

The Singularity Business

Toward a Realistic, Fine-grained Economics for an AI-Infused World*

Selmer Bringsjord

Dept of Cognitive Science; Dept of Computer Science

Lally School of Management

Rensselaer Polytechnic Institute (RPI)¹

Troy NY 12180 USA

Alexander Bringsjord

PricewaterhouseCoopers LLP²

300 Atlantic Street

Stamford CT 06901 USA

v of 070715NY

Contents

1	Introduction	1
2	The Singularity Causes Extreme Turbulence in Economics	3
2.1	Simon, AI, and Employment	5
2.2	Leontief and Duchin, AI, and Employment	5
2.3	Miller, AI, and Employment	6
3	The MiniMaxularity	7
4	The Economics of the MiniMaxularity	8
4.1	Answer to Q1b: The Shallow AI-Data Cycle	8
4.2	Answer to Q2: Creativity Seemingly Untouchable	11
4.3	Answer to Q3: Repair is Gold	12
5	Recommendations	13
	References	17

*We are greatly indebted to anonymous referees and the editors for helpful comment and criticism, including that which catalyzed our rather more circumspect position on the hypothetical state of economics in a world with machines that have either near-human intelligence in some spheres (our $AI^{<HI}$), human-level intelligence (our $AI^{=HI}$), or super-human intelligence (our $AI^{>HI}$).

1 Introduction

This is an essay on the Singularity business. Contrary to what many might expect upon parsing the previous sentence, we don't use 'the Singularity business' to refer to the general and multi-faceted discussion and debate surrounding the Singularity. In that usage, 'business' is approximately 'intellectual brouhaha.' Instead, we're concerned literally with business and economic questions relating to the Singularity, and to events that would be marked by the arrival of machine intelligence at various levels, including levels below human intelligence. We have elsewhere expressed and defended our claim that belief in the Singularity is fideistic (Bringsjord, Bringsjord & Bello 2013),¹ In the present essay, we are principally concerned with the arrival of a high-but-sub-human level of artificial intelligence that, barring some catastrophe, will *inevitably* materialize. We refer to this level of AI as $AI^{<HI}$. This level of machine intelligence is much lower than that anticipated by Chalmers (2010), who, following Good (1965), believes that *super-human* machine intelligence ($AI^{>HI}$) will in fact arrive and usher in the Singularity (= **Sing**).² Given that 'AI' *simpliciter* refers to *today's* level of machine intelligence, which powers the likes of Google's search engine and IBM's famous *Jeopardy!*-winning Watson (Ferrucci et al. 2010), we are herein most interested in some initial economic and business questions related to $AI^{<HI}$, and the foreseeable road to there from AI. The emergence of $AI^{<HI}$ from AI is a development we — for reasons to be explained — refer to as the *MiniMaxularity*. Placing our emphasis on the MiniMaxularity, or on **MiniMax** for short, ensures that the present chapter is within the realm of the customary realism in which economists and business experts operate. Figure 1 sums up the landscape that underlies the present paragraph. (The "claimed" portion of this landscape will be refined below, when discussing aspects of (Miller 2012).)

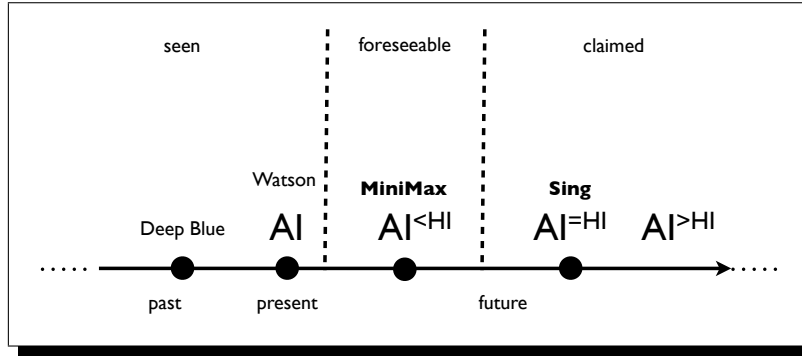


Figure 1: The Basic Landscape/Context of Our Investigation

What sort of economics/business questions do we address? As this paper is intended to be not

¹A view that is e.g. pretty much in line with (Floridi 2015), a piece that, like our chapter, is clearly part of the 'Singularity business' in the philosophical sense.

²We don't spend time discussing herein the possibility of $AI^{=HI}$, i.e., machine intelligence at the level of human intelligence. Chalmers (2010) and Good (1965) assume that were this class of computing machines to arrive, these machines would promptly build $AI^{>HI}$, so by their lights it would make little sense to contemplate an economy in which, in persisting equilibrium, humans and machines, intellectually on a roughly equal plane, work together somehow, perhaps with each group tending to play to its strengths. While we believe that Chalmers and Good are provably wrong in making this assumption (the purported proofs can be found in (Bringsjord 2012)), the reason why consideration of $AI^{=HI}$ is left aside herein is simply that the future event we soon label '**MiniMax**' assumes machine intelligence *short* of the human level.

only realistic, but also primarily an efficient prolegomenon to a sustained science of the economics of artificial intelligence we restrict our attention to only the following three questions. Note that in order to express these questions accurately, we need to have a designator for a generic, all-encompassing concept of machine intelligence that is independent of any relative standing with respect to human intelligence, and independent of any timeline for increasingly intelligent machines. This need is in line with the fact that the vast majority of those economists who have considered what is commonly called ‘automation’ have not had in mind anything like the more fine-grained division shown in Figure 1, but rather at most a very general, informal notion of machine intelligence.³ This means that none of the following designators will do: AI , $AI^{<HI}$, $AI^{=HI}$, $AI^{>HI}$; that is, none of the designators used to form the landscape pictured in Figure 1 will do. Accordingly, we use ‘ $\mathcal{AI}$ ’ to denote the generic concept of machine intelligence. Armed with concept, here now the three questions:

- Q1 (a) What *currently is* the overall state of business and economics in connection with $\mathcal{AI}$, when the more fine-grained landscape of Figure 1 is factored in; and (b) specifically what is the overall state of business and economics in connection with **MiniMax**, impressionistically put?
- Q2 What *will be* the overall state of business and the economy when **MiniMax** arrives, in terms of employment/unemployment?
- Q3 What kind of business strategies make sense today, and in the near term, in light of the road to **MiniMax** that promises to unfold into our future?

Obviously, in the space of one short essay, we can give neither detailed nor ironclad answers even to this hand-picked trio. But our answers, we believe, can be reasonable, non-trivial, and make a contribution to a more systematic treatment of the economics and business of machine intelligence than what is found in today’s literature.

The sequel unfolds as follows. In the next section (2) we provide an answer to question Q1a. To do so, we first simply note the brute fact that $AI^{=HI}$ (and *a fortiori* $AI^{>HI}$) creates what might be called severe “turbulence” in a number of schools of and approaches to economics. The overall reason is that the advent of such powerful machines, by the core formalisms of some prominent schools of economics, will alter today’s economic landscape so severely as to leave in its place a state-of-affairs outside the basic conceptions that are core to these very schools. The section ends with some discussion of three economists (two of whom are late nobelists in economics) who have explicitly considered the connection between $\mathcal{AI}$ and human employment; this discussion serves as the remainder of our answer to Q1a, and as a springboard to rest of the paper. We next (§3) characterize the aforementioned **MiniMax**, which, unlike **Sing**, can be unproblematically pondered from the perspective of perhaps all modern schools of, and approaches to, economics and business. Then, in section 4, we do some brief pondering: we share our answer to Q1b, and then our answers to the other two driving questions. We connect these answers to the stimulating thought of some AI researchers who have specifically considered $AI^{=HI}$ and/or $AI^{>HI}$, in connection with the future of human employment. We end (§5) with a recommendation offered in light of the coming **MiniMax**, and of the answers to Q1–Q3 that we have by then provided.

³E.g., Kelseo & Adler (1958) wrote only two years after the dawn of the discipline of AI (in 1956 at the famous DARPA-sponsored Dartmouth conference; a nice recounting is provided in Russell & Norvig 2009), and had no notion whatsoever of a hierarchy of machine intelligence relative to the human case. This is of course not a criticism of their work; we just report an uncontroversial historical fact.

2 The Singularity Causes Extreme Turbulence in Economics

Certain schools of economics are seemingly unprepared for **Sing**, so much so that these schools would undergo severe turbulence were this event to obtain. In a nutshell, the reason is simply that certain schools of, and approaches to, economics are predicated on a conception of automation that makes room for only *sub*-human intelligence that proceeds under the direct guidance provided by human intelligence, and which works in concert with human labor. In this conception, the productivity of human labor is increased or decreased, as the case may be, by $AI^{<HI}$. But what happens when automation literally *eliminates* the need for human labor, which is presumably a state-of-affairs entailed by the arrival of **Sing**? The answer is that certain schools of and approaches to economics, as many know them and define them today, would be rendered obsolete, for reasons soon to be given.⁴

Please note that we don't say any such thing as that the entire field of modern neoclassical economics will disappear were **Sing** to come to pass. Economists are a very creative, diverse bunch; and they have at their disposal a vast array of formalisms. Even in a post-**Sing** world where no wealth whatsoever is produced by human labor, and where, in light of that fact, certain schools of economics are — as we soon indicate — wholly inapplicable, there would remain the possibility of creating new schools and new formalisms in order to sustain economics as a science. Speculation about the fruits of such creativity is beyond our present purposes; we simply report that nothing we say herein implies that such creativity cannot be exercised.

What, then, is the turbulence we have in mind? It's easy to give a list of different types of turbulence; here goes: Nobelist Milton Friedman's form of free-market capitalism⁵ is predicated on the work of clever humans in search of the wealth that that work provides; but if there is no work that can be done for compensation to humans, his basic model apparently explodes.⁶ Sowell (2010) provides an elegant, laic distillation of the nature of economics, at least as he sees it; but since that distillation defines economics to apply only to systems and processes for maximizing wealth given both scarcity of goods and the need for human insight, ingenuity, and labor in the face of that scarcity, his basic model seems to explode as well in a post-**Sing** world.⁷ Next on our list is the micro-simulation school inaugurated by Orcutt et al. (1961), which presupposes dynamic computational processes for modeling and simulating representative human agents. The problem for this school is that if there isn't anything for the human agents to do, there is nothing to model, and hence nothing to understand and predict, and we have explosion once again. Next, agent-based approaches to economics⁸, are driven by the fundamental notion that the individually simple agents in the relevant simulations represent human agents performing economically relevant behaviors like working, buying, and selling, such that in the aggregate interesting phenomena emerge (e.g. see

⁴We cannot lay claim to being the first to observe the turbulence in question: At the very least, AI researcher Nilsson can be read as entertaining this turbulence; see (Nilsson 1984), and see in particular his sustained argument for the evaporation of human labor in a future in which $AI^{=HI}$ arrives.

⁵In large measure, Friedman's nobel prize was awarded for his work on monetary history and theory (e.g. see Friedman 1963), but as is widely known, he was an arch defender of free-market forces and — as he saw it — human freedom itself as the key driver of wealth (e.g. see Friedman 2002).

⁶Presumably a parallel diagnosis would need to be made of Hayek's (1976) position that search for profit on the part of individuals is a wondrously productive mechanism for both individual and corporate wealth and well-being. In fact, a parallel diagnosis would seem to be in order not just for Hayek, but for the so-called Austrian school, period.

⁷Older influential-and-popular treatments of economics suffer the same fate as Sowell's. E.g., (Hazlitt 1948) makes no room for the "**Sing**-ish" concept of machines capable of everything and more than humans can muster, labor-wise.

⁸As good a place as any to start and work backwards from is (Chen 2016).

Epstein & Axtell 1996); but with no human behavior to model, the approach presumably explodes. Parallel observations can be offered *ad indefinitum*, for school after school, and for approach after approach. Yet, our charitable position remains, and we reiterate it: Despite these — as we have called them — explosions, we don’t assume that the spirit of a given school cannot live on in some newly created model.⁹

Nonetheless, it’s important to understand that while we refuse to speculate about the reaction of schools to the turbulence we cite, there can be no denying that **Sing** creates such turbulence. The inapplicability of many schools of, and approaches to, economics, in a post-**Sing** world, is not speculative in the least. Indeed, the turbulence in question isn’t seen in the light of empirical evidence. Rather, the turbulence can be seen via purely conceptual analysis. This is so because, in certain schools and approaches, the formal schemes for modeling, assessing, and predicting the effects of automation on employment are ones in which the presence of human labor, and the augmentation/diminution of the productivity of that labor, are ineliminable. Put in terms of formal logic, not only are the underlying formal languages presupposed by some economists who study the effects of automation on human employment ones in which, invariably, the relevant predicate and function symbols, and constants, are available to denote humans and human-composed firms, but when axioms in these languages are invoked, they are populated by the use of these symbols and constants. The skeptical reader is encouraged to consult the relevant literature in order to certify our observation. Journal-wise, a nice place to start is (Mortensen & Pissarides 1998); as to textbooks, all of them, whether treating microeconomics or macroeconomics, will be seen to align with our observation, and a nice first step to confirming this is the aforementioned readable and lively (Sowell 2010), and the venerable, bigger (and dear) (Mankew (2014)).¹⁰

The economic implications of **Sing** can be seen as well by closer analysis, including analysis of a select group of economists who have braved explicit consideration of extremely intelligent machines. Importantly, some analysis of the thought of these economists provides an answer to Q1a, and pays dividends when we thereafter characterize and consider **MiniMax**, and on the heels of that consider Q1b–Q3. We now specifically discuss, briefly, three economists: Simon, Leontief (with Duchin), and

⁹Due to both space limitations and constraints on scope, we don’t consider whether our own new paradigm in economics, **logicist agent-based economics** [(Bringsjord, Govindarajulu, Licato, Sen, Johnson, Bringsjord & Taylor 2015); see also (Johnson, Govindarajulu & Bringsjord 2014)], explodes under the assumption that **Sing** has happened.

¹⁰It’s somewhat remarkable that the typical formalisms of both microeconomics and macroeconomics preclude explicit consideration of all-out replacement of human labor with the computing-machine variety, since economists have long had occasion to reflect upon this phenomenon. It seems that when such reflection is engaged, things quickly and mysteriously turn away from any genuine attempt to formalize the prospect. A classic case is an attempt at prophecy by Keynes, who, in trying to look out 100 years into the future from 1930, wrote:

We are being afflicted with a new disease of which some readers may not yet have heard the name, but of which they will hear a great deal in the years to come—namely, *technological unemployment*. This means unemployment due to our discovery of means of economising the use of labour outrunning the pace at which we can find new uses for labour. (Keynes 1931, p. 360; italics his)

But then, unaccountably, Keynes immediately dismisses the concern:

But this is only a temporary phase of maladjustment. All this means in the long run is that *mankind is solving its economic problem*. I would predict that the standard of life in progressive countries one hundred years hence will be between four and eight times as high as it is to-day. (Keynes 1931, p. 360; italics his)

Those coming after Keynes had the benefit of seeing at least some of the wonders of the information-processing age — but still there was no attempt to devise a formal framework allowing for the modeling of outright replacement of human labor by the machine variety.

Miller.

2.1 Simon, AI, and Employment

Herbert Simon was one of the founders of AI, and also a nobel laureate in economics.¹¹ In AI, Simon’s seminal and inaugural work revolved around automated theorem proving; in particular, Simon’s LOGIC THEORIST was the first AI program able to prove a theorem of interest to human beings. In economics, Simon is regaled primarily as the inventor of “bounded rationality” (e.g. see Simon 1972), but the purposes at hand lead us to zero in on Simon’s (1977) optimism about human employment in the face of AI — optimism that has been noted as well by Nilsson (1984), who like us finds the optimism unjustified. Simon believed that in light of the truth of a certain equation, humans would not be disemployed by intelligent machines. For the equation in question, let r_w be the (human) labor wage rate, t_w be the average (human) labor time needed to produce one unit of output, r_i be the interest rate, and c_{1uo} be the average capital required to produce one unit of output. Then:

$$\mathbf{E}: \quad r_w \cdot t_w + (1 + r_i) \cdot c_{1uo} = 1$$

Given $\mathbf{E}$, the basis of Simon’s optimism is easy to state. The basis is simply that automation serves simply to lower t_w , and that as long as r_i remains essentially constant, human wages, as a matter of ironclad arithmetic using $\mathbf{E}$, will continue to rise.

But this sanguinity seems quite problematic, given the more fine-grained landscape of $\mathcal{AI}$ that we have introduced. To see this, we ask: What level of $\mathcal{AI}$ did Simon have in mind when he expressed his optimism about the future of human employment, in light of $\mathbf{E}$? Was he specifically reflecting upon one or more of AI , $\text{AI}^{<\text{HI}}$, $\text{AI}^{=\text{HI}}$, $\text{AI}^{>\text{HI}}$? There is no way of knowing for sure. But our since Simon infamously declared at the dawn of AI that computing machines would soon match human intelligence (Russell & Norvig 2009), it seems quite safe to say that he had in mind $\text{AI}^{=\text{HI}}$.¹² But that makes Simon’s optimism on the strength of $\mathbf{E}$ unfounded. For if we assume that $\text{AI}^{=\text{HI}}$ has arrived, how do we know that the wages paid to humans don’t go effectively to zero, with a compensatory reward to capital rather than labor? One can indeed interpret Kelseo & Adler (1958) to be pressing this question. We see, then, that at the very least, once a more fine-grained landscape is applied to Simon’s position, the basis for his optimism is at best inconclusive.

2.2 Leontief and Duchin, AI, and Employment

Nobelist Leontief, joined by Duchin, are not quite as optimistic as Simon about the prospects for human employment in the face of automation, but on the other hand they are certainly not pessimistic. In their *The Future Impacts of Automation on Workers* (1986),¹³ using the particular methodology of input-output economics (Leontief 1966), they conclude that by 2000 the intensive

¹¹This remarkable (indeed, to our knowledge, singular) combination is discussed at some length in (Johnson et al. 2014).

¹²Presumably Simon would not have agreed with Chalmers that were that level of machine intelligence to arrive, $\text{AI}^{>\text{HI}}$ would inevitably quickly follow.

¹³Economists and other scholars who need to be meticulous about such things, should note that (Leontief & Duchin 1986) is a polished and expanded version of a prior report (i.e. Leontief & Duchin 1984). Nilsson (1984) cites the report only, but the citation appears to be incorrect. The one given here to the report is accompanied by a working URL.

use of automation will lead to a situation in which the same bill of goods can be produced by human labor that is reduced by 10%. The year 2000 has of course come and gone, and whether or not Leontief and Duchin were exactly right, certainly even today, in 2015, human labor remains in relatively high demand, especially in countries that are highly automated.

Yet we must ask what L&D mean by the term ‘automation.’ Recall that we introduced the concept of $\mathcal{AI}$ in part to have on hand a concept that is similarly generic. But what happens to the kind of analysis L&F provide in the context of a more fine-grained view of machine intelligence, relative to the human case? L&D foresaw a significant decline in the need for humans to operate as “clerical workers,” but they appear to understand ‘automation,’ or — in our terminology — $\mathcal{AI}$, as not even at the level of AI. It would be very interesting to see a fresh input-output analysis and forecast based on the identification of ‘automation’ with AI and $\text{AI}^{<\text{HI}}$, but currently Duchin is more interested in the study of other issues.¹⁴ As we cannot find other input-output researchers whose work takes account of the more fine-grained progression of AI through $\text{AI}^{>\text{HI}}$, we are provided with no additional assistance from this school of economics in the search for answers to Q1–Q3.

2.3 Miller, AI, and Employment

To his considerable credit, and in this regard certainly exceeding the reach, relevance, and precision of Simon and L&D in the context of the present paper, economist Miller (2012) considers more directly the issues we have laid at hand. He argues that the Singularity doesn’t entail the disemployment of humans, under certain assumptions, and he makes crucial use of neoclassical economics in his argument. The kernel of his reasoning is an adaptation of the famous and widely affirmed case for free trade, which employs a framework that can be traced back to Ricardo (1821), and in fact ultimately (at least in spirit) to Adam Smith (2011). Miller’s (2012) instantiation of the framework is conveyed via a clever example in which — to use our notation — $\text{AI}^{>\text{HI}}$ are able to produce both flying cars and donuts. The former are beyond the reach of humans, whereas the latter are just made more slowly by humans. Under predictable assumptions, it’s advantageous for the superior machines to forego their production of donuts by purchasing them from the humans. Doing so frees up time they can spend on more profitable pursuits (like making flying cars).¹⁵

We find this line of argument to be important and worthy of further investigation, and to be commendably hopeful, but on the other hand regard one of its key assumptions to be very questionable. This assumption is revealed once one understands that $\text{AI}^{>\text{HI}}$ can be construed in two fundamentally different ways (Bringsjord 2012). In the first way, the essence of $\text{AI}^{>\text{HI}}$ is merely the pure speed of its information-processing, rather than the *nature* of that processing. In short, the idea here is that the intellectually super-human machines are super-human because they can compute Turing-computable functions much faster than we can. We too can compute the functions, albeit slowly. We label this sub-class of machine intelligence $\text{AI}_T^{>\text{HI}}$. The other sub-class is composed of computing machines that are able to carry out information-processing that is qualitatively more powerful than what can be carried out by a Turing machine or one of its equivalents (e.g., a register machine).¹⁶ We denote the sub-class of machines in this category by $\text{AI}_H^{>\text{HI}}$. Now let’s turn back to

¹⁴Personal communication.

¹⁵Miller gives an intimately related argument from chess: He claims that the current superiority of hybrid human-machine chessplaying over both independent human chessplaying and independent machine chessplaying opens up the possibility that post-**Sing** machines will collaborate with us. However, chess is a Turing-solvable game, and fundamentally easy (Bringsjord 1998).

¹⁶There is now a mature mathematics of information-processing beyond what Turing machines can muster. A nice

Miller’s example, and see the crucial assumption therein.

The crucial assumption is that innovations achievable by $AI^{>HI}$ are not so radical and valuable as to break *utterly* outside the range of what is humanly understandable. But this would hold only if $AI^{>HI} = AI_T^{>HI}$; it would not hold if $AI^{>HI} = AI_H^{>HI}$. In short, if Miller’s post-**Sing** machines managed to parallelize the production of everything within human reach to the point of infinitesimal time and effort (including the “production” of long-term strategies in chess; see note 15), it’s very hard to see how Ricardo’s (1821) rationale would have any bite at all. (If fully formalized, Ricardo’s (1821) framework will be seen to presuppose an at-once linear and finitary conception of time and effort on the part of the agents in question.)¹⁷ We thus conclude that even the status of Miller’s analysis, in light of the more fine-grained progression of machine intelligence that we have brought to bear, is unclear.

Let us take stock of where we find ourselves at this juncture, with respect to an answer to question Q1a: We have seen that the implications of **Sing**, for certain schools of economics, are severe and disruptive. When we turned specifically to consideration of the work of illustrious economists who grappled with the issue of human employment in the context of machine intelligence, we found that no conclusive answer can be given to Q1a, in the light of the more fine-grained progression of machine intelligence that we have introduced to frame the discussion. We turn now to consideration of **MiniMax**, and questions Q1b, Q2, and Q3.

3 The MiniMaxularity

MiniMax, unlike what might well be the case with respect to its lofty cousin **Sing**, is no pipe dream. It will — barring some out-of-the-blue asteroid or some such thing — definitely come. What is it? It is the arrival of machine intelligence that is on the one hand “minimized” with respect to all dimensions of human-level cognition that as of yet are, for all we know, beyond the reach of standard computation, and beyond the reach of AI assumed to be “frozen” at its current point of development, logico-mathematically speaking. But on the other hand, the level of machine intelligence in **MiniMax** is “maximized” with respect to all aspects of cognition that are at present within the range of the relevant human science and engineering. This can be put another way, with help from the dominant, encyclopedic handbook of modern AI: **MiniMax** is the future point at which the techniques in *Artificial Intelligence: A Modern Approach*¹⁸ are not only further refined (without the introduction of fundamentally new formal paradigms), but are run on hardware that is many orders of magnitude faster than what is available today, *and* are fully integrated with each other and placed interoperably within particular artificial agents. In terms of our progression, The machines that characterize **MiniMax** are $AI^{<HI}$; see again Figure 1.

Put yet another way, in granting that the MiniMaxularity will come, we grant that machine intelligence will indeed reach great heights, but will be “minimal” relative to The Singularity (e.g., machines will not have subjective awareness or self-consciousness; see Bringsjord 1992), yet “maximal” with respect to certain logico-mathematical constraints. These constraints, put impressionistically, amount to saying that computing machines will reach a level of intelligence that is maximal along the lines of the smartest such machines we have so far seen. A paradigmatic example of such a machine

example is (Hamkins & Lewis 2000).

¹⁷Also, the assumption that there is an intersection of utility maximization between $AI^{>HI}$ and humans seems to us tenuous.

¹⁸I.e., (Russell & Norvig 2009).

is IBM’s Watson, a QA (QuestionAnswering) system that famously defeated the two best *Jeopardy!* players on our planet. Watson was engineered in surprisingly short order, by a tiny (but brilliant and brilliantly led) team, at a tiny cost relative to the combined size of today’s AI companies, so clearly it portends great power for AI systems — but on the other hand, the logico-mathematical nature of Watson is severely limited. Specifically, from the standpoint of the knowledge and reasoning that enabled Watson to win, it is at a level below first-order logic (FOL); and this turns out to be precisely the level at which, as we explain below (§4.1), today’s AI technology, from today’s AI companies, is at most operating at. **MiniMax** occurs when computing machines much more powerful than Watson arrive, but are constrained by this same logico-mathematical nature. (A formal analysis of Watson in the context of a somewhat elaborate framework for assessing the logico-mathematical nature of AI systems, past, present, and future, is provided in (Bringsjord & Govindarajulu forthcoming).)

4 The Economics of the MiniMaxularity

We now proceed to extrapolate from today to **MiniMax** tomorrow. We do so by taking up the three driving questions enumerated (and partially answered) above, reiterated here for convenience:

- Q1 (a) What *currently* is the overall state of business and economics in connection with $\mathcal{AI}$, when the more fine-grained landscape of Figure 1 is factored in; and (b) specifically what is the overall state of business and economics in connection with **MiniMax**, impressionistically put?
- Q2 What *will be* the overall state of business and the economy when **MiniMax** arrives, in terms of employment/unemployment?
- Q3 What kind of business strategies make sense today, and in the near term, in light of the road to **MiniMax** that promises to unfold into our future?

4.1 Answer to Q1b: The Shallow AI-Data Cycle

The “cloud,” “big data,” “machine learning,” “cognitive computing,” “natural language processing,” ... are all flowing glibly from the tongues of every high-powered executive within every sector, of every sub-industry, under every industry. Whether it be financial services, healthcare, real estate, social networking, retail, alternative energy, or cloud/data providers themselves, massive corporations, and for that matter non-massive ones, all are seeking to harness the power of these technologies to optimize their business processes and add value to their customers.

We ask two sub-questions about this state-of-affairs:

- Q1a₁ What is the role and nature of $\mathcal{AI}$ in the acquisition and exploitation of the data that stands at the heart of this state-of-affairs?
- Q1a₂ What is the nature of the data that this $\mathcal{AI}$ is designed to handle?

In broad strokes, the answer to Q1a₁ is simply that $\mathcal{AI}$ plays an absolutely essential role. This is so for the simple reason that the amount of data being collected on humans, products, services, mobile devices, etc. is multiplying at such a rapid rate that manual processing is of course impossible. The possibility for the human brain to keep up, in a purely computational sense, is gone. We have vast and seemingly endless data centers, server farms, data warehouses; and analyzing all of this data can only happen on the strength of $\mathcal{AI}$. Hence, machines are assisting machines assisting

machines assisting ... machines in order to assist far-removed humans. All of this we take to quite undeniable.¹⁹

But this processing is really just one of three stages in a cycle that will increasingly dominate industrial economies. The cycle also includes a stage in which computing machines *acquire* the data, for acquisition too is beyond the capacity of manual efforts on the part of humans. And then the cycle is completed by the third stage: actions performed on the strength of the analysis of the relevant data. This three-stage cycle, note, is simply a generic description of the perceive-process-act cycle that is the essence of an artificial intelligent agent (Russell & Norvig 2009). Hence, what is happening before our eyes is that the machines in the class AI are becoming ubiquitous, powerful, and (at least when it comes to running the cycle in question in real time) unto themselves. Before long, this three-stage cycle will be entirely driven by the AI machines in many, many domains; and the cycle will as time passes accelerate to higher and higher speeds, and be in a real sense inaccessible to human cognition. In fact, the march to **MiniMax** arguably consists in both the application of the three-stage cycle to more and more domains, and the acceleration of the cycle.

We believe that it's important to realize that the cycle in question directly relates to technologies for automated perception, and automated actions. This is why companies like Google and Apple are so interested in technologies like speech recognition. Even though speech recognition isn't at the heart of human general intelligence (after all, one can be a genius, yet be unable to hear and for that matter unable to speak, from birth), such capability is incredibly valuable for enhancing cognition — or to put the point in terms of the more practical business/economic matters with which we are now concerned: the three-phase data cycle we have singled out is, and will continue to be, greatly amplified by speech recognition technology (and other technologies at the perception-action level). This ties back to the earlier discussion of Miller's example of human-machine coöperation, and specifically to S. Bringsjord's (2012) point that one perhaps-defensible fleshing-out of the nature of $AI^{>HI}$ is that these machines compute NP-complete functions within economically meaningful periods of time. That is, to use the notation introduced in §2.3, the $AI^{>HI}$ may be $AI_T^{>HI}$.

Certainly the cycle we describe will for instance be concretized on our road system, where sooner rather than later vehicles will be machine-controlled on the basis of this cycle, largely independent of humans, who will be firmly positioned outside the cycle.²⁰

And what of the second sub-question, Q1a₂? Our answer, in short, is that the data in the cycle discussed immediately above is both inexpressive and (relative to the machines intended to “understand” this data) semantically shallow. The data is inexpressive in the rigorous sense that the formal languages needed to express the data are themselves inexpressive, as a matter of formal fact. It's well-known and well-documented (e.g. see the mere use of RDF and Owl in Antoniou & van Harmelen 2004) that even the Semantic Web is associated with, and indeed increasingly based upon, languages that don't even reach the level of FOL: description logics (Baader, Calvanese & McGuinness 2007), for example. Since even basic arithmetic requires more than FOL,²¹ it's very

¹⁹Notice that we have used the generic label ‘*AI*’ to characterize the situation. Clearly, in this situation it's specifically AI that is in use.

²⁰Our readers may be well-served by reading the discussion, in (Brynjolfsson & McAfee 2011), of the speed with which self-driving vehicles came upon us, which completely overturned “expert” opinion that such technology would only arrive in the distant future.

²¹Peano Arithmetic (**PA**), which captures all of basic arithmetic that young students routinely master, is a set of axioms, where each member of the set is a formulae in FOL (a nice presentation is in Ebbinghaus, Flum & Thomas 1994). But **PA** is infinite, and is specified by the use of beyond-FOL machinery able to express axiom *schemata*.

hard to see how the current data cycle is getting at the heart of human intelligence, since such intelligence routinely handles not just “big data,” but infinite data (Bringsjord & Bringsjord 2014).

We also claimed that the data in the three-stage cycle is “semantically shallow.” What does this claim amount to? It’s easy enough to quickly explain, at least to a degree, what we are referring to, by appealing to two simple examples.

For the first example, imagine that a husband, Ronald, gives a Valentine’s-Day gift to his wife, Rhonda. This gift is a bouquet of roses, accompanied by a box of chocolates. When he hands both the bouquet and box to her, she says, with exaggerated warmth in her voice: “Ah, I see that your usual level of thoughtfulness is reflected in the age of these flowers!” Rhonda then also promptly opens the box and eats a chocolate nugget from therein. She then says: “More thoughtfulness! O Ronald, such love! These morsels are wonderfully stale. And again, really and truly, wilted roses are every woman’s favorite!” What is the semantic meaning of what Rhonda has said? Well, she has communicated a number of propositions, but certainly one of them is that Ronald is actually *thoughtless*, because he hasn’t even managed to give fresh-cut roses (or at least ones that haven’t already wilted), and because the chocolates are well beyond the expiration date on their box. The meaning of what Rhonda has uttered is quite beyond the mere syntax she has used. The data cycle that drives present-day companies is completely separate from the ability of computing machines to understand the semantic nature of this data. And this is true not only in the case of natural language, but in mathematics. This can be seen via a second example:

Consider the set $\mathcal{O}$, which is Kleene’s (1938) famous collection of recursive notations for every recursive ordinal. Put in barbarically intuitive fashion, one might say that $\mathcal{O}$ points to all that can be obtained mechanically and in finitary fashion regarding the nature of mathematics. That by any metric implies that $\mathcal{O}$ carries a staggering amount of semantic information — all delivered via a single symbol. We are hoping that exceedingly few of our readers have deep understanding of this set, for that lack of understanding, combined with our reporting the facts that (i) $\mathcal{O}$ carries an enormous amount of semantic content, and (ii) that that content is beyond what a computing machine can understand in — to harken back to Figure 1 — the foreseeable future serves to make our point.²² The data that computers process in the three-stage data cycle is inexpressive and carries no deep semantic value, as confirmed by the fact that nothing like $\mathcal{O}$ is in the data in question.

Assuming that we’re right about the nature of the data that is processed by the three-stage cycle that stands at the heart of the current age of AI, and the coming age of AI^{HI} , what are the implications? We mention only one immediate implication, which fits well our overarching confidence that **MiniMax** will arrive, but not **Sing**: Since human intelligence, at least in its “highest” forms (e.g., in the form of sophisticated natural-language communication, and in logic and mathematics), is provably joined to formal languages that are highly expressive and deeply semantic, the machines in our present and our foreseeable future will be limited relative to the human case. Note that since the mathematics we’re talking about is the basis for vast amounts of human endeavor, not just mathematics itself, computing machines in the foreseeable future will be unable to not only do mathematics, but unable to do things that employ mathematics. Ironically, a good example is macroeconomics itself, which employs real analysis.²³

²²It’s often said these days that what humans find hard, computers find easy; and that what humans find easy, computers find hard. In logic and mathematics, this doesn’t hold true, for in these fields one doesn’t contribute by doing mere *calculation* (which is, we of course concede, hard for humans and easy for computing machines), and the real contributions are still being made only by humans. This will continue into and through the foreseeable future.

²³E.g., see the intriguing case in favor of Keynesian spending articulated in (Woodford 2011), in which economies are modeled as infinitely-lived “households” that maximize utility through infinite time series, under for instance

4.2 Answer to Q2: Creativity Seemingly Untouchable

Expressed with brutal brevity, our answer to Q2 is that **MiniMax** will “knock out” human jobs that aren’t genuinely creative; but jobs that *are* of this type will be secure as granite. This reply is likely to strike many readers as disturbingly uninformative, as it stands. Obviously the notion of *genuinely creative* must be, at least to some degree, explained.

In order to begin to unpack this concept, we start by drawing the reader’s attention to what we see as a fatal flaw in the pessimism of the late James Albus with respect to human disemployment. Though an illustrious engineer of intelligent, “brain-inspired” systems (see e.g. the interesting Albus 1981) and not a professional economist, Albus bravely wrote that what he in 1976 called ‘super-automation’ had the potential to cause massive human unemployment (Albus 1976).²⁴ He specifically perceived a simple but profoundly disturbing “paradoxical situation where automation is generally conceded to be a major source of our national wealth, yet new advances in automation are widely feared and often actively opposed by a large segment of the population.” (Albus 1976, p. 40) In a bit more detail:

As machines grow more efficient, they produce more wealth, and the human workers’ wages rise accordingly. [N.B.: Simon’s **E**, or some close relative, seems here presupposed by Albus.] Eventually, however, the machines become proficient enough to function without human assistance. At that point, human workers serve no function, and their inflated salaries make them a costly liability. (Albus 1976, p. 40)

The part of this that partakes of modern economics and its standard formalisms for modeling automation and job creation/destruction is uncontroversial. For instance, Mortensen & Pissarides (1998) deploy a formal framework in which machines, through automation, do increase productivity and human wealth; but — in line with what we pointed out in section 2 — they don’t consider the scenario Albus points to: machines becoming “proficient enough to function without human assistance.” This scenario is of course entailed by **Sing**; but not by mere **MiniMax**.

We now point out that Albus’ views are, for formal reasons, problematic. To see this, we begin by quoting Albus as follows:

[O]nly a relatively small number of workers will be needed to run high-tech businesses and operate largely automated production facilities. Eventually, as machine intelligence grows, these jobs will be vulnerable as well. Information technology will be the next sector where intelligent systems technology will have an impact on productivity. Improved software will be much more user-friendly, easier to install, and simpler to maintain. ... *Automatic software-development systems will do most of the future programming tasks* ... (Albus 2011, Chapter 3; emphasis ours)

What we have here emphasized in the quote is the prediction that computing machines will automatically program computing machines. This is a challenge traditionally referred to as *automatic programming*. In bare-bones terms, the challenge consists in receiving some semi-formal description (e.g., what might be found in a textbook or journal article, which uses a mixture of natural language

the constraint that the specific, underlying function u , which returns the utility produced by the consumption of a good, must be such that its first derivative is greater than zero, while its second derivative is less than zero. In how Woodford here models an economy, he follows a longstanding neoclassical approach articulated e.g. by Barro & King (1984).

²⁴There can be little doubt that here (and in the antidote he prescribed) Albus echoed Kelseo & Adler (1958). And indeed as Duchin has informed us in personal communication, Leontief too was apparently quite concerned about massive disemployment, and found attractive the prospect of conveying ownership of robots to the populace.

and formal notation) D_f of a number-theoretic²⁵ function f , and having to produce as output at least one computer program P_f that verifiably²⁶ computes f . At the time Albus wrote (Albus 2011), there was still considerable optimism about the possibility of programs that can write programs (under the constraints we have given); now we know better. Today, despite many millions of dollars spent on the problem, essentially no progress has been made (Bringsjord & Arkoudas 2009), and next to no one claims to see a time in the near future when programs program themselves. (And since there's vanishingly small funding available for attacking automatic programming, things are for economic reasons unlikely to change any time soon.) Therefore, we have here a kind of creativity absent from the current digital marketplace, and from the foreseeable future.²⁷ In other words, this creativity is not within the reach of $\text{AI}^{<\text{HI}}$, ergo not part of **MiniMax**.

4.3 Answer to Q3: Repair is Gold

We humbly concede that there are an infinite number of business strategies²⁸ that will be successful along the path to reaching **MiniMax**. Obviously a sizable cluster of these center around the construction and sale of the very machines that fall into the class $\text{AI}^{<\text{HI}}$; but this basic strategy isn't our concern here, in part because it's utterly obvious that this sub-class of machines exists, and includes many that, if implemented, will be very successful. (After all, many are *already* successful, where the machines are not $\text{AI}^{<\text{HI}}$, but AI .) What less-obvious strategies make sense, in light of the steady march to **MiniMax**? We reply with one word: repair. We refer here to repair of the very computing machines that will travel the road to **MiniMax**.

Sparing the reader any of the relevant formal logic, we simply point out that all of the aforementioned extreme intrinsic difficulty of generating a program P_f that computes some given-as-input (number-theoretic) function f carries over to the domain of repair. This is so because, given a program $P_{\tilde{f}}$ known to fail to compute the function f , along with the task of modifying $P_{\tilde{f}}$ so as to yield P_f , is to face a problem that incorporates the original challenge. Intelligent programs and

²⁵These are here (total) functions f from the natural numbers $\mathbf{N}$ (or from $\mathbf{N} \times \dots \times \mathbf{N}$) to the natural numbers $\mathbf{N}$.

²⁶Without this requirement, a machine could be programmed to simply spit out mere possibilities, and perhaps by sheer serendipity one or more of them might compute the function f described in D_f . Given the requirement, approaches to automatic programming based entirely on genetic programming, since they are bereft of proofs that programs that are produced by evolution do compute f , are inadequate.

²⁷Formally and computationally speaking, how hard is the automatic programming problem? *Very*. Not only is the problem Turing-uncomputable, but it's more difficult than the halting problem. Hence only machines in the class $\text{AI}_H^{>\text{HI}}$ would have a chance.

²⁸It's probably worth pointing out that the concept of a business strategy is to our knowledge happily allowed to be informal within the academic area known as *corporate strategy*, and we follow suit herein. (Were *business strategy* to be formalized, presumably an intensional logic would be required, since such things as propositional attitudes are central to this concept. For presentation of a platform intended to allow such formalization, and associated computational simulation, see (Bringsjord 2008).) That said, some characterization is possible. For example: A business strategy can be invented and considered in advance of invention and innovation at the level of technology. Put more concretely, the CEO can decide that his/her company X Inc is going to enter a completely new but projected-to-grow market (this decision is to adopt a certain strategy), even if at the time of this decision X Inc has no manufacturing capacity whatsoever in this market. In particular, our CEO could decide that X Inc should get into the business of — to anticipate what we are about to say — repairing humanoid robots, even if X Inc has no capability in robotics at all. Our distinction between strategies on the one hand, and ideas, inventions, entrepreneurship, and business models on the other, is supported by the literature. E.g., in the *locus classicus* for what invention, innovation, and entrepreneurship consist in, viz. (Schumpeter 1934), one finds that a business strategy is at a radically different, higher level.

intelligent robots (where the intelligence of the latter consists in their having “minds” based on the former), then, at least if we base our diagnosis on the general case, can’t create themselves, but they *also* can’t repair themselves, and it seems to stand to reason that businesses devoted to the repair domains have a very good chance of thriving in the approach to **MiniMax**, and in the coming age of **MiniMax**, when machines in $AI^{<HI}$ are with us.

5 Recommendations

It has become fashionable to offer pontifical and far-reaching recommendations about “what to do” in light of the implacable advance of $\mathcal{AI}$, and the attendant bleak prospects for human employment as the attempt to progress toward $AI^{=HI}$ unfolds. Perhaps²⁹ the primogenitor in this tradition is Albus (1976), who recommends that as protection against **Sing** (and, by trivial inference, against the less exotic **MiniMax**), steps should be taken well ahead of time³⁰ to establish shared human ownership of intelligent computing machines and robots. The idea is that while you might not have a job, at least you’d be part owner of the machines that render you superfluous, and would as a result receive money deriving from their paid-for prowess. By ‘shared human ownership,’ Albus means that *all* humans on the planet, whether in possession of capital or education or relevant talent or energy, would enjoy appreciable ownership; hence his use of the label ‘People’s Capitalism’ for his recommendation. Albus’ recommendation may well be “spiritually” sound, and commendably other-regarding, but there is no reason to think that it’s logico-mathematically sound. What formalisms, frameworks, and mechanisms rigorously define his recommendation, and allow us to see how it might get launched, and sustained? The answer, alas, is that no such things were devised by Albus, nor by anyone else of his persuasion. We aren’t asking for math for math’s sake, as ideologically driven formalists. No, it’s *natural* to ask to see the math, in light of obvious challenges. For example, if shared ownership of robotics firms was by fiat granted across the globe in 2015, how would this arrangement be *sustained*, in light of the inevitable vicissitudes of corporations as generations pass on and new ones are born? And given that the human population is bound to increase, how does ownership established for humans alive in 2015 cover the case of humans alive in 2035?

More recently, Brynjolfsson & McAfee (2014) continue the tradition that Albus initiated: In Chapter 14 of their *The Second Machine Age*, “Long-Term Recommendations,” these authors recommend explicit and serious consideration of the like-minded idea that, as insurance against any serious machine-caused disemployment, humans should all receive a “basic income,” a lump sum of money that they don’t earn, but which is gifted to them. B&F don’t stop here. They also recommend that a *negative* income tax, which, as they remind the reader, nobelist Milton Friedman found appealing, be implemented, and they also want to see a dramatic reduction in positive tax levels, in light of the advance of intelligent machines, and of the option that companies increasingly have of turning to mechanical labor rather than human labor. For example, they write:

As digital technologies keep acquiring new skills and capabilities, these same organizations will increasingly have another option [other than moving away from domestic human employees to

²⁹We say ‘Perhaps’ because, again, (Kelseo & Adler 1958) is a complicating factor. Certainly Albus is the first AI engineer in the tradition.

³⁰One of us (S.) knows from personal interaction with Albus that on his “timetable,” the planet is now well beyond the point in time at which global ownership should (according to him) be established. But given e.g. that Google has of late acquired a number of robotics companies, and that neither of us was automatically granted some ownership in the course of these transactions, we feel rather confident in asserting that Albus’ recommendation has yet to take root.

human employees in other countries]: they’ll be able to make use of digital laborers rather than humans. The more expensive human labor is, the more readily employers will switch over to machines. And since payroll taxes make human labor more expensive, they’ll very likely have the effect of hastening this switch. (Brynjolfsson & McAfee 2014, Chap. 14, §“Better Than Basic: The Negative Income Tax”)

The final recommendation B&F make is that humans seek to partner with machines, rather than compete with them. As they put it: “[I]t’s not the case that people cease to be valuable the instant computers surpass them in a domain. They can be enormously useful once they’ve paired up to race *with* machines, instead of against them.” (Brynjolfsson & McAfee 2014, Chap. 14, §“The Peer Economy and Artificial Intelligence”) Of course, as the reader knows by now, Miller (2012) has already made this point in the past; recall section 2.3.

We here observe but two propositions regarding such recommendations: One, they are, by any metric, simply impracticable, because following them would require sea-changes in the social and political landscapes. Such monumental changes, in a pluralistic democracy usually capable of no more than incremental change born out of compromise between competing paradigms, philosophies, and ideologies, strike us as exceedingly unlikely; accordingly, those recommending them strike us as irrationally hopeful. Two, no formalisms, and *a fortiori* no theorems, are provided in order to make clear and justify such recommendations.

In this context, we conclude by issuing but a single recommendation, one that, compared with the ambitious, sweeping recommendations of Albus, Brynjolfsson, McAfee, and other like-minded thinkers, appears to be eminently easy for any industrialized economy to activate:

$\mathcal{R}$ We recommend that a logico-mathematically sophisticated investigation of the economics of an AI^{<HI}-infused planet,³¹ virtually non-existent in the intellectual landscape of today, be inaugurated, on the strength of federal government-sponsored support issued through extant agencies (e.g., in the U.S., the National Science Foundation).

If, courtesy of the present paper, we have made even the slightest headway toward a time when $\mathcal{R}$ is executed, our labor may not have been in vain.

References

- Albus, J. (1976), *People’s Capitalism: The Economics of the Robot Revolution*, New World Books, Kensington, MD.
- Albus, J. (1981), *Brains, Behavior, Robotics*, BYTE Books, Peterborough, NH.
- Albus, J. (2011), *Path to a Better World: a Plan for Prosperity, Opportunity, and Economic Justice*, iUniverse, Bloomington, IN. A Kindle book available from Amazon. ISBN: 978-1-4620-3534-2.
- Antoniou, G. & van Harmelen, F. (2004), *A Semantic Web Primer*, MIT Press, Cambridge, MA.
- Baader, F., Calvanese, D. & McGuinness, D., eds (2007), *The Description Logic Handbook: Theory, Implementation (Second Edition)*, Cambridge University Press, Cambridge, UK.
- Barro, R. & King, R. (1984), ‘Time-Separable Preferences and Intertemporal-Substitution Models of Business Cycles’, *Quarterly Journal of Economics* **99**, 817–839.

³¹A key aim of this investigation would be to erect a formal hierarchy of intelligence, taking about of the progression in Figure 1, that can be used in specifying levels of human-machine collaboration. Recall our comments at the conclusion of §2.3.

- Bringsjord, E. & Bringsjord, S. (2014), Education and ... Big Data Versus Big-But-Buried Data, in J. Lane, ed., 'Building a Smarter University', SUNY Press, Albany, NY, pp. 57–89. This url goes to a preprint only.
URL: http://kryten.mm.rpi.edu/SB_EB_BBBD_0201141900NY.pdf
- Bringsjord, S. (1992), *What Robots Can and Can't Be*, Kluwer, Dordrecht, The Netherlands.
- Bringsjord, S. (1998), 'Chess is Too Easy', *Technology Review* **101**(2), 23–28.
URL: <http://www.mm.rpi.edu/SELPAP/CHESSSEASY/chessistooeasy.pdf>
- Bringsjord, S. (2008), Declarative/Logic-Based Cognitive Modeling, in R. Sun, ed., 'The Handbook of Computational Psychology', Cambridge University Press, Cambridge, UK, pp. 127–169.
URL: http://kryten.mm.rpi.edu/sb_lccm_ab-toc_031607.pdf
- Bringsjord, S. (2012), 'Belief in The Singularity is Logically Brittle', *Journal of Consciousness Studies* **19**(7), 14–20.
URL: http://kryten.mm.rpi.edu/SB_singularity_math_final.pdf
- Bringsjord, S. & Arkoudas, K. (2009), A Cognitively Informed Approach to Automatic Programming. This report, prepared for the U.S. National Science Foundation, was made possible by crucial contributions from J. Li to S. Bringsjord, for which Bringsjord is deeply grateful.
URL: http://kryten.mm.rpi.edu/KA_SB_CreativeIT_SGER_finrep_120908.pdf
- Bringsjord, S., Bringsjord, A. & Bello, P. (2013), Belief in the Singularity is Fideistic, in A. Eden, J. Moor, J. Søraker & E. Steinhardt, eds, 'The Singularity Hypothesis', Springer, New York, NY, pp. 395–408.
- Bringsjord, S. & Govindarajulu, N. S. (forthcoming), Leibniz's Art of Infallibility, Watson, and the Philosophy, Theory, & Future of AI, in V. Müller, ed., 'Synthese Library', Springer.
URL: http://kryten.mm.rpi.edu/SB_NSg_Watson_Leibniz_PT-AI_061414.pdf
- Bringsjord, S., Govindarajulu, N. S., Licato, J., Sen, A., Johnson, J., Bringsjord, A. & Taylor, J. (2015), On Logicist Agent-Based Economics, in 'Proceedings of Artificial Economics 2015 (AE 2015)', University of Porto, Porto, Portugal.
URL: http://kryten.mm.rpi.edu/SB_JL_AS_JJ_AB_NS_JT_AE2015_0605152315NY.pdf
- Brynjolfsson, E. & McAfee, A. (2011), *Race Against the Machine*, Digital Frontier Press. Sold by Amazon Digital Services.
- Brynjolfsson, E. & McAfee, A. (2014), *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*, Norton, New York, NY. Kindle edition.
- Chalmers, D. (2010), 'The Singularity: A Philosophical Analysis', *Journal of Consciousness Studies* **17**, 7–65.
- Chen, S.-H. (2016), *Agent-Based Computational Economics: How the Idea Originated, and Where it is Going*, Routledge, New York, NY.
- Ebbinghaus, H. D., Flum, J. & Thomas, W. (1994), *Mathematical Logic (second edition)*, Springer-Verlag, New York, NY.
- Epstein, J. & Axtell, R. (1996), *Growing Artificial Societies*, Bradford Books, Cambridge, MA.
- Ferrucci, D., Brown, E., Chu-Carroll, J., Fan, J., Gondek, D., Kalyanpur, A., Lally, A., Murdock, W., Nyberg, E., Prager, J., Schlaef, N. & Welty, C. (2010), 'Building Watson: An Overview of the DeepQA Project', *AI Magazine* pp. 59–79.
URL: <http://www.stanford.edu/class/cs124/AIMazine-DeepQA.pdf>
- Floridi, L. (2015), 'Singularitarians, AItheists, and Why the Problem with Artificial Intelligence is H.A.L. (Humanity At Large), not HAL', *Philosophy and Computers* **14**(2), 8–11. This is part of NEWSLETTER published by the American Philosophical Association.

- Friedman, M. (1963), *A Monetary History of the United States, 1867–1960*, Princeton University Press, Princeton, NY.
- Friedman, M. (2002), *Capitalism and Freedom: Fortieth Anniversary Edition*, University of Chicago Press, Chicago, IL. This is a re-issue. The book was originally published in 1962.
- Good, I. J. (1965), Speculations Concerning the First Ultraintelligent Machines, in F. Alt & M. Rubinoff, eds, ‘Advances in Computing’, Vol. 6, Academic Press, New York, NY, pp. 31–38.
- Hamkins, J. D. & Lewis, A. (2000), ‘Infinite Time Turing Machines’, *Journal of Symbolic Logic* **65**(2), 567–604.
- Hayek, F. A. (1976), *The Mirage of Social Justice*, Law, Legislation and Liberty, Volume 2, University of Chicago Press, Chicago, IL.
- Hazlitt, H. (1948), *Economics in One Lesson*, Pocket Books, New York, NY. This little classic can be obtained free of charge online in multiple places; e.g., see http://www.fee.org/pdf/books/Economics_in_one_lesson.pdf.
- Johnson, J., Govindarajulu, N. & Bringsjord, S. (2014), ‘A Three-Pronged Simonesque Approach to Modeling and Simulation in Deviant ‘Bi-Pay’ Auctions, and Beyond’, *Mind and Society* **13**(1), 59–82.
URL: http://kryten.mm.rpi.edu/JJ_NSG_SB_bounded_rationality_031214.pdf
- Kelseo, L. & Adler, M. (1958), *The Capitalist Manifesto*, Random House, New York, NY.
- Keynes, J. M. (1931), Economic Possibilities for our Grandchildren, in ‘Essays in Persuasion’, Macmillan, London, UK, pp. 358–373. The essay itself was apparently penned in 1930.
- Kleene, S. (1938), ‘On Notation for Ordinal Numbers’, **3**(4), 150–155.
- Leontief, W. (1966), *Input-Output Economics*, Oxford University Press, Oxford, UK.
- Leontief, W. & Duchin, F. (1984), The Impacts of Automation on Employment, 1963–2000. Final Report, Technical report, New York University: Institute for Economic Analysis.
URL: <http://eric.ed.gov/?id=ED241743>
- Leontief, W. & Duchin, F. (1986), *The Future Impacts of Automation on Workers*, Oxford University Press, Oxford, UK.
- Mankiw, N. G. (2014), *Principles of Economics (7th ed)*, Cengage Learning, Independence, KY.
- Miller, J. (2012), *Singularity Rising: Surviving and Thriving in a Smarter, Richer, and More Dangerous World*, BenBella Books, Dallas, TX.
- Mortensen, D. & Pissarides, C. (1998), ‘Technological Progress, Job Creation, and Job Destruction’, *Review of Economic Dynamics* **1**, 733–753.
- Nilsson, N. (1984), ‘Artificial Intelligence, Employment and Income’, *AI Magazine* pp. 5–14.
- Orcutt, G., Greenberger, M., Korb, J., Korb, J. & Rivlin, A. (1961), *Microanalysis of Socioeconomic Systems: A Simulation Study*, Harper and Brothers, New York, NY.
- Ricardo, D. (1821), *On the Principles of Political Economy, and Taxation*, Murray, London, UK. This is the 3rd edition.
- Russell, S. & Norvig, P. (2009), *Artificial Intelligence: A Modern Approach*, Prentice Hall, Upper Saddle River, NJ. Third edition.
- Schumpeter, J. (1934), *The Theory of Economic Development*, Transaction Publishers, New Brunswick, NJ.
- Simon, H. (1972), Theories of Bounded Rationality, in C. McGuire & R. Radner, eds, ‘Decision and Organization’, North-Holland, Amsterdam, The Netherlands, pp. 361–176.
- Simon, H. (1977), *The New Science of Management Decision*, Prentice-Hall, Englewood Cliffs, NJ.

- Smith, A. (2011), *An Inquiry Into the Nature and Causes of the Wealth of Nations*, Simon and Brown. First published in 1776.
- Sowell, T. (2010), *Basic Economics: A Common Sense Guide to the Economy*, Basic Books, New York, NY.
- Woodford, M. (2011), ‘Simple Analytics of the Government Expenditure Multiplier’, *American Economic Journal: Macroeconomics* **3**(1), 1–35.