy. Generality in Artificial Intelligence. *Communications of the ACM*, 30(12):1030–1035, 1987.

Nils J Nilsson. Shakey the robot. Technical report, SRI International, Menlo Park, CA, 1984.

Luc De Raedt, Kristian Kersting, and Sriraam Natarajan. *Statistical Relational Artificial Intel-*

ligence: Logic, Probability, and Computation. Morgan & Claypool Publishers, 2016.

S. Russell and P. Norvig. *Artificial Intelligence: A Modern Approach*. Pearson, 2009.

Max Roser and Esteban Ortiz-Ospina. Financing education, 2018. <https://ourworldindata.org/financing-education>, Accessed March 22, 2018.

Peter Stone, Rodney Brooks, Erik Brynjolfsson, Ryan Calo, Oren Etzioni, Greg Hager, Julia Hirschberg, Shivaram Kalyan Krishnan, Ece Kamar, Sarit Kraus, Kevin Leyton-Brown, David Parkes, William Press, AnnaLee Saxenian, Julie Shah, Milind Tambe, and Astro Teller. *Artificial intelligence and life in 2030. One Hundred Year Study on Artificial Intelligence: Report of the 2015-2016 Study Panel*, 2016.

D. Silver, A. Huang, C. Maddison, A. Guez, L. Sifre, G. van den Driessche, J. Schrittwieser, I. Antonoglou, V. Panneershelvam, M. Lanctot, S. Dieleman, D. Grewe, J. Nham, N. Kalchbrenner, I. Sutskever, T. Lillicrap, M. Leach, K. Kavukcuoglu, T. Graepel, and D. Hassabis. Mastering the game of Go with deep neural networks and tree search. *Nature*, 529:484–489, 2016.

D. Silver, T. Hubert, J. Schrittwieser, I. Antonoglou, M. Lai, A. Guez, M. Lanctot, L. Sifre, D. Kumaran, T. Graepel, T. Lillicrap, K. Simonyan, and D. Hassabis. Mastering Chess and Shogi by self-play with a general reinforcement learning algorithm. ArXiv Preprint arXiv:1712.01815, 2017.

Arnold WM Smeulders, Marcel Worring, Simone Santini, Amarnath Gupta, and Ramesh Jain. Content-based image retrieval at the end of the early years. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 22(12):1349–1380, 2000.

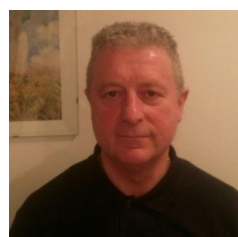
Inbal Wiesel-Kapah, Gal A. Kaminka, Guy Hachmon, Noa Agmon, and Ido Bachelet. Rule-based programming of molecular robot swarms for biomedical applications. In *Proc. International Joint Conference on Artificial Intelligence (IJCAI)*, pages 3505–3512, 2016.



Maria Gini is a Professor in the Department of Computer Science and Engineering at the University of Minnesota. She studies decision making for autonomous agents for distributed task allocation for robots, robotic exploration, teamwork for search and rescue, and navigation in dense crowds. She is a AAAI and a IEEE Fellow, and current president of IFAAMAS. She is Editor-in-Chief of Robotics and Autonomous Systems, and is on the editorial board of numerous journals, including Artificial Intelligence, Autonomous Agents and Multi-Agent Systems, and Integrated Computer-Aided Engineering.



Noa Agmon is an Assistant Professor in the Department of Computer Science at Bar-Ilan University, Israel. Her research focuses on various aspects of single and multi robot systems, including multi-robot patrolling, formation, coverage and navigation, with specific interest in robotic behavior in adversarial environments. She is a board member of IFAAMAS, and on the editorial board of JAIR and IEEE Transactions on Robotics.



Fausto Giunchiglia is a Professor in the Department of Computer Science and Information Engineering (DISI) at the University of Trento, Italy. His current research is focused on the integration of perception and knowledge representation as the basis for understanding how this generates the diversity of knowledge. His ultimate goal is to build machines which reason like humans. HE is an ECCAI Fellow. He was a President and a Trustee of IJCAI, an Associate Editor of JAIR, and a member of the Editorial Board of many journals, including the Journal of Data Semantics and the Journal of Autonomous Agents and Multi-Agent systems.



Sven Koenig is a professor in computer science at the University of Southern California. Most of his research centers around techniques for decision making that enable single situated agents (such as robots or decision-support systems) and teams of agents to act intelligently in their environments and

exhibit goal-directed behavior in real-time, even if they have only incomplete knowledge of their environment, imperfect abilities to manipulate it, limited or noisy perception or insufficient reasoning speed. He is a AAAI and AAAS fellow and chair of ACM SIGAI. Additional information about Sven can be found on his webpages: idm-lab.org.



Kevin Leyton-Brown is a Professor of Computer Science at the University of British Columbia in Vancouver, Canada. He studies the intersection of computer science and microeconomics, addressing computational problems in economic contexts and incentive issues in multi-agent systems. He also applies machine learning to

various problems in artificial intelligence, notably the automated design and analysis of algorithms for solving hard computational problems. He is a AAAI Fellow, chair of ACM SIGecom, an IFAAMAS board member, and has served as an Associate Editor for AIJ, JAIR, and ACM-TEAC.
