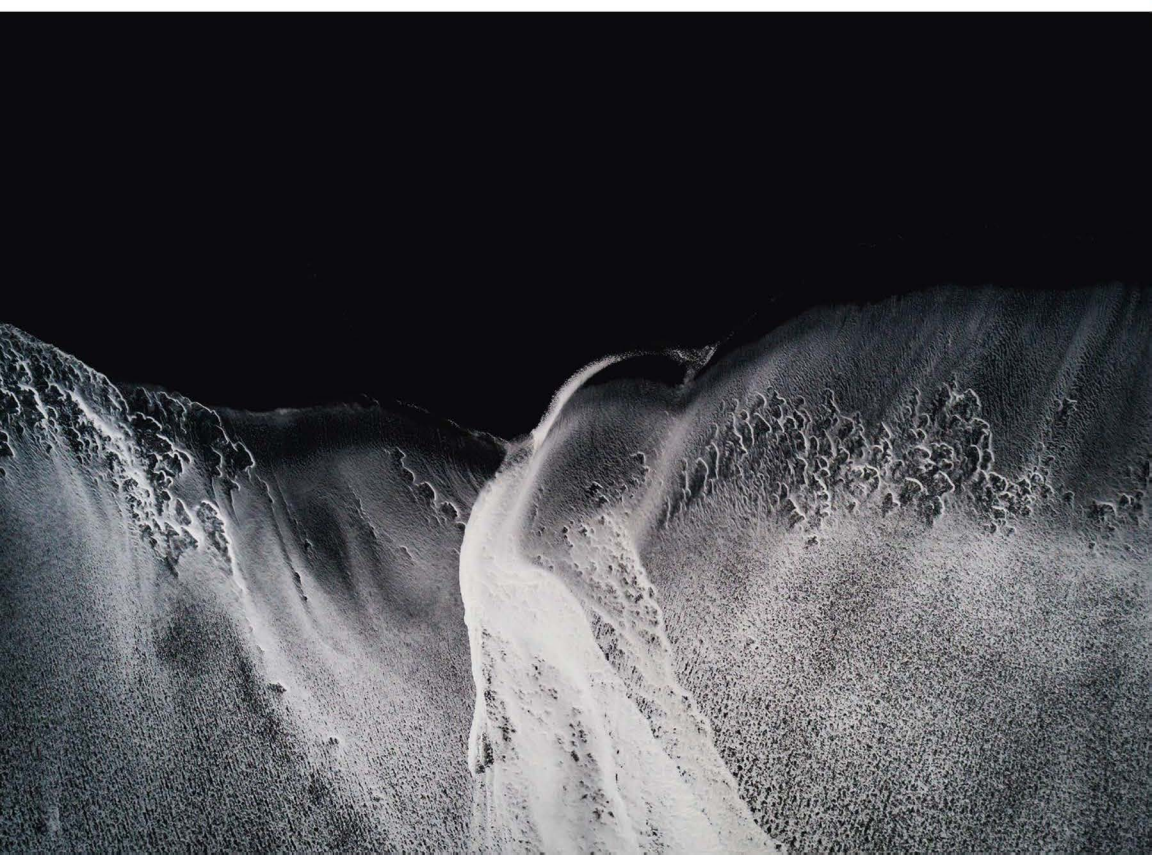


Why Machines Will Never Rule the World

Artificial Intelligence without Fear

**JOBST LANDGREBE
AND BARRY SMITH**



WHY MACHINES WILL NEVER RULE THE WORLD

Available from August 12, 2022

To order via publisher

To order via Amazon go here

‘It’s a highly impressive piece of work that makes a new and vital contribution to the literature on AI and AGI. The rigor and depth with which the authors make their case is compelling, and the range of disciplinary and scientific knowledge they draw upon is particularly remarkable and truly novel.’

Shannon Vallor, Edinburgh Futures Institute,
The University of Edinburgh

‘The alluring nightmare in which machines take over running the planet and humans are reduced to drudges is not just far off or improbable: the authors argue that it is mathematically impossible. While drawing on a remarkable array of disciplines for their evidence, the argument of Landgrebe and Smith is in essence simple. There can be no models and no algorithms of the complexity required to run machines which can come close to emulating human linguistic and social skills. Far from decrying AI, they laud its achievements and encourage its development; but they pour cold water on those who fail to recognise its inherent limitations. Compulsory reading for those who fear the worst, but also for those inadvertently trying to bring it about.’

Peter M. Simons FBA, Department of Philosophy
Trinity College Dublin

‘Just one year ago, Elon Musk claimed that AI will overtake humans “in less than five years”. Not so, say Landgrebe and Smith, who argue forcefully that it is mathematically impossible for machines to emulate the human mind. This is a timely, important, and thought-provoking contribution to the contemporary debate about AI’s consequences for the future of humanity.’

Berit Brogaard, Department of Philosophy
University of Miami

‘This book challenges much linguistically underinformed AI optimism, documenting many foundational aspects of language that are seemingly intractable to computation, including its capacity for vagueness, its interrelation with context, and its vast underpinning of implicit assumptions about physical, interpersonal, and societal phenomena.’

Len Talmy, Center for Cognitive Science
University at Buffalo

‘Interdisciplinarity is often touted, more or less successfully, as a cure for many disciplinary blind-spots. Landgrebe and Smith orchestrate a battery of arguments from philosophy, biology, computer science, linguistics, mathematics and physics in order to argue effectively and with great brio that AI has been oversold. The result is a model of how to bring together results from many different fields and argue for an important thesis. It shows that, along many dimensions, AI has not paid close enough attention to the foundations on which it rests.’

Kevin Mulligan, Department of Philosophy
University of Geneva

‘Why, after 50 years, are AI systems so bad at communicating with human beings? This book provides an entirely original answer to this question—but it can be summed up in just one short phrase: there is too much haphazardness in our language use. AI works by identifying patterns, for instance in dialogue, and then applying those same patterns to new dialogue. But every dialogue is different. The old patterns never work. AI must fail. If you care, read the book.’

Ernie Lepore, Department of Philosophy
Rutgers University

WHY MACHINES WILL NEVER RULE THE WORLD

The book's core argument is that an artificial intelligence that could equal or exceed human intelligence—sometimes called artificial *general* intelligence (AGI)—is for mathematical reasons impossible. It offers two specific reasons for this claim:

1. Human intelligence is a capability of a complex dynamic system—the human brain and central nervous system.
2. Systems of this sort cannot be modelled mathematically in a way that allows them to operate inside a computer.

In supporting their claim, the authors, Jobst Landgrebe and Barry Smith, marshal evidence from mathematics, physics, computer science, philosophy, linguistics, and biology, setting up their book around three central questions: What are the essential marks of human intelligence? What is it that researchers try to do when they attempt to achieve “artificial intelligence” (AI)? And why, after more than 50 years, are our most common interactions with AI, for example with our bank's computers, still so unsatisfactory?

Landgrebe and Smith show how a widespread fear about AI's potential to bring about radical changes in the nature of human beings and in the human social order is founded on an error. There is still, as they demonstrate in a final chapter, a great deal that AI can achieve which will benefit humanity. But these benefits will be achieved without the aid of systems that are more powerful than humans, which are as impossible as AI systems that are intrinsically “evil” or able to “will” a takeover of human society.

Jobst Landgrebe is a scientist and entrepreneur with a background in philosophy, mathematics, neuroscience, and bioinformatics. Landgrebe is also the founder of Cognotekt, a German AI company which has since 2013 provided working systems used by companies in areas such as insurance claims management, real estate management, and medical billing. After more than 10 years in the AI industry, he has developed an exceptional understanding of the limits and potential of AI in the future.

Barry Smith is one of the most widely cited contemporary philosophers. He has made influential contributions to the foundations of ontology and data science, especially in the biomedical domain. Most recently, his work has led to the creation of an international standard in the ontology field (ISO/IEC 21838), which is the first example of a piece of philosophy that has been subjected to the ISO standardization process.

Cover image: © Abstract Aerial Art / Getty Images

First published 2023

by Routledge

605 Third Avenue, New York, NY 10158

and by Routledge

4 Park Square, Milton Park, Abingdon, Oxon, OX14 4RN

Routledge is an imprint of the Taylor & Francis Group, an informa business

© 2023 Jobst Landgrebe and Barry Smith

The right of Jobst Landgrebe and Barry Smith to be identified as authors of this work has been asserted in accordance with sections 77 and 78 of the Copyright, Designs and Patents Act 1988.

All rights reserved. No part of this book may be reprinted or reproduced or utilised in any form or by any electronic, mechanical, or other means, now known or hereafter invented, including photocopying and recording, or in any information storage or retrieval system, without permission in writing from the publishers.

Trademark notice: Product or corporate names may be trademarks or registered trademarks, and are used only for identification and explanation without intent to infringe.

Library of Congress Cataloging-in-Publication Data

A catalog record for this title has been requested

ISBN: 978-1-032-31516-4 (hbk)

ISBN: 978-1-032-30993-4 (pbk)

ISBN: 978-1-003-31010-5 (ebk)

DOI: 10.4324/9781003310105

Typeset in Bembo

by Apex CoVantage, LLC

CONTENTS

<i>Foreword</i>	ix
1 Introduction	1
1.1 <i>The Singularity</i>	1
1.2 <i>Approach</i>	3
1.3 <i>Limits to the modelling of animate nature</i>	7
1.4 <i>The AI hype cycle</i>	9
1.5 <i>Why machines will not inherit the earth</i>	11
1.6 <i>How to read this book</i>	18
 PART I	
Properties of the human mind	21
2 The human mind	23
2.1 <i>Basic characteristics of the human mind</i>	23
2.2 <i>The mind-body problem: monism and its varieties</i>	24
3 Human and machine intelligence	37
3.1 <i>Capabilities and dispositions</i>	37
3.2 <i>Intelligence</i>	41
3.3 <i>AI and human intelligence</i>	48

4	The nature of human language	63
4.1	<i>Why conversation matters</i>	63
4.2	<i>Aspects of human language</i>	64
5	The variance and complexity of human language	74
5.1	<i>Conversations: an overview</i>	74
5.2	<i>Levels of language production and interpretation</i>	77
5.3	<i>Conversation contexts</i>	77
5.4	<i>Discourse economy: implicit meaning</i>	82
5.5	<i>Structural elements of conversation</i>	85
5.6	<i>How humans pass the Turing test</i>	88
6	Social and ethical behaviour	90
6.1	<i>Can we engineer social capabilities?</i>	91
6.2	<i>Intersubjectivity</i>	93
6.3	<i>Social norms</i>	95
6.4	<i>Moral norms</i>	98
6.5	<i>Power</i>	106

PART II

The limits of mathematical models 107

7	Complex systems	109
7.1	<i>Models</i>	109
7.2	<i>Computability</i>	115
7.3	<i>Systems</i>	117
7.4	<i>The scope of extended Newtonian mathematics</i>	119
7.5	<i>Complex systems</i>	124
7.6	<i>Examples of complex systems</i>	140
8	Mathematical models of complex systems	144
8.1	<i>Multivariate distributions</i>	144
8.2	<i>Deterministic and stochastic computable system models</i>	146
8.3	<i>Newtonian limits of stochastic models of complex systems</i>	149
8.4	<i>Descriptive and interpretative models of complex systems</i>	153
8.5	<i>Predictive models of complex systems</i>	158

- 8.6 *Naïve approaches to complex system modelling* 160
- 8.7 *Refined approaches* 180
- 8.8 *The future of complex system modelling* 187

Part III

The limits and potential of AI **193**

- 9 Why there will be no machine intelligence 195
 - 9.1 *Brain emulation and machine evolution* 195
 - 9.2 *Intentions and drivenness* 203
 - 9.3 *Consciousness* 205
 - 9.4 *Philosophy of mind, computation, and AI* 213
 - 9.5 *Objectifying intelligence and theoretical thinking* 214
- 10 Why machines will not master human language 217
 - 10.1 *Language as a necessary condition for AGI* 217
 - 10.2 *Why machine language production always falls short* 219
 - 10.3 *AI conversation emulation* 226
 - 10.4 *Mathematical models of human conversations* 235
 - 10.5 *Why conversation machines are doomed to fail* 242
- 11 Why machines will not master social interaction 245
 - 11.1 *No AI emulation of social behaviour* 245
 - 11.2 *AI and legal norms* 248
 - 11.3 *No machine emulation of morality* 250
- 12 Digital immortality 259
 - 12.1 *Infinity stones* 259
 - 12.2 *What is a mind?* 261
 - 12.3 *Transhumanism* 282
 - 12.4 *Back to Bostrom* 287
- 13 AI spring eternal 288
 - 13.1 *AI for non-complex systems* 288
 - 13.2 *AI for complex systems* 295
 - 13.3 *AI boundaries* 298
 - 13.4 *How AI will change the world* 301

viii Contents

<i>Appendix: Turbulence: Mathematical details</i>	302
<i>Glossary</i>	304
<i>References</i>	313
<i>Index</i>	335

FOREWORD

Rationale for this book

This book is about artificial intelligence (AI), which we conceive as the application of mathematics to the modelling (primarily) of the functions of the human brain. We focus specifically on the question of whether modelling of this sort has limits, or whether—as proposed by the advocates of what is called the ‘Singularity’—AI modelling might one day lead to an irreversible and uncontrollable explosion of ever more intelligent machines.

As concerns the current state of the art, AI researchers are, for understandable reasons, immensely proud of their amazing technical discoveries. It therefore seems obvious to all that there is an almost limitless potential for further, equally significant AI discoveries in the future.

Enormous amounts of funding are accordingly being invested in advancing the frontiers of AI in medical research, national defense, and many other areas. If our arguments hold water, then a significant fraction of this funding may be money down the drain. For this reason alone, therefore, it is probably no bad thing for the assumption of limitless potential for AI progress to be subjected to the sort of critical examination that we have here attempted.

The result, we must confess, is not always easy reading. To do our job properly, we found it necessary to move to a rather drastic degree beyond the usual disciplinary borders, drawing not merely on philosophy, mathematics, and computer science, but also on linguistics, psychology, anthropology, sociology, physics, and biology. In the “Approach” section of the Introduction we provide the rationale for this methodology and, where this is needed, for our choice of literature. In the “Glossary” (pp. 304ff.) we provide what we hope are reader-friendly definitions of the technical terms used in the main text.

We raise what we believe are powerful arguments against the possibility of engineering machines that would possess an intelligence that would equal or surpass that of humans. These arguments have immediate implications for claims, such as those of Elon Musk, according to whom AI could become ‘an immortal dictator from which we would never escape’. Relax. Machines will not rule the world.

At the same time, our arguments throw light on the question of which varieties of AI *are* achievable. In this respect we are fervent optimists, and one of us is indeed contributing to the creation of new AI products being applied in industry as this book is being written. In the final chapter of the book we outline some of the positive consequences of our arguments for practical applications of AI in the future.

A new affaire Dreyfus?

This book is concerned not with the tremendous successes of artificial intelligence along certain narrow lanes, such as text translation or image recognition. Rather, our concern is with what is called *general* AI and with the ability of computers to emulate, and indeed to go beyond, the general intelligence manifested by humans.

We will show that it is not possible (and this means: not ever) to engineer machines with a general cognitive performance even at the level of vertebrates such as crows. When we have presented our arguments in favour of this view to friendly audiences, the most common reaction has been that we are surely just repeating the mistake of earlier technology sceptics and that our arguments, too, are doomed to be refuted by the inevitable advances of AI in the future.

Hubert Dreyfus was one of the first serious critics of AI research. His book *What Computers Can't Do*, first published in 1972, explains that symbolic (logic-based) AI, which was at that time the main paradigm in AI research, was bound to fail, because the mental processes of humans do not follow a logical pattern. As Dreyfus correctly pointed out, the logical formulation of our thoughts is merely the end-product of a tiny fraction of our mental activities—an idea which seems to be undergoing a mild revival (Fjelland 2020).

In the third edition of his book, Dreyfus (1992) was still claiming that he had been right from the beginning. And so he was; though he did not provide the sorts of arguments we give in this book, which are grounded not on Heideggerian philosophy but on the mathematical implications of the theory of complex systems.

We start out from the assumption that all complex systems are such that they obey the laws of physics. However, we then show that for mathematical reasons we cannot use these laws to analyse the behaviours of complex systems because the complexity of such systems goes beyond our mathematical modelling abilities. The human brain, certainly, is a complex system of this sort; and while there are some of today's AI proponents who believe that the currently

fashionable AI paradigm of ‘deep neural networks’—connectionist as opposed to symbolic AI (Minsky 1991)—can mimic the way the brain functions, we will show in what follows that, again for mathematical reasons, this is not so, not only for deep neural networks but for any other type of AI software that might be invented in the future.

We define artificial general intelligence (AGI) as an AI that has a level of intelligence that is either equivalent to or greater than that of human beings or is able to cope with problems that arise in the world that surrounds human beings with a degree of adequacy at least similar to that of human beings (a precise definition of what this means is given in sections 3.3.3–3.3.4).

Our argument can be presented here in a somewhat simplified form as follows:

- A1. To build an AGI we would need technology with an intelligence that is at least comparable to that of human beings (from the definition of AGI just provided).
- A2. The only way to engineer such technology is to create a software emulation of the human neurocognitive system. (Alternative strategies designed to bring about an AGI without emulating human intelligence are considered and rejected in section 3.3.3 and chapter 12.)

However,

- B1. To create a software emulation of the behaviour of a system we would need to create a mathematical model of this system that enables prediction of the system’s behaviour.¹
- B2. It is impossible to build mathematical models of this sort for complex systems. (This is shown in sections 8.4–8.7.)
- B3. The human neurocognitive system is a complex system (see chapter 7).
- B4. *Therefore*, we cannot create a software emulation of the human neurocognitive system.

From (A2.) and (B4.) it now follows that:

- C. An AGI is impossible.

An analogy from physics

We conceive thesis (C.) to be analogous to the thesis that it is impossible to create a perpetual motion machine.

¹ The requirements which such predictions would need to satisfy are outlined in 3.3.4, with mathematical details in 7.1.1.4 and 8.5.

Someone might, now, argue that our current understanding of the laws of physics might one day be superseded by a new understanding, according to which a perpetual motion machine *is* possible after all.

And similarly someone might argue against the thesis of this book that our current understanding of the laws of *mathematics* might one day change. New ways of modelling complex dynamic systems may be discovered that would indeed allow the mathematical modelling of, for example, the workings of the human mind.

To see why this, too, is impossible, we show that it would have to involve discoveries even more far-reaching than the invention by Newton and Leibniz of the differential calculus. And it would require that those who have tried in the past to model complex systems mathematically, including Feynman (Feynman et al. 2010) and Heisenberg (Marshak et al. 2005, p. 76), were wrong to draw the conclusion that such an advance will never be possible. This conclusion was drawn not only by the best minds in the past. There exist also today no proposals even on the horizon of current physics or mathematics to surmount the obstacles to the modelling of complex systems identified in what follows.²

Acknowledgements

We thank especially Kevin Mulligan, Ingvar Johansson, and Johannes Weiß for their critical review of the manuscript. In addition, we thank Yuri Basharov, David Braun, Janna Hastings, David Hershenov, David Limbaugh, Jens Kipper, Bill Rapaport, Alan Ruttenberg, Len Talmy, Erik Thomsen, Leo Zaibert, and our anonymous reviewers for their valuable comments. We alone are responsible for any errors that remain.

The book is dedicated to our families.

² Two papers on *turbulence* written in 1941 by Kolmogorov (1941b, 1941a) raised some hopes that at least this complex system phenomenon could be understood mathematically. But these hopes, too, were abandoned, as we show in 8.7.1.1, with mathematical details provided in the Appendix.

1

INTRODUCTION

Since the research field of AI was first conceived in the late 1940s, the idea of an artificial general intelligence (AGI) has been put forward repeatedly. Advocates of AGI hold that it will one day be possible to build a computer that can emulate and indeed exceed all expressions of human-specific intelligence, including not only reasoning and memory, but also consciousness, including feelings and emotions, and even the will and moral thinking. This idea has been elaborated and cultivated in many different ways, but we note that AGI is far from being realised (Cohen et al. 2019).

Pennachin et al. (2007, p. 1) assert that ‘AGI appears by all known science to be quite possible. Like nanotechnology, it is “merely an engineering problem”, though certainly a very difficult one’. As we shall see, assertions of this sort are common in the literature on AGI. As the examination of this literature reveals, however, this is not because the thesis that there might be *fundamental* (which means: mathematical) obstacles to the achievement of AGI has been investigated and ruled out. Rather, it is because this possibility has simply been ignored.

1.1 The Singularity

Closely related to the concept of AGI is the idea of the ‘Singularity’, a term first applied in the AI field by Vinge (1993) and then popularised by Ray Kurzweil (2005), a pioneer of second generation AI. The term is used to refer to a point in time after which the development of AI technology becomes irreversible and uncontrollable,¹ with unknown consequences for the future of humanity.

1 See the Glossary for a definition. The decrease and increase of the values of a function close to a singularity is hyperbolic, which is why the term ‘singularity’ was repurposed to describe an

2 Introduction

These developments are seen by Kurzweil as an inevitable consequence of the achievement of AGI, and he too believes that we are approaching ever closer to the point where AGI will in fact be achieved (Eden et al. 2012). Proponents of the Singularity idea believe that once the Singularity is reached, AGI machines will develop their own will and begin to act autonomously, potentially detaching themselves from their human creators in ways that will threaten human civilisation (Weinstein et al. 2017). The Singularity idea features widely in debates around AGI, and it has led scientists and philosophers (as well as politicians and science fiction authors) to explore the ethical implications of this idea, for instance by postulating norms and principles that would need to be somehow built into AGIs in the future in order to counteract their potentially negative effects. Visions are projected according to which AI machines, because of their superior ethical and reasoning powers, will one day supplant existing human-based institutions such as the legal system and democratic elections. Moor (2009) talks of ‘full ethical agents’, which are for him the highest form of machine morality, though at the same time he believes that agents of this sort will not appear any time soon.

Visions of AGI have been associated with lavishly promoted ideas according to which we are moving towards a time when it will be possible to find cures for human diseases by applying ever more powerful computers to ‘big’ biological data. In the wake of the successful sequencing of the human genome and the related advent of DNA microarrays and of mass spectrometry, mass cytometry, and other sophisticated methods for performing mass assays of the fundamental components of organic matter, considerable funds have been invested in big ‘ome’ (transcriptome, proteome, connectome) and similar projects. Yet at the same time, even after some 20 years of research (in which both of us have participated), there is a haunting awareness of the paucity of results with significant implications for human health and disease achieved along these lines.

But we already know enough from what we have learned in the foregoing that the Singularity will not arise, given that

- D1. Such a Singularity would require the engineering of an AI with the capability to engineer another machine more intelligent than itself.
- D2. The exercise of this capability, at least in its early stages, would require assistance from and thus persuasive communication with human beings in bringing about the realisation of a series of highly complex goals (section 12.2.5).
- D3. Only an AGI could succeed in the realisation of such goals (section 3.3).
- D4. Therefore, the Singularity would require an AGI.

imagined rapid realisation of ‘superintelligence’ once a certain point is reached in the development of AI.

Now, however, using the proposition C (from p. xi), that an AGI is impossible and (D4.) we can infer:

E. The Singularity is impossible.

1.2 Approach

In the pages that follow we will analyse the scope and potential of AI in the future and show why the dark scenarios projected by Nick Bostrom (Bostrom 2003), Elon Musk, and others will not be realised. First, however, we set forth the sources and methods we will use. The reader interested more in content than methods may accordingly skip this section and proceed directly to page 9.

1.2.1 *Realism*

The overarching doctrine which binds together the different parts of this book is a non-reductivist commitment to the existence of physical, biological, social, and mental reality, combined with a realist philosophy about the world of common sense or in other words the world of ‘primary theory’ as expounded by the anthropologist Robin Horton (1982).² Thus we hold that our common-sense beliefs—to the effect that we are conscious, have minds and a will, and that we have access through perception to objects in reality—are both true and consistent with the thesis that everything in reality is subject to the laws of physics. To understand how scientific realism and common-sense realism can be reconciled, we need to take careful account of the way in which systems are determined according to the granularity of their elements. (This book is essentially a book about systems, and about how systems can be modelled scientifically.)

Central components of our realist view include the following:³

1. The universe consists of matter, which is made of elementary particles: quarks, leptons, and bosons.⁴ Entities of various supernumerary sorts exist in
- 2 Those parts of primary theory which concern human mental activities—for example thinking, believing, wanting—correspond to what elsewhere in this book we refer to as the common-sense ontology of the mental, and which is (sometimes disparagingly) referred to by analytic philosophers as ‘folk psychology’.
- 3 The view in question is inspired by Johansson’s ‘irreductive materialism’ (Johansson 2012). It is similar also to the liberal naturalism expounded by De Caro (2015), which attempts ‘to reconcile common sense and scientific realism in a non-Cartesian pluralist ontological perspective’ and which explicitly includes as first-class entities not only material things such as you and me but also entities, such as debts, that have a history and yet are non-physical. See also Haack (2011).
- 4 This is the current view, which is likely to change as physics progresses. Changes on this level will not affect any of the arguments in this book.

4 Introduction

those parts of the universe where animals and human beings congregate. (See item 6 in this list.)

2. All interactions of matter are governed by the four fundamental forces (interactions) described by physics (electromagnetism, gravity, the strong interaction, the weak interaction), yielding all of the phenomena of nature that we perceive, including conscious human beings.
3. Fundamental entities should not be multiplied without necessity.⁵ No counterpart of this maxim applies, however, to the vast realms of entities created as the products of human action and of human convention. The kilometre exists; but so also does the Arabic mile, the Burmese league, and the Mesopotamian cubit—and so do all the ‘ordinary objects’ discussed by Lowe (2005).
4. We thus hold that the totality of what exists can be viewed from multiple different, mutually consistent granular perspectives. From one perspective, this totality includes quarks, leptons, and the like. From another perspective it includes organisms, portions of inorganic matter, and (almost) empty space.⁶
5. At all levels we can distinguish both types and instances of these types. In addition we can distinguish at all levels continuants (such as molecules) and occurrents (such as protein folding processes).
6. Some organisms, for instance we ourselves and (we presume) you also, dear reader, are conscious. Conscious processes, which always involve an observer, are what we shall call emanations from complex systems (specifically: from organisms).⁷ When viewed from the outside, they can be observed only indirectly (they can, though, be viewed directly via introspection).
7. In the world made by conscious organisms, there exist not only tapestries and cathedrals, dollar bills and drivers’ licenses, but also social norms, poems, nation states, cryptocurrencies, Olympic records, mathematical equations, and data.

1.2.2 General remarks on methods

To answer the question of whether AGI is possible, we draw on results from a wide range of disciplines, in particular on technical results of mathematics and theoretical physics, on empirical results from molecular biology and other hard science domains, and (to illustrate the implications for AI of our views when

5 This is Schaffer’s Laser (Schaffer 2015).

6 On the underlying theory of granular partitions see Bittner et al. (2001, 2003).

7 We adapt the term ‘emanation’ from its usage in physics to mean any type of electromagnetic radiation, for example, thermal radiation, or other form of energy propagation (for example, sound), which is observable, but which is produced by a system process which is hidden (Parzen 2015) (cannot be observed); for a detailed definition see 2.1.

applied to the phenomenon of human conversation) on descriptive results from linguistics. In addition to the standard peer-reviewed literature, our sources in these fields include authoritative textbooks—above all the *Introduction to the Theory of Complex Systems* by Thurner et al. (2018)—and salient writings of Alan Turing, Jürgen Schmidhuber, and other leaders of AI research.

We also deal with contributions to the Singularity debate made by contemporary philosophers, above all David Chalmers (see section 9.1), and—by way of light relief—on the writings of the so-called transhumanists on the prospects for what they call ‘digital immortality’ (chapter 12).

1.2.3 Formal and material ontology

For reasons set forth in (Smith 2005), most leading figures in the early phases of the development of analytic philosophy adhered to an overly simplistic view of the world, which left no room for entities of many different sorts. Thus they developed assays of reality which left no room for, *inter alia*, norms, beliefs, feelings, values, claims, obligations, intentions, dispositions, capabilities, communities, societies, organisations, authority, energy, works of music, scientific theories, physical systems, events, natural kinds, and entities of many other sorts. Many analytic philosophers embraced further an overly simplistic view of the mind/brain, often taking the form of an assertion to the effect that the mind operates like (or indeed that it is itself) a computer. Computer scientists often think the opposite, namely that a computer ‘acts’ like the human brain and that the differences between these two types of machines will one day be overcome with the development of AGI.⁸ But a computer does not *act*, and the human brain is not a machine, as we shall see in the course of this book.

In recent years, on the other hand, analytic philosophers have made considerable strides in expanding the coverage domain of their ontologies, in many cases by rediscovering ideas that had been advanced already in other traditions, as, for example, in phenomenology. They continue still, however, to resist the idea of a comprehensive realist approach to ontology. This reflects a more general view, shared by almost all analytic philosophers, to the effect that philosophy should not seek the sort of systematic and all-encompassing coverage that is characteristic of science, but rather seek point solutions to certain sorts of puzzles, often based on ‘reduction’ of one type of entity to another (Mulligan et al. 2006).

There is however one group of philosophers—forming what we can call the school of realist phenomenologists (Smith 1997)—who embraced this sort of comprehensive realist approach to ontology, starting out from the methodological guidelines sketched by Husserl in his *Logical Investigations* (Husserl 2000).

8 Mathematicians who have to deal with computers take the view that computers are mere servants (*Rechenknechte*) which exist merely to perform calculations.

6 Introduction

Like Frege, and in contrast to, for example, Heidegger or Derrida, the principal members of this school employed a clear and systematic style. This is especially true of Adolf Reinach, who anticipated in his masterwork of 1913—‘The *A Priori* Foundations of the Civil Law’ (Reinach 2012)—major elements of the theory of speech acts reintroduced by Austin and Searle in the 1960s.⁹ Husserl’s method was applied by Reinach to the ontology of law, by Roman Ingarden to the ontology of literature and music, and to the realm of human values in general (Ingarden 1986, 1973, 2013, 2019).¹⁰

Two other important figures of the first generation of realist phenomenologists were Max Scheler and Edith Stein, who applied this same method, respectively, to ethics and anthropology on the one hand, and to social and political ontology on the other. Ingarden is of interest also because he established a branch of realist phenomenology in Poland.¹¹

The most salient members of the second generation of this school were Nicolai Hartmann, whose systematisation of Scheler’s ideas on the ontology of value will concern us in chapter 6, and Arnold Gehlen, a philosopher, sociologist, and anthropologist working in the tradition of the German school of philosophical anthropology founded by Scheler.¹²

For questions of perception, person, act, will, intention, obligation, sociality, and value, accordingly, we draw on the accounts of the realist phenomenologists, especially Reinach, Scheler, and Hartmann. This is both because their ideas were groundbreaking and because their central findings remain valid today. For questions relating to human nature, psychology, and language, we use Scheler and Gehlen, though extended by the writings of J. J. Gibson and of the ecological school in psychology which he founded.

1.2.3.1 An ecological approach to mental processing

A subsidiary goal of this book is to show the relevance of environments (settings, contexts) to the understanding of human and machine behaviour, and in this we are inspired by another, less familiar branch of the already mentioned ecological

9 (Smith 1990; Mulligan 1987) It is significant that Reinach was one of the first German philosophers to take notice, in 1914, of the work of Frege (Reinach 1969).

10 Ingarden’s massive three-volume work on formal, existential, and material ontology is only now being translated into English. This work, along with Husserl’s *Logical Investigations*, provides the foundation for our treatment of the principal ontological categories (such as continuant, occurrent, role, function, and disposition) that are used in this book.

11 One prominent member of the latter was Karol Wojtyła, himself an expert on the ethics and anthropology of Scheler (Wojtyła 1979), and it is an interesting feature of the school of phenomenological realists, perhaps especially so when we come to gauge the value of its contribution to ethics, that two of its members—namely Stein (St. Teresa Benedicta of the Cross) and Wojtyła (St. John Paul II)—were canonised.

12 Gehlen was one of the first to explore theoretically the question of the nature and function of human language from the evolutionary perspective in his main work *Man. His Nature and Place in the World* (Gehlen 1988), first published in German in 1940 (Gehlen, 1993 [1940]).

school in psychology, which gave rise to a remarkable volume entitled *One Boy's Day: A Specimen Record of Behavior* by Barker et al. (1951). This documents over some 450 pages all the natural and social environments¹³ occupied by Raymond, a typical 7-year-old schoolboy on a typical day (April 26, 1949) in a typical small Kansas town.

Barker shows how each of the acts, including the mental acts, performed by Raymond in the course of the day is tied to some specific environment, and he reminds us thereby that any system designed to emulate human mental activity inside the machine will have to include a subsystem (or better: systems) dealing with the vast and ever-changing totality of environments within which such activity may take place.

1.2.3.2 Sociology and social ontology

A further feature of the analytic tradition in philosophy was its neglect of sociality and of social interaction as a topic of philosophical concern. Matters began to change with the rediscovery of speech acts in the 1960s by Austin and Searle, a development which has in recent years given rise to a whole new sub-discipline of analytic social ontology, focusing on topics such as ‘shared’ or ‘collective agency’ (Ziv Konzelmann 2013). Many 20th-century analytic philosophers, however, have adopted an overly simplistic approach also to the phenomena of sociology and social ontology¹⁴, and to counteract this we move once again outside the analytic mainstream, drawing first on the classical works of Max Weber and Talcott Parsons, on the writings on sociality of Gibson and his school (Heft 2017), and also on contemporary anthropologists, especially those in the tradition of Richerson et al. (2005).

1.3 Limits to the modelling of animate nature

It is well known that the utility of science depends (in increasing order) on its ability to *describe*, *explain*, and *predict* natural processes. We can *describe* the foraging behaviour of a parrot, for example, by using simple English. But to *explain* the parrot's behaviour we need something more (defined in section 7.1.1.4). And for *prediction* we need causal models, and these causal models must be mathematical.

13 These are called ‘settings’ in Barker's terminology (Barker 1968), (Smith 2000). Schoggen (1989) gives an overview of the work of Barker and his school, and Heft (2001) describes the philosophical background to Barker's work and his relations to the broader community of ecological psychologists.

14 The legal philosopher Scott Shapiro points to two major limitations of much current work on shared agency by analytic philosophers: ‘first, that it applies only to ventures characterised by a rough equality of power and second, that it applies only to small-scale projects among similarly committed individuals’ (Shapiro 2014). Examples mentioned by Shapiro are: singing a duet and painting a house together.

For this reason, however, it will prove that the lack of success in creating a general AI is not, as some claim, something that can and will be overcome by increasing the processing power and memory size of computers. Rather, this lack of success flows not only from a lack of understanding of many biological matters, but also from an inability to create mathematical models of how brains and other organic systems work.

In biology, valid mathematical models aiming at system explanation are hard or impossible to obtain, because there are in almost every case very large numbers of causal factors involved, making it well-nigh impossible to create models that would be predictive.¹⁵ The lack of models is most striking in neuroscience, which is the science dealing with the physical explanation of how the brain functions in giving rise not only to consciousness, intelligence, language and social behaviour but also to neurological disorders such as dementia, schizophrenia, and autism.

The achievement of AGI would require models whose construction would presuppose solutions of some of the most intractable problems in all of science (see sections 8.5 to 8.8). It is thus disconcerting that optimism as concerns the potential of AI has been most vigorous precisely in the promotion of visions relating to enhancement, extension, and even total emulation of the human brain. We will see that such optimism rests in part on the tenacity of the view according to which the human brain is a type of computer (a view still embraced explicitly by some of our best philosophers), which on closer analysis betrays ignorance not only of the biology of the brain but also of the nature of computers. And for a further part it rests on naïve views as to the presumed powers of deep neural networks to deal with emanations from complex systems in ways which go beyond what is possible for traditional mathematical models.

1.3.1 Impossibility

Throughout this book, we will defend the thesis that it is impossible to obtain what we shall call *synoptic* and *adequate* mathematical models of complex systems, which means: models that would allow us to engineer AI systems that can fulfill the requirements such systems must satisfy if they are to emulate intelligence.

Because the proper understanding of the term *impossible* as it is used in this sentence is so important to all that follows, we start with an elucidation of how we are using it. First, we use the term in three different senses, which we refer to as *technical*, *physical*, and *mathematical* impossibility, respectively.

To say that something is *technically impossible*—for example, controlled nuclear fusion with a positive energy output—is to draw attention to the fact that it is

15 There are important exceptions in some specific subfields, for example models of certain features of monogenetic and of infectious diseases, or of the pharmacodynamics of antibiotics. See 8.4.

impossible *given the technology we have today*. We find it useful to document the technically impossible here only where (as is all too often) proponents of trans-humanist and similar concepts seek to promote their ideas on the basis of claims which are, and may for all time remain, technically impossible.¹⁶

To say that something is *physically impossible* is to say that it is impossible because it would contravene the laws of physics. To give an example: in highly viscous fluids (low Reynolds numbers), no type of swimming object can achieve net displacement (this is the scallop theorem [Purcell 1977]).¹⁷

To speak of *mathematical impossibility*, finally, is to assert that a solution to some mathematically specified problem—for example, an analytical solution of the n -body problem (see p. 189) or an algorithmic solution of the halting problem (see section 7.2)—cannot be found; not because of any shortcomings in the data or hardware or software or human brains, but rather for *a priori* reasons of mathematics. This is the primary sense in which we use the term *impossible* in this book.

1.4 The AI hype cycle

Despite the lack of success in brain modelling, and fired by a naïve understanding of human brain functioning, optimism as to future advances in AI feeds most conspicuously into what is now called ‘transhumanism’, the idea that technologies to enhance human capabilities will lead to the emergence of new ‘post-human’ beings. On one scenario, humans themselves will become immortal because they will be able to abandon their current biological bodies and live on, forever, in digital form. On another scenario, machines will develop their own will and subdue mankind into slavery with their superintelligence while they draw on their immortality to go forth into the galaxy and colonise space.

Speculations such as this are at the same time fascinating and disturbing. But we believe that, like some earlier pronouncements from the AI community, they must be taken with a heavy pillar of salt, for they reflect enthusiasm triggered by successes of AI research that does not factor in the fact that these successes have been achieved only along certain tightly defined paths.

In 2016 it became known that the company DeepMind had used AI in their AlphaGo automaton to partially solve the game of Go.¹⁸ Given the complexity of the game, this must be recognised as a stunning achievement. But it is an achievement whose underlying methodology can be applied only to a narrow set of problems. What it shows is that, in certain completely rule-determined

16 For example: ‘Twenty-first-century software makes it technologically possible to separate our minds from our biological bodies.’ (Rothblatt 2013, p. 317). We return to this example in chapter 11.

17 We return to this example in section 12.3.3.1.

18 Solving a game *fully* means ‘determining the final result in a game with no mistakes made by either player’. This has been achieved for some games, but not for either chess or GO (Schaeffer et al. 2007).

confined settings with a low-dimensional phase space such as abstract games, a variant of machine learning known as reinforcement learning (see 8.6.7.3) can be used to create algorithms that outperform humans. Importantly, this is done in ways that do not rely on any knowledge on the part of the machine of the rules of the games involved. This does not, however, mean that DeepMind can ‘discover’ the rules governing just *any* kind of activity. DeepMind’s engineers provided the software with carefully packaged examples of activity satisfying just this set of rules and allowed it to implicitly generate new playing strategies not present in the supplied examples by using purely computational adversarial settings (two algorithms playing against each other).¹⁹ The software is not in a position to go out and identify packages of this sort on its own behalf. It is cognizant neither of the rules nor of the new strategies which we, its human observers, conceive it to be applying.

Yet the successes of DeepMind and of other AI engineering virtuosi have led once again to over-generalised claims on behalf of the machine learning approach, which gave new energy to the idea of an AI that would be *general* in the same sort of way that human intelligence is general, to the extent that it could go forth into the world unsupervised and achieve ever more startling results.

Parallel bursts of enthusiasm have arisen also in connection with the great strides made in recent years in the field of image recognition. But there are already signs that there, too, the potential of AI technology has once again been overestimated (Marcus 2018; Landgrebe et al. 2021).

Why is this so? Why, in other words, is AI once again facing a wave of dampening enthusiasm²⁰ representing the third major AI sobering episode after the mid-1970s and late 1980s, both of which ended in AI winters? There are, certainly, many reasons for this cyclical phenomenon. One such reason is that genuine advances in AI fall from public view as they become embedded in innumerable everyday products and services. Many contributions of working (narrow) AI are thereby hidden. But a further reason is the weak foundation of AI enthusiasm itself, which involves in each cycle an initial exaggeration of the potential of AI under the assumption that impressive success along a single front will be generalisable into diverse unrelated fields (taking us, for instance, from *Jeopardy!* to curing cancer [Strickland 2019]).²¹

19 This is an excellent example of the use of synthetic data which is appropriate and adequate to the problem at hand.

20 This is not yet so visible in academia and in the public prints; but it is well established among potential commercial users, for example Bloomberg is clearly indicating this in fall 2021 (<https://www.bloomberg.com/opinion/articles/2021-10-04/artificial-intelligence-ain-t-that-smart-look-at-tesla-facebook-healthcare>). Further documentation of a breakdown in AI enthusiasm is provided by Larson (2021, pp. 74ff.).

21 The consequences of this assumption are thoroughly documented by Larson (2021), who explains why it is so difficult to re-engineer AI systems built for one purpose to address a different purpose.

Assumptions of this sort are made, we believe, because AI enthusiasts often do not have the interdisciplinary scientific knowledge that is needed to recognise the obstacles that will stand in the way of envisaged new AI applications. It is part of our aim here to show that crucial lessons concerning both the limits and the potential of AI can be learned through application of the right sort of interdisciplinary knowledge.

1.5 Why machines will not inherit the earth

In this book, we will argue that it is not an accident that so little progress has been made towards AGI over successive cycles of AI research. The lack of progress reflects, rather, certain narrow, structural limits to what can be achieved in this field, limits which we document in detail.

The human tendency to anthropomorphise is very powerful (Ekbja 2008). When our computer pauses while executing some especially complicated operation, we are tempted to say, ‘it’s thinking’. But it is not so. The processes inside the computer are physical through and through and, as we shall see in section 7.2, limited to certain narrowly defined operations defined in the 1930s by Turing and Church. The fact that we describe them in mental terms turns on the fact that the computer has been built to imitate (*inter alia*) operations that human beings would perform by using their minds. We thus impute the corresponding mental capabilities to the machine itself, as we impute happiness to a cartoon clown.

As Searle (1992) argued, computation has a physical, but no mental, reality because the significance that we impute to what the computer does (what we perceive as its mental reality) is observer dependent. If we take away all observers, then only the physical dimensions of its operations would remain. To see what is involved here, compare the difference between a dollar bill as a piece of paper and a dollar bill as money. If we take away all observers, then only the former would remain, because the latter is, again, ‘observer dependent’. While we *impute* consciousness to computers, we ourselves *are* conscious.

Computers also will not be able to *gain* consciousness in the future, since as we will show, whatever remarkable feats they might be engineered to perform, the aspect of those feats we are referring to when we ascribe consciousness or mentality to the computer will remain forever a product of observer dependence.

As we discuss in more detail in chapter 9, any process that machines can execute in order to *emulate* consciousness would have to be such that the feature of consciousness that is imputed to it would be observer dependent. From the fact that a certain green piece of paper is imputed to have the observer-dependent value of one dollar, we can infer with high likelihood that this piece of paper has the value of one dollar. As we will show in chapter 9, no analogous inference is possible from the fact that a process in a

12 Introduction

machine is imputed to have the feature of consciousness (or awareness, or excitedness, or happiness, or wariness, or desire). And thus we will never be able to create an AI with the faculty of consciousness in the sense in which we understand this term when referring to humans or animals.

But if we cannot *create* consciousness in the machine, the machine might still surely be able to *emulate* consciousness? This, and the related question of the limits of computer emulation of human *intelligence*, is one of the main questions addressed in this book.

1.5.1 The nature of the human mind

As mentioned earlier, in the eyes of many philosophers working in the theory of mind, the mind works like a universal Turing machine: it takes sensory inputs and processes these to yield some behavioural output. But how does it do this? When it comes to answering this question, there are three major schools: connectionists (Elman et al. 1996), computationalists (Fodor et al. 1988), and the defenders of hybrid approaches (Garson 2018). All of them think that the mind works like a machine. Connectionists believe that the mind works like a neural network as we know it from what is called ‘artificial neural network’ research.²² Computationalists believe that the mind operates by performing purely formal operations on symbols.²³

Most important for us here are the hybrid approaches as pioneered by Smolensky (1990), which seek to reconcile both schools by proposing that neural networks can themselves be implemented by universal Turing machines, an idea that was indeed technically realised in 2014 by Graves et al. (2014), who used a neural network to implement a classical von Neumann architecture. Their result proves that the initial dispute between connectionists and computationalists was mathematically nonsensical, because it shows that a universal Turing machine can implement *both* connectionist *and* symbolic logic. In other words, both types of computational procedures can be expressed using the basic recursive functions defined by Alonzo Church (1936). That both symbolic and perceptron (neural network) logic are Turing-computable has been known to mathematicians and computer scientists since the 1950s, and this makes the whole debate look naïve at best.

However there is a deeper problem with all ideas according to which the functioning of the mind (or brain) can be understood and modelled as the functioning of one or other type of machine, namely that such ideas are completely detached from the standpoint of biology and physics.²⁴ We will show that the

22 An artificial neural network is an implicit mathematical model generated by constraining an optimisation algorithm using training data and optimisation parameters; see further in chapter 8.

23 The relation between these two schools from an AI research perspective is summarised by Minsky (1991), who made important contributions to connectionist AI.

24 We shall see in detail why this is so in chapter 2 and section 9.4 and 12.2.

mentioned alternatives fail, because the mind (or brain) does not operate like a machine, and those who propose that it does do not acknowledge the results of neuroscience. For while we do not know how the mind works exactly, what we do know from neuroscience is that the workings of the mind resist mathematical modelling.²⁵ Therefore, we cannot emulate the mind using a machine, nor can we engineer other non-machine kinds of complex systems to obtain so-far undescribed kinds of non-human intelligence, and we will understand why in the course of this book.

1.5.2 *There will be no AGI*

The aim of AGI research, and of those who fund it, is to obtain something useful, and this will imply that an AGI needs to fulfill certain requirements—broadly, that it is able to cope with the reality in which humans live with a level of competence that is at least equivalent to that of human beings (see sections 3.3.3 and 3.3.4). We show that this is not possible, because there is an upper bound to what can be processed by machines. This boundary is set, not by technical limitations of computers, but rather by the limits on the possibilities for mathematical modelling.

There can be no ‘artificial general intelligence’, and therefore no ‘Singularity’ and no ‘full ethical agents’, because all of these would lie way beyond the boundary of what is even in principle achievable by means of a machine.

As we show at length in chapters 7 and 8, this upper bound is not a matter of computer storage or processing speed—factors which may perhaps continue to advance impressively from generation to generation. Rather, it is a matter of mathematics, at least given the presupposition that the aim is to realise AGI using computers.²⁶ For every computational AI is, after all, just a set of Turing-computable mathematical models taking input and computing output in a deterministic manner. Even ‘self-learning’ stochastic models behave deterministically once they have been trained and deployed to operate in a computer.

We shall deal in this book with all types of models that can currently be used to create computer-based AI systems, and we present each in great detail in chapter 8. Our arguments are completely general; they apply to all these types of models in the same way, and we are confident that these same arguments will apply also to any new types of models that will be developed in the future. At the same time, however, we note that these arguments potentially provide a boon to our adversaries, who can use them as a guide to the sorts of obstacles that would need to be overcome in order to engineer an AI that is both more useful than what we already have, and feasible from the point of view of engineering.

25 The 1,696 pages of *Principles of Neural Science* by Kandel et al. (2021), which is the gold standard textbook in the field and summarises some 100 years of neuroscientific research, contain almost no mathematics. And this is not about to change.

26 Other approaches, for example resting on the surgical enhancement of human brains, are considered in section 12.2.4.

1.5.3 *Prior arguments against artificial human-level intelligence*

We are not alone in believing that the idea of AGI, and of the Singularity which will follow in its wake, is at least to some degree a reflection of overconfidence among some members of the AI research community, and a number of AI proponents have expressed views which anticipate at least part of what we have to say here. For example, and most usefully, Walsh (2017). Walsh does indeed believe that AI with human-level intelligence will be achieved within the next 30–40 years; but he holds at the same time that there are a number of reasons why the Singularity will not arise:

1. intelligence is much more than thinking faster,
2. humans may not be intelligent enough to design superintelligence,
3. there is no evidence at all that an ML (machine learning) algorithm which achieves human level intelligence would thereby somehow proceed to becoming *more* intelligent (what David Chalmers [2010] calls ‘AI+’),
4. there are diminishing returns from AI performance, so that performance improvements to the level of a Singularity may be stymied,
5. systems have physical limits, and there are ‘empirical laws that can be observed emerging out of complex systems’. Intelligence itself as ‘a complex phenomenon may also have such limits that emerge from this complexity. Any improvements in machine intelligence, whether it runs away or happens more slowly, may run into such limits’ (op. cit., p. 61)²⁷,
6. the computational complexity required to go beyond human level intelligence may not be physically realisable.

Other important reservations concerning the possibility of the Singularity and the limits of AI in general have been brought forward by:

- Yann LeCun, who addresses the claims made by some researchers concerning an anticipated exponential growth in the powers of AI and points out that,

the first part of a sigmoid looks a lot like an exponential. It’s another way of saying that what currently looks like exponential progress is very likely to hit some limit—physical, economical, societal—then go

²⁷ We note in passing that this may be one reason for the apparent contradiction between the lack of evidence for extraterrestrial civilisations and various high estimates for their probability (Fermi’s paradox). Why do we see no evidence of alien superintelligences? Because the same limits to the increase in power of AI would (we believe) apply also to any technology developed by other intelligent life forms. This has implications also for the idea, favoured by Elon Musk, according to which the world in which we live is a simulation.

through an inflection point, and then saturate. I'm an optimist, but I'm also a realist.

(LeCun 2015)

- Yoshua Bengio, who makes the point that it is impossible to teach machines moral judgement: 'People need to understand that current AI—and the AI that we can foresee in the reasonable future—does not, and will not, have a moral sense or moral understanding of what is right and what is wrong' (Ford 2018, p. 31).
- Judea Pearl, who emphasises that the currently fashionable stochastics-based 'opaque learning machines' (Pearl 2020) lack an important feature of human-level intelligence in that they cannot answer questions related to causality and thus they cannot develop understanding about how things work. Pearl does not exclude the possibility of creating an AGI. He insists only that 'human-level AI cannot emerge solely from model-blind learning machines; it requires the symbiotic collaboration of data and models'.²⁸
- Brian Cantwell Smith (2019), who states that

neither deep learning nor other forms of second-wave AI, nor any proposals yet advanced for third-wave, will lead to genuine intelligence. Systems currently being imagined will achieve formidable reckoning prowess, but human-level intelligence and judgment, honed over millennia, is of a different order.

(*The Promise of Artificial Intelligence*, Introduction)

- For Shannon Vallor:

Those who are predicting an imminent 'rise of the robots' or an 'AI singularity,' in which artificially intelligent beings decide to dispense with humanity or enslave us, in my view serve as an unhelpful distraction from the far more plausible but less cinematic dangers of artificial intelligence. These mostly involve unexpected interactions between people and software systems that aren't smart enough to avoid wreaking havoc on complex human institutions, rather than robot overlords with 'superintelligence' dwarfing our own.

(Vallor 2016, p. 250)

- Steven Pinker argues that the threats to freedom in the future lie not so much in the advent of any putative Singularity, but rather in the way

²⁸ We shall see what this means in chapter 8; essentially, that the AI we can realise is determined by us.

societies choose to use technology. He draws what we shall recognise later as the crucial distinction between intelligence and motivation. And while he is ready to accept that we might technically realise something like the former, he points out that ‘there is no law of complex systems that says that intelligent agents must turn into ruthless megalomaniacs’. He also clearly sees that intelligence is not a boundless continuum with no limits to its potency (a point which we discuss in chapter 12); he recognises that stochastic models do not create knowledge and that AI is just a technology like any other, which is ‘constantly tweaked for efficacy and safety’ (Pinker 2020).

- Darwiche stated in (2018) that

what just happened in AI is nothing close to a breakthrough that justifies worrying about doomsday scenarios.... The current negative discussions by the public on the AI Singularity, also called super intelligence, can only be attributed to the lack of accurate ... characterisations of recent progress.
(p. 66)

Although such expressions of AGI pessimism are rarely encountered in the public prints, we suspect that the passages just cited in fact represent the views of a majority of AI experts. But they are all arguments to the effect that the Singularity *might not happen*. Here, in contrast, we will present arguments to the effect that already the creation of AI with an intelligence comparable to that of a human being is *impossible to achieve*, and thus that the Singularity, too, *will never happen*.

One notable exception is François Chollet, who argues that the idea of an ‘intelligence explosion comes from a profound misunderstanding of both the nature of intelligence and the behavior of recursively self-augmenting systems’.²⁹ His main hypotheses are:

- that AGI theorists employ an erroneous definition of intelligence,
- that human intelligence depends on innate dispositions, on interaction with the environment (sensorimotor affordances), and on socialisation; it can be exemplified only by a human being who is part of a society,
- that complex real-world systems cannot be modelled using the Markov assumption.

Chollet points out further that the ‘no free lunch theorem’ (8.6.6.3) implies that if ‘intelligence is a problem-solving algorithm, then it can only be understood with respect to a specific problem’. In sum, Chollet defends a view of AGI very much in the spirit of this book, but he provides only limited arguments on behalf of this view, as contrasted with the sort of detailed discussion that we present in

²⁹ Retrieved at <https://medium.com/@francois.chollet/the-impossibility-of-intelligence-explosion-5be4a9eda6ec>

chapters 7 and 8. For we will demonstrate that it is impossible to create the sorts of mathematical models even of vertebrate intelligence that would be needed in order to engineer its counterpart in a computer.

1.5.3.1 On abduction

A more recent, and for our purposes more significant, contribution to the debate on AGI is the book by Larson (2021). Larson hedges his bets as to whether human-level AI will or will not be achieved in the future, though he points out that ‘no one has the slightest clue how to build an artificial general intelligence’ (p. 275). But he emphasises that we do not have today, even on the horizon, anything like human-level AI. This is so, he argues, because of the current dominance of the assumption that the arrival of AGI is only a matter of time, because ‘we have already embarked on the path that will lead to human-level AI, and then superintelligence’. He calls this assumption ‘the myth of AI’, arguing that the assumption of inevitability is so deeply entrenched that—as we ourselves have discovered in many of our encounters with AI scientists—arguing against it is taken as a form of Luddism. Larson points out in this connection that there are after all strong incentives for proponents of AI to keep its limitations in the dark, where a healthy culture for innovation ‘emphasises exploring unknowns, not hyping extensions of existing methods—especially when these methods have been shown to be inadequate to take us much further’.

As we shall see in great detail in what follows, human and machine intelligence are radically different. The myth of AI insists that the differences are only temporary, in the sense that, step-by-step, more powerful AI systems will erase them. Yet, as Larson points out, the success achieved by focusing on narrow AI applications such as game-playing or protein folding ‘gets us not one step closer to general intelligence. ... No algorithm exists for general intelligence. And we have good reason to be skeptical that such an algorithm will emerge through further efforts on deep learning systems or any other approach popular today’. To identify one potential alternative approach, Larson points to what he sees as the three different types of inference: *deduction*, which is explored by classic symbolic AI; *induction*, which he classifies as the province of modern stochastic AI³⁰; and a third type which, following the American pragmatist philosopher Peirce, he calls *abduction*. Peirce’s term is nowadays used in different contexts as another word for ‘hypothesis formation’ or also just plain ‘guessing’.³¹

30 This is not correct, as we shall see in chapter 8.6.6.1. Machines do not engage in inductive reasoning; they rather compute local minima for loss functions, which can be seen as a very primitive emulation of induction from data because a functional is indeed obtained from observations (individual data). However, machines do not perform the induction themselves; they merely compute human-designed optimisation algorithms which emulate a narrow form of human induction.

31 For an account of problems we might face in formalising Peirce’s notion, see Frankfurt (1958).

It is abduction, Larson argues, which is at the core of human intelligence, and thus engineering a counterpart of abduction—a combination of intuition and guessing—would be needed for human-level AI. His book provides a thorough and convincing account of why this is so. Yet at the same time he complains that ‘no one is working on it—at all’.

His explanation for this lack of interest is that the myth of AI is holding back AI researchers. Yet this surely underestimates the degree to which the AI field is and has always been unrestrainedly opportunistic. For if modeling abduction truly provided even the beginnings of a feasible path toward modeling human-level intelligence, would there not be contrarian AI researchers who would have started off already down this path?

The fact, if it is a fact, that there is no one who is exploring a strategy along these lines leads us to postulate that this is not for reasons having to do with the culture of AI research. Rather, it is because attempts to engineer the types of abductive inference characteristic of human reasoning have in every case failed to reach even first base. The reasons for this are explored in what follows. For where, already on the first page of his book, Larson asserts that ‘the future of AI is a scientific unknown’, we show that there are in fact many things that we know about the future of AI, all of which derive from the premise that any AI algorithm must be Church-Turing computable.

1.6 How to read this book

If you have not done so already, please go back and read the Foreword.

This provides an account of how the AGI scepticism defended in this, book differs from earlier varieties of scepticism, in that it is based in mathematics, physics, and biology.

For those who want to go straight to the technical details of our argument against the possibility of AGI, read chapters 7 and 8 first. The earlier chapters are there to set the scene, especially as concerns the reasons why human dialogue and human ethics cannot be modelled in a neural network because of the impossibility of collecting representative samples that can be used for training.

For everyone else: read chapter 2 to understand our view of the relationship between the mental and the physical: mental events are a special type of physical event in the brain and are subject to the same laws. We argue that this view is consistent with a common-sense understanding of human mental activity (of how it feels from the inside to be a conscious human being).

Read chapter 3 to understand what the ‘intelligence’ is that AI researchers are seeking to emulate. We introduce a distinction between two types of intelligence: the basic kind, which we share with higher animals; and the type of intelligence that is unique to humans and is closely associated with our ability to use language. We then examine the definition of ‘intelligence’

used by AI researchers and show that this definition does justice to neither of these.

Read chapters 4 to 6 to get an idea of the complex systems formed when human beings interact. These chapters survey our social capabilities as humans, including our capability to use language, to follow social (including ethical) norms, and to engage in social interactions. Human languages and human societies are complex systems—in fact they are complex systems of complex systems.

Read chapters 7 and 8 to understand what complex systems are and why their behaviour cannot be modelled mathematically and therefore cannot be emulated by using computers. We survey attempts to model complex systems in medicine, psychology, and economics. We survey the entire mathematical repertoire of available approaches to the emulation of complex systems, from recursive neural networks through evolutionary process models to entropy models. And we show why they all fail.

Read chapter 9 to find out how the results obtained so far throw light on philosophers' attempts to demonstrate that an AGI, and with it the 'Singularity', can be achieved 'before long'. We focus especially on the attempts by David Chalmers to show how the Singularity might be achieved, either by emulating human intelligence in a machine or by creating a machine intelligence that would emulate the entire course of evolution.

Read chapters 10 and 11 to find out why machines cannot emulate human conversation or moral behaviour. We cover in detail why machines will neither conduct conversations nor interpret text as humans can for a variety of reasons again having to do with the properties of complex systems. We then show why this same complexity rules out the possibility of an AI ethics.

Read chapter 12 if you are interested in 'transhumanism' and in what some are pleased to call 'digital immortality'. Here we address some of the more outlandish speculations that have grown up in the hinterlands of the Singularity. We demonstrate that we can neither create a machine emulation of anything like the human mind nor transcend our human condition as mortal organisms with organic bodies in order to enjoy immortal life in digital form. We also show, along the way, that there will be no AGI, and no Singularity.

Read chapter 13 if you are interested in what can still be achieved by AI in the future, even after taking account of the limits identified in this book. For there are still many grounds for optimism as concerns the potential uses of AI. This chapter is entitled 'AI spring eternal', and it describes how narrow AI will intensify and further broaden the technosphere that mankind has been creating since the beginning of urbanisation and the advent of the first high cultures. Even though there will be no AGI, and no Singularity, AI in the narrow sense will prove itself able to bring about new and still unconceived enhancements and extensions to the texture of our industrialised societies.

References

- Abbott, Barbara. 2017. Reference. In *The Oxford handbook of pragmatics*, edited by Yan Huang , 240–258. London: Oxford University Press.
- Abbott, Ryan. 2019. Everything is obvious. *UCLA Law Review* 66: 2.
- Adams, Eldridge S. 2016. Territoriality in ants (hymenoptera: formicidae): a review. *Myrmecological News* 23: 101–118.
- Aellen, Melisande , Judith M. Burkart , and Redouan Bshary . 2021. No evidence for general intelligence in a fish. *bioRxiv*. <https://doi.org/10.1101/2021.01.08.425841>.
- Alexander, Samuel . 2004. *Space, time, and deity (1920)*. Whitefish, MT: Kessinger Publications.
- Aljalbout, Elie , Vladimir Golkov , 2018. Clustering with deep learning: taxonomy and new methods. *arXiv:1801.07648*.
- Ambady, Nalini , and Robert Rosenthal . 1992. Thin slices of expressive behavior as predictors of interpersonal consequences: a meta-analysis. *Psychological Bulletin* 111 (2): 256–274.
- Antweiler, Christoph. 2019. On the human addiction to norms: social norms and cultural universals of normativity. In *The normative animal? On the anthropological significance of social, moral, and linguistic norms*, edited by Neil Roughley and Kurt Bayertz , 83–100. London: Oxford University Press.
- Armstrong, Rachel. 2013. Alternative biologies. In *The transhumanist reader*, edited by Max More and Natasha Vita-More , 100–109. Oxford: Wiley-Blackwell.
- Armstrong, Rachel , and Neill Spiller . 2010. Synthetic biology: living quarters. *Nature* 467 (7318): 916–918.
- Arora, Saurabh , and Prashant Doshi . 2018. A survey of inverse reinforcement learning: challenges, methods and progress. *arXiv:1806.06877*.
- Arp, Robert , Barry Smith , and Andrew D. Spear . 2015. *Building ontologies with Basic Formal Ontology*. Cambridge, MA: MIT Press.
- Ashburner, Michael , Catherine A. Ball , 2000. Gene Ontology: tool for the unification of biology. *Nature Genetics* 25 (1): 25–29.
- Asimov, Isaac. 1950. *I, robot*. New York: Doubleday.
- Auer, Peter. 2009. On-line syntax: thoughts on the temporality of spoken language. *Language Sciences* 31 (1): 1–13.
- Austin, John L. 1962. *How to do things with words*. Oxford: Clarendon Press.
- Babaie, Hassan , Armita Davarpanah , and Nirajan Dhakal . 2019. Projecting pathways to food-energy-water systems sustainability through ontology. *Environmental Engineering Science* 36 (7): 808–819.
- Baberowski, Jörg. 2015. *Räume der Gewalt*. Frankfurt a. M.: S. Fischer Verlag.
- Band, Yehuda B. , and Yshai Avishai . 2012. *Quantum mechanics with applications to nanotechnology and information science*. Cambridge, MA: Academic Press.
- Barker, Roger G. 1968. *Ecological psychology*. Stanford: Stanford University Press.
- Barker, Roger G. , and Herbert F. Wright . 1951. *One boy's day: a specimen record of behavior*. New York: Harper & Brothers.
- Bassenge, Friedrich. 1930. Hexis und akt. Eine phänomenologische skizze. *Philosophischer Anzeiger* 4: 163–168.
- Bekoff, Marc , Colin Allen , Gordon M. Burghardt , 2002. *The cognitive animal: empirical and theoretical perspectives on animal cognition*. Cambridge, MA: MIT Press.
- Ben-David, Shai , Pavel Hrubeš , 2019. Learnability can be undecidable. *Nature Machine Intelligence* 1 (1): 44.
- Berger, Peter , and Thomas Luckmann . 1966. *The social construction of reality*. New York: Anchor Books.
- Bergson, Henri. 1911. *Creative evolution [1907]*. New York: Henry Holt & Co.
- Bernstein, Ethan , and Umesh Vazirani . 1997. Quantum complexity theory. *SIAM Journal on Computing* 26 (5): 1411–1473.
- Bertsekas, Dimitri P. 2016. *Nonlinear programming*. Belmont, MA: Athena Scientific.

Bhattacharya, Sanchita , Patrick Dunn , 2018. Immport, toward repurposing of open access immunological assay data for translational and clinical research. *Scientific Data* 5: 180015.

Bickle, John. 2020. Multiple realizability. In *The Stanford encyclopedia of philosophy*, Winter 2020, <https://plato.stanford.edu/entries/multiple-realizability/>.

Bitterman, M. E. 2006. Classical conditioning since Pavlov. *Review of General Psychology* 10 (4): 365–376.

Bittner, Thomas , and Barry Smith . 2001. A taxonomy of granular partitions. In *Spatial information theory. Foundations of geographic information science*, edited by Daniel Montello , 28–43. Berlin: Springer.

Bittner, Thomas , and Barry Smith . 2003. A theory of granular partitions. In *Foundations of geographic information science*, edited by M. Duckham , M. F. Goodchild , and M. F. Worboys , 117–151. London: Taylor & Francis.

Block, Ned. 1978. Troubles with functionalism. In *Perception and cognition*, edited by W. Savage , 261–325. Minneapolis: University of Minnesota Press.

Block, Ned. . 1995. On a confusion about a function of consciousness. *Behavioral and Brain Sciences*, 18 (2): 227–247.

Boden, Margaret. 1977. *Artificial intelligence and natural man*. New York: Branch Line.

Bommasani, Rishi , Drew A. Hudson , 2021. On the opportunities and risks of foundation models. [arXiv:abs/2108.07258](https://arxiv.org/abs/2108.07258) [cs.LG].

Boolos, G. S. , J. P. Burgess , and R. C. Jeffrey . 2007. *Computability and logic*. Cambridge: Cambridge University Press.

Borghini, Andrea , and Neil E. Williams . 2008. A dispositional theory of possibility. *Dialectica* 61: 21–41.

Bostrom, Nick. 2003. *Superintelligence: paths, dangers, strategies*. London: Oxford University Press.

Bostrom, Nick. . 2013. Why I want to be a posthuman when I grow up. In *The transhumanist reader*, edited by Max More and Natasha Vita-More , 28–53. Oxford: Wiley-Blackwell.

Boyd, Robert , and Peter J. Richerson . 1996. Why culture is common, but cultural evolution is rare. *Proceedings of the British Academy: Evolution of Social Behaviour Patterns in Primates and Man* 88: 77–93.

Boyle, Evan A. , Yand I. Li , and Jonathan K. Pritchard . 2017. An expanded view of complex traits: from polygenic to omnigenic. *Cell* 169: 1177–1186.

Bringer, Eran , Abraham Israeli , 2019. Osprey: weak supervision of imbalanced extraction problems without code. In *Proceedings of the 3rd international workshop on data management for end-to-end machine learning*, 1–11. <https://doi.org/10.1145/3329486.3329492>.

Bringsjord, Selmer. 2015. A refutation of Searle on Bostrom (re: malicious machines) and Floridi (re: information). *APA Newsletter on Philosophy and Computation* 15 (1): 7–9.

Brogaard, Berit. 2017. The publicity of meaning and the perceptual approach to speech comprehension. *ProtoSociology* 34: 144–162.

Brogaard, Berit. 2019. Seeing and hearing meanings: a non-inferential approach to speech comprehension. In *Inference and consciousness*, edited by Timothy Chan and Anders Nes , 99–124. London: Routledge.

Brooks, Rodney A. 1991. Intelligence without representation. *Artificial Intelligence* 47 (1–3): 139–159.

Brown, Noam , and Tuomas Sandholm . 2019. Superhuman AI for multiplayer poker. *Science* 365 (6456): 885–890.

Brown, Tom B. , Benjamin Pickman Mann , 2020. Language models are few-shot learners. [arXiv abs/2005.14165](https://arxiv.org/abs/2005.14165).

Brundage, Miles. 2014. Limitations and risks of machine ethics. *Journal of Experimental & Theoretical Artificial Intelligence* 26 (3): 355–372.

Bühler, Karl. 1927. *Die Krise der Psychologie*. Jena: Gustav Fischer.

Bühler, Karl. . 1990. *Theory of language: the representational function of language*. Amsterdam: John Benjamins Publishing Company.

Bulgakov, Mikhail. 1967. *The master and Margarita*. New York: Grove Press.

Cabessa, Jérémie , and Hava Siegelmann . 2014. The super-Turing computational power of plastic recurrent neural networks. *International Journal of Neural Systems* 24: 1450029.

Cantwell Smith, Brian. 2019. *The promise of artificial intelligence: reckoning and judgment*. Cambridge, MA: MIT Press.

Carey, Susan , and Fei Xu . 2001. Infants' knowledge of objects: beyond object files and object tracking. *Cognition* 80: 179–213.

Cavalli-Sforza, Luigi Luca. 2000. *Genes, peoples, and languages*. New York: Farrar, Straus & Giroux.

Chaitanyaa, Lakshmi , Krystal Breslinb , 2018. The HlrisPlex-S system for eye, hair and skin colour prediction from DNA: introduction and forensic developmental validation. *Forensic Science International: Genetics* 35: 125–135.

Chalmers, David J. 1996. *The conscious mind: in search of a fundamental theory*. Oxford: Oxford University Press.

Chalmers, David J. . 2010. The singularity: a philosophical analysis. *Journal of Consciousness Studies* 17: 7–65.

Chalmers, David J. 2012. The singularity: a reply to commentators. *Journal of Consciousness Studies*: 141–167.

Chalmers, David J. 2016. The singularity: a philosophical analysis. In *The singularity: could artificial intelligence really out-think us (and would we want it to)?*, edited by Uziel Awret , 12–88. Exeter: Imprint Academic.

Chan, Ronald Ping Man , Karl A. Stol , and C. Roger Halkyard . 2013. Review of modelling and control of two-wheeled robots. *Annual Reviews in Control* 37 (1): 89–103.

Chapman, Jeremy R. , David Kasmier , 2021. Conceptual Spaces For Space Event Characterization Via Hard And Soft Data Fusion. In *American Institute of Aeronautics and Astronautics (AIAA) Scitech Forum*, 11–15. <https://doi.org/10.2514/6.2021-1163>.

Chatterjee, Deen K. 2003. Moral distance: introduction. *The Monist*: 327–332.

Chen, Ritchie , Felicity Gore , 2020. Deep brain optogenetics without intracranial surgery. *Nature Biotechnology*: 1–4.

Chollet, François. 2017. *Deep learning with Python*. Shelter Island, NY: Manning Publications Company.

Chu, Dominique , Roger Strand , and Ragnar Fjelland . 2003. Theories of complexity. *Complexity* 8 (3): 19–30.

Church, Alonzo. 1936. A note on the Entscheidungsproblem. *Journal of Symbolic Logic* 1: 40–41.

Churchland, Paul M. 1981. Eliminative materialism and the propositional attitudes. *The Journal of Philosophy* 78: 67–90.

Clark, Andy. 2013. Re-inventing ourselves: the plasticity of embodiment. In *The transhumanist reader*, edited by Max More and Natasha Vita-More , 111–127. Oxford: Wiley-Blackwell.

Clarke, Edmund M. , Orna Grumberg , 1999. State space reduction using partial order techniques. *International Journal on Software Tools for Technology Transfer* 2 (3): 279–287.

Cohen, Michael K. , Badri N. Vellambi , and Marcus Hutter . 2019. Asymptotically unambitious artificial general intelligence. *CoRR abs/1905.12186*.

Coleman, Jonathan R. I. , Julien Bryois , 2019. Biological annotation of genetic loci associated with intelligence in a meta-analysis of 87,740 individuals. *Molecular Psychiatry* 24 (2): 182–197.

Conerton, Paul. 1979. *How societies remember*. New York: Cambridge University Press.

Conrad, Rolf. 1999. Contribution of hydrogen to methane production and control of hydrogen concentrations in methanogenic soils and sediments. *FEMS Microbiology Ecology* 28: 193–202.

Cosmides, Leda , and John Tooby . 2005. Neurocognitive adaptations designed for social exchange. In *Handbook of evolutionary psychology*, edited by D. M. Buss , 584–627. Hoboken, NJ: Wiley.

Cowell, Robert G. , A. P. Dawid , 2007. *Probabilistic networks and expert systems*. Berlin: Springer.

Crandall, Richard E. 1996. Topics in advanced scientific computation. Chapter on Nonlinear and Complex Systems. Berlin: Springer.

Crosby, Matthew , Benjamin Beyret , and Marta Halina . 2019. The animal-AI Olympics. *Nature Machine Intelligence* 1 (257).

Damasio, Antonio R. 1999. The feeling of what happens: body and emotion in the making of consciousness. Boston, MA: Houghton Mifflin Harcourt.

D'Amour, Alexander , Katherine Heller , 2020. Underspecification presents challenges for credibility in modern machine learning. arXiv, eprint: 1906.01563.

Darwiche, Adnan . 2018. Human-level intelligence or animal-like abilities? *Communications of the ACM* 61 (10): 56–67.

Darwin, Charles. 1872. The expression of emotions in man and animals. London: John Murray.

Dasgupta, Sakyasingha , and Takayuki Osogami . 2017. Nonlinear dynamic Boltzmann machines for time-series prediction. In *Proceedings of the 31st AAAI Conference on Artificial Intelligence*, 1833–1839. <https://ojs.aaai.org/index.php/AAAI/article/view/10806>.

Davidson, Donald. 1970. Mental events. In *Experience and theory*, edited by L. Foster and J. W. Swanson , 207–224. Oxford: Clarendon Press.

Davidson, Donald. 1987. Knowing one's own mind. *Proceedings and Addresses of the American Philosophical Association* 60: 441–458.

Davidson, Donald. 2009. Truth and predication. Cambridge, MA: Harvard University Press.

Davis, Martin. 2004. The myth of hypercomputation. In *Alan Turing: life and legacy of a great thinker*, edited by Christof Teuscher , 195–211. Heidelberg: Springer.

Davis, Martin , Yuri Matijasevic , and Julia Robinson . 1976. Hilbert's tenth problem. Diophantine equations: positive aspects of a negative solution. *Proceedings of Symposia in Pure Mathematics* 28: 323–378.

Davis, Zachary. 2017. Max Scheler and pragmatism. In *Pragmatic perspectives in phenomenology*, edited by O. Švec and Jakub Čapek , 158–172. New York: Routledge.

Deamer, David. 2005. A giant step towards artificial life? *Trends in Biotechnology* 23 (7): 336–338.

De Caro, Mario. 2015. Realism, common sense, and science. *The Monist* 98 (2): 197–214.

Degenaar, Jan , and J. Kevin O'Regan . 2015. Sensorimotor theory of consciousness. *Scholarpedia* 10: 4952.

Dennett, Daniel C. 2018. Facing up to the hard question of consciousness. *Philosophical Transactions of the Royal Society B: Biological Sciences* 373.

De Soto, Hernando. 2000. The mystery of capital: why capitalism triumphs in the West and fails everywhere else. New York: Civitas Books.

Dignum, Virginia. 2018. Ethics in artificial intelligence: introduction to the special issue. *Ethics and Information Technology* 20 (1): 1–3.

Drace, Sasa. 2013. Evidence for the role of affect in mood congruent recall of autobiographic memories. *Motivation and Emotion* 37: 623–628.

Dreyfus, Hubert L. 1992. What computers still can't do: a critique of artificial reason. Cambridge, MA: MIT Press.

Dreyfus, Hubert L. , Stuart E. Dreyfus , and Tom Athanasiou . 2000. Mind over machine. New York: Simon/Schuster.

Dubhashi, Devdatt , and Shalom Lappin . 2017. AI dangers: imagined and real. *Communications of the ACM* 60 (2): 43–45.

Eden, Amnon H. , and James H. Moor , editors. 2012. Singularity hypotheses: a scientific and philosophical assessment. Dordrecht: Springer.

Efron, Bradley. 1979. Bootstrap methods: another look at the jackknife. *The Annals of Statistics* 7 (1): 1–26.

Einstein, Albert. 1905. Über die von der molekularkinetischen Theorie der Wärme geforderte Bewegung von in ruhenden Flüssigkeiten suspendierten Teilchen. *Annalen der Physik* 322 (8): 549–560.

Eisenstein, Elizabeth L. 1980. The printing press as an agent of change. Cambridge: Cambridge University Press.

Ekbia, H. R. 2008. Artificial dreams. Cambridge: Cambridge University Press.

Elman, Jeffrey L. , Elizabeth A. Bates , 1996. Rethinking innateness: a connectionist perspective on development. Cambridge, MA: MIT Press.

Enderton, Herbert B. 2010. Computability theory: an introduction to recursion theory. Cambridge, MA: Academic Press.

Eysenck, Michael W. , and Christine Eysenck . 2021. AI vs Humans. Abingdon: Routledge.

Falkovich, Gregory. 2011. Fluid mechanics. Cambridge: Cambridge University Press.

Falkovich, Gregory , and Katepalli R. Sreenivasan . 2006. Lessons from hydrodynamic turbulence. *Physics Today* 59 (4): 43.

Fernández-Villaverde, Jesús. 2020. Simple rules for a complex world with artificial intelligence. <https://economics.sas.upenn.edu/pier/working-paper/2020/simple-rules-complex-world-artificial-intelligence>.

Fetzer, Anita. 2017. Context. In *The Oxford handbook of pragmatics*, edited by Yan Huang , 259–276. London: Oxford University Press.

Feynman, Richard P. , Robert B. Leighton , and Matthew Sands . 2010. *The Feynman lectures on physics (1964)*. Boston, MA: Addison-Wesley.

Fieguth, Paul. 2017. Complex systems. In *An introduction to complex systems*, 245–269. Berlin: Springer.

Fjelland, Ragnar. 2020. Why general artificial intelligence will not be realized. *Humanities and Social Sciences Communications* 7 (1): 1–9.

Fodor, Jerry A. 1981. The mind-body problem. *Scientific American* 244 (1): 114–123.

Fodor, Jerry A. 2005. Reply to Steven Pinker 'So how does the mind work?' *Mind and Language* 20 (1): 25–32.

Fodor, Jerry A. , and Zenon W. Pylyshyn . 1988. Connectionism and cognitive architecture: a critical analysis. *Cognition* 28 (1–2): 3–71.

Ford, Martin. 2018. *Architects of intelligence: the truth about AI from the people building it*. Birmingham: Packt Publishing Ltd.

Forguson, Lynd. 1989. *Common sense*. London: Routledge.

Forguson, Lynd , and Alison Gopnik . 1988. The ontogeny of common sense. In *Developing theories of mind*, 226–243. London and New York: Cambridge University Press.

Frankfurt, Harry G. 1958. Peirce's notion of abduction. *The Journal of Philosophy* 55 (14): 593–597.

Frankfurt, Harry G. . 1971. Freedom of the will and the concept of a person. *Journal of Philosophy* 68 (1): 5–20.

Freitas, Robert A. 2009. Welcome to the future of medicine. In *Studies in health technology and informatics*, edited by Renata G. Bushko , 149:251–256. Amsterdam, NL: IOS Press.

Friston, Karl. 2018. Does predictive coding have a future? *Nature Neuroscience* 21 (8): 1019–1021.

Fuxjager, Matthew J. , and Barney A. Schlinger . 2015. Perspectives on the evolution of animal dancing: a case study of manakins. *Current Opinion in Behavioral Sciences* 6: 7–12.

Gabriel, Iason. 2020. Artificial intelligence, values, and alignment. *Minds and Machines* 30 (3): 411–437.

Gamut, L. T. F. 1991a. *Logic, language and meaning*. Vol. 1. Chicago and London: The University of Chicago Press.

Gamut, L. T. F. 1991b. *Logic, language and meaning*. Vol. 2. Chicago and London: The University of Chicago Press.

Gando, A. , Y. Gando , 2011. Partial radiogenic heat model for Earth revealed by geoneutrino measurements. *Nature Geoscience* 4 (9): 647.

Gao, Jianfeng , Michel Galley , and Lihong Li . 2018. Neural approaches to conversational AI. *arXiv abs/1809.08267*.

Garson, James. 2018. Connectionism. In *The Stanford encyclopedia of philosophy*, Spring 2018, <https://plato.stanford.edu/entries/connectionism/>.

Gavaldà, Ricard , and Hava T. Siegelmann . 1999. Discontinuities in recurrent neural networks. *Neural Computation* 11 (3): 715–745.

Gavrilets, Sergey , and Peter J. Richerson . 2017. Collective action and the evolution of social norm internalization. *Proceedings of the National Academy of Sciences* 114 (23):

6068–6073.

- Gehlen, Arnold. 1988. *Man: his nature and place in the world* [1940]. New York: Columbia University Press.
- Gehlen, Arnold. 1993 (1940). *Der Mensch. Seine Natur und seine Stellung in der Welt*. Frankfurt am Main: Vittorio Klostermann.
- Gelman, Susan A. , and Henry M. Wellman . 1991. Insides and essences: early understandings of the non-obvious. *Cognition* 38 (3): 213–244.
- Ghazvininejad, Marjan , Chris Brockett , 2017. A knowledge-grounded neural conversation model. *arXiv abs/1702.01932*.
- Gibson, James J. 1963. The useful dimensions of sensitivity. *American Psychologist* 18 (1): 1–15.
- Gibson, James J. 1966. *The senses considered as perceptual systems*. Boston, MA: Houghton Mifflin.
- Gibson, James J. 2015. *An ecological theory of perception* (1979). Boston, MA: Houghton Mifflin.
- Gilman, Sander L. , Carole Blair , and David J. Parent . 1990. *Friedrich Nietzsche on rhetoric and language*. Oxford: Oxford University Press.
- Goebel, Randy , Ajay Chander , Katharina Holzinger , 2018. Explainable AI: the New 42? In *Machine learning and knowledge extraction. Lecture notes in computer science*. 11015, edited by A. Holzinger , P. Kieseberg , 295–303. Berlin: Springer.
- Goertzel, Ben , and Cassio Pennachin , editors. 2007a. *Artificial general intelligence*. Berlin and Heidelberg: Springer-Verlag.
- Goertzel, Ben , and Cassio Pennachin . 2007b. The Novamente artificial intelligence engine. In *Artificial general intelligence*, edited by Ben Goertzel and Cassio Pennachin , 76–129. Berlin and Heidelberg: Springer-Verlag.
- Gómez-Vilda, Pedro , A. R. M. Londral , 2013. Characterization of speech from amyotrophic lateral sclerosis by neuromorphic processing. In *Natural and artificial models in computation and biology (IWINAC 2013)*, edited by J. M. Ferrández-Vicente , J. R. Álvarez Sánchez , Berlin/Heidelberg: Springer.
- Goodfellow, Ian J. , Yoshua Bengio , 2014. Generative adversarial networks. *arXivabs/1406.2661*.
- Goodfellow, Ian J. , Yoshua Bengio , and Aaron Courville . 2016. *Deep learning*. Cambridge, MA: MIT Press.
- Gopnik, Alison , and Andrew N. Meltzoff . 1997. *Words, thoughts, and theories*. Cambridge, MA: MIT Press.
- Goriounova, Natalia A. , and Huibert D. Mansvelder . 2019. Genes, cells and brain areas of intelligence. *Frontiers in Human Neuroscience* 13: 44.
- Görlach, Manfred. 1996. And is it english? *English World-Wide* 17 (2): 153–174.
- Gottfredson, Linda S. 1997. Mainstream science on intelligence. *Intelligence* 24: 13–23.
- Graves, Alex , Greg Wayne , and Ivo Danihelka . 2014. Neural Turing machines. *arXiv*, eprint: 1410.5401.
- Green, Patrick A. , Nicholas C. Brandley , and Stephen Nowicki . 2020. Categorical perception in animal communication and decision-making. *Behavioral Ecology* 31 (4): 859–867.
- Grice, H. Paul . 1957. Meaning. *The Philosophical Review* 66: 377–388.
- Grice, H. Paul . 1989. Logic and conversation. In *Studies in the way of words*, 22–40. Cambridge, MA: Harvard University Press.
- Gunter, Pete A. Y. 1991. Bergson and non-linear non-equilibrium thermodynamics: an application of method. *Revue Internationale de Philosophie* 45 (177): 108–121.
- Haack, Susan. 2011. *Defending science—within reason: between scientism and cynicism*. Prometheus Books.
- Hagendorff, Thilo. 2020. The ethics of AI ethics: an evaluation of guidelines. *Minds and Machines* 30 (1): 99–120.
- Hampshire, Stuart. 1959. *Thought and action*. London: University of Notre Dame Press.
- Hanks, William. 1996. Language form and communicative practices. In *Rethinking linguistic relativity*, edited by J. J. Gumpertz and S. C. Levinson . Cambridge, MA: Cambridge

University Press.

Harré, Rom , and Edward H. Madden . 1975. *Causal powers: theory of natural necessity*. Oxford: Blackwell Publishers.

Hartmann, Nicolai. 2014. *Ethics* (3 volumes) [1949]. London: Routledge.

Hastie, Trevor , Robert Tibshirani , and Jerome Friedman . 2008. *The elements of statistical learning*. 2nd ed. Berlin: Springer.

Hastings, J. , W. Ceusters , 2011. Dispositions and processes in the Emotion Ontology. In *Proceedings of the 2nd international conference on biomedical ontology*, 71–78, <http://ceur-ws.org/Vol-833/paper10.pdf>.

Haugeland, John. 1985. *Artificial intelligence: the very idea*. Cambridge, MA: MIT Press.

Haugeland, John. 1993. Mind embodied and embedded. In *Mind and cognition: 1993 International Symposium*, edited by Yu-Houng H. Houng and J. Ho , 233–267. Taipei: Academica Sinica.

Hayek, Friedrich August von. 1937. Economics and knowledge. *Economica* 4 (13): 33–54.

Hayek, Friedrich August von. 1945. The use of knowledge in society. *The American Economic Review* 35 (4): 519–530.

Hayek, Friedrich August von. 1952. *The sensory order. An inquiry into the foundations of theoretical psychology*. Chicao, IL: Chicago University Press.

Hayek, Friedrich August von. 1967. Notes on the evolution of systems of rules of conduct. *Studies in Philosophy, Politics and Economics*: 66–81.

Hayek, Friedrich August von. 1996. *Studies in philosophy, politics and economics*. New York: Touchstone.

Hayek, Friedrich August von. 2014a. *The market and other orders*. Edited by Bruce Caldwell . Chicago: University of Chicago Press.

Hayek, Friedrich August von. 2014b. The pretence of knowledge. In *The market and other orders*, edited by Bruce Caldwell , 362–372. Chicago: University of Chicago Press.

Hayes, P. J. 1985. The second naive physics manifesto. In *Formal theories of the commonsense world*, edited by J. R. Hobbs and R. C. Moore . Norwood, NJ: Ablex Publishing.

Heft, Harry. 2001. *Ecological psychology in context: James Gibson, Roger Barker, and the legacy of William James's radical empiricism*. Mahwah, NJ: Lawrence Erlbaum.

Heft, Harry. 2017. Perceptual information of an entirely different order: the cultural environment in The senses considered as perceptual systems. *Ecological Psychology* 29: 122–145.

Heineman, George T. , Gary Pollice , and Stanley Selkow . 2008. *Algorithms in a nutshell*. Sebastopol, CA: O'Reilly.

Hempel, Carl. 1969. Reduction: ontological and linguistic facets. In *Philosophy, science, and method: essays in honor of Ernest Nagel*, edited by M. White , S. Morgenbesser , and P. Suppes . New York: St Martin's Press.

Hertz, Heinrich. 1899. *The principles of mechanics presented in a new form*. London: Macmillan/Company.

Herzig, Andreas , Laurent Perrussel , 2016. Refinement of intentions. In *Logics in artificial intelligence*, edited by Loizos Michael and Antonis Kakas , 558–563. Cham: Springer International Publishing. ISBN: 978-3-319-48758-8.

Hesse, Mary. 1963. *Models and analogies in science*. London: Sheed/Ward.

Hesse, Mary. 1980. *Revolutions and reconstructions in the philosophy of science*. Brighton, Sussex: The Harvester Press.

Hibbard, Bill. 2012. Model-based utility functions. *Journal of Artificial General Intelligence* 3 (1): 1–24.

Hirschberg, Julia , and Christopher D . Manning. 2015. Advances in natural language processing. *Science* 349: 261–266.

Hobaiter, Catherine , and Richard W. Byrne . 2014. The meanings of chimpanzee gestures. *Current Biology* 24: 1596–1600.

Hochreiter, Sepp , and Jürgen Schmidhuber . 1997. Long short-term memory. *Neural Computation* 9 (8): 1735–1780.

Hohwy, Jakob. 2013. *The predictive mind*. Oxford: Oxford University Press.

Holton, Robert , and Bryan S. Turner . 2010. *Max Weber on economy and society*. London and New York: Routledge.

Horgan, John G. , Max Taylor , 2017. From cubs to lions: a six stage model of child socialization into the Islamic State. *Studies in Conflict & Terrorism* 40 (7): 645–664.

Hornik, Kurt. 1991. Approximation capabilities of multilayer feedforward networks. *Neural Networks* 4 (2): 251–257.

Horton, Robin. 1982. Tradition and modernity revisited. In *Rationality and relativism*, edited by Martin Hollis and Steven Lukes , 201–260. Cambridge, MA: MIT Press.

Hu, Zicheng , Alice Tang , 2020. A robust and interpretable, end-to-end deep learning model for cytometry data. *Proceedings of the National Academy of Sciences* 117 (35): 21373–21380.

Huang, Yan. 2017. Implicature. In *The Oxford handbook of pragmatics*, edited by Yan Huang , 155–179. London: Oxford University Press.

Humphries, Nicolas E. , Daniel W. Fuller , 2010. Environmental context explains Lévy and Brownian movement patterns of marine predators. *Nature* 465 (7301): 1066–1069.

Husserl, Edmund. 1989. *The crisis of European sciences and transcendental phenomenology: an introduction to phenomenological philosophy [1936]*. Evanston, IL: Northwestern University Press.

Husserl, Edmund. 2000. *Logical Investigations [1901]*. Abingdon: Routledge.

Hutchinson, G Evelyn. 1957. *A treatise on limnology. vol 1: geography, physics and chemistry*. New York: John Wiley & Sons.

Hutchison, Clyde A. , Ray-Yuan Chuang , 2016. Design and synthesis of a minimal bacterial genome. *Science* 351 (6280).

Hutter, Marcus. 2012. Can intelligence explode? *Journal of Consciousness Studies* 19 (1–2): 143–166.

Hwangbo, Jemin , Joonho Lee , 2019. Learning agile and dynamic motor skills for legged robots. *Science Robotics* 4 (26).

Ingarden, Roman. 1973. *The literary work of art. Investigations on the borderlines of ontology, logic and the theory of literature*. Evanston, IL: Northwestern University Press.

Ingarden, Roman. 1983. *Man and value*. Munich: Philosophia.

Ingarden, Roman. 1986. *The work of music and the problem of its identity*. Berkeley, CA: University of California Press.

Ingarden, Roman. 2013. *Controversy over the existence of the world. Vol. I*. Frankfurt a. M.: Peter Lang Edition.

Ingarden, Roman. 2019. *Controversy over the existence of the world. Vol. II*. Frankfurt a. M.: Peter Lang Edition.

Jackson, Pete. 1998. *Introduction to expert systems*. Boston, MA: Addison Wesley.

Jacob, Pierre. 2008. What do mirror neurons contribute to human social cognition? *Mind and Language* 23: 190–223.

Jaderberg, Max , Wojciech M. Czarnecki , and Iain Dunning . 2019. Human-level performance in first-person multiplayer games with population-based deep reinforcement learning. *Science* 364: 859–865.

Jeamblanc, Monique , Marc Yor , and Marc Chesney . 2009. *Mathematical methods for financial markets*. Berlin and New York: Springer.

Jelbert, Sarah A. , Alex H. Taylor , 2014. Using the Aesop's fable paradigm to investigate causal understanding of water displacement by New Caledonian crows. *PLoS ONE* 9 (3): e92895.

Jensen, Arthur R. 1998. *The g factor: the science of mental ability*. Westport, CT: Greenwood.

Jo, Jason , and Yoshua Bengio . 2017. Measuring the tendency of CNNs to learn surface statistical regularities. *arXiv abs/1711.11561*.

Johansson, Ingvar. 1998. Pattern as an ontological category. In *Formal ontology in information systems*, edited by N. Guarino , 86–94. Amsterdam, NL: IOS Press.

Johansson, Ingvar. 2012. *Ontological investigations*. Frankfurt am Main: de Gruyter.

Jumper, John , Richard Evans , 2020. High accuracy protein structure prediction using deep learning. In *Critical assessment of techniques for protein structure prediction (CASP-14)*, edited by John Moult , 22. <https://predictioncenter.org/casp14/>.

Kandel, Eric , John D. Koester , 2021. *Principles of neural science*. New York: McGraw Hill.

Kant, Immanuel . 1991. *The metaphysics of morals (part 1): the philosophy of law—an exposition of the fundamental principles of jurisprudence as the science of right (1797)*. Cambridge: Cambridge University Press.

Kant, Immanuel . 1998. *Groundwork of the metaphysics of morals [1795]*. Cambridge: Cambridge University Press.

Kant, Immanuel . 2000. *Critique of the power of judgment*. The Cambridge Edition of the Works of Immanuel Kant. Cambridge: Cambridge University Press.

Karras, Tero , Samuli Laine , and Timo Aila . 2018. A style-based generator architecture for Generative Adversarial Networks. *arXiv abs/1812.04948*.

Karray, Mohamed , Neil Otte , 2021. The Industrial Ontologies Foundry (IOF) perspectives. International conference on interoperability for enterprise systems and applications. Tarbes, France: IOF—Achieving Data Interoperability Workshop.

Katzen, Abraham , Hui-Kuan Chung , 2021. The nematode worm *C. elegans* chooses between bacterial foods exactly as if maximizing economic utility. *bioRxiv*, 2021.04.25.441352.

Keil, Frank C. 1989. *Concepts, kinds and cognitive development*. Cambridge, MA: MIT Press.

Keil, Frank C. 1994. Explanation, association, and the acquisition of word meaning. *Lingua* 92 (1–4): 169–196.

Kempes, Christopher P. , Geoffrey B. West , and Mimi Koehl . 2019. The scales that limit: the physical boundaries of evolution. *Frontiers in Ecology and Evolution* 7: 242.

Kim, In-Kyeong , and Elizabeth S. Spelke . 1999. Perception and understanding of effects of gravity and inertia on object motion. *Developmental Science* 2 (3): 339–362.

Kim, Jaegwon. 1984. Concepts of supervenience. *Philosophy and Phenomenological Research* 45 (2): 153–176.

King, Julia , Michele Insanally , 2015. Rodent auditory perception: critical band limitations and plasticity. *Neuroscience* 296: 55–65.

Klenke, Achim . 2013. *Probability theory: a comprehensive course*. 2nd ed. New York and Berlin: Springer.

Kogo, Naoki , and Raymond van Ee . 2014. Neural mechanisms of figure-ground organization: border-ownership, competition and perceptual switching, 352–372. *Oxford: Oxford Handbook of Perceptual Organization*.

Kolmogorov, A. N. 1941a. Dissipation of energy in the locally isotropic turbulence. *Doklady Akademii Nauk SSSR* 31: 538–540.

Kolmogorov, A. N. . 1941b. The local structure of turbulence in incompressible viscous fluid for very large Reynolds numbers. *Doklady Akademii Nauk SSSR* 30: 301–305.

Kross, E. , A. Duckworth , 2011. The effect of self-distancing on adaptive versus mal-adaptive self-reflection in children. *Emotion* 11: 1032–1039.

Kurzweil, Ray. 2005. *The singularity is near*. New York: Viking Press.

Kwiatkowski, Tom , Jennimaria Palomaki , 2019. Natural questions: a benchmark for question answering research. *Transactions of the Association for Computational Linguistics* 7: 453–466.

Ladyman, James , James Lambert , and Karoline Wiesner . 2013. What is a complex system? *European Journal for Philosophy of Science* 3 (1): 33–67.

Ladyman, James , Stuart Presnell , 2007. The connection between logical and thermodynamic irreversibility. *Studies in History and Philosophy of Science* 38 (1): 58–79.

Lai, Guokun , Wei-Cheng Chang , 2017. Modeling long- and short-term temporal patterns with Deep Neural Networks. *arXiv abs/1703.07015*.

La Mettrie, Julien Offray de. 1748. *L'homme machine*. Leiden: Elie Luzac.

Landauer, Rolf. 1961. Irreversibility and heat generation in the computing process. *IBM Journal of Research and Development* 5 (3): 183–191.

Landgrebe, Jobst. 2022. Certifiable AI. *Applied Science* 12 (3): 1050. <https://doi.org/10.3390/app12031050>.

Landgrebe, Jobst , and Barry Smith . 2021. Making AI meaningful again. *Synthese* 198: 2061–2081. <https://doi.org/10.1007/s11229-019-02192-y>.

La Porta, A. , G. A. Voth , 2001. Fluid particle accelerations in fully developed turbulence. *Nature* 409: 1017–1019.

Lapuschkin, Sebastian , Stephan Wäldchen , 2019. Unmasking clever Hans predictors and assessing what machines really learn. *Nature Communications* 10 (1): 1–8.

Larsen, Rasmus Rosenberg. 2020. Psychopathy as moral blindness: a qualifying exploration of the blindness-analogy in psychopathy theory and research. *Philosophical Explorations* 23 (3): 214–233.

Larson, Erik J. 2021. *The myth of artificial intelligence: why computers can't think the way we do*. Cambridge, MA: Harvard University Press.

Latrémouille, Christian , Alain Carpentier , 2018. A bioprosthetic total artificial heart for end-stage heart failure: results from a pilot study. *The Journal of Heart and Lung Transplantation* 37 (1): 33–37.

Lau, Ellen F. , Colin Phillips , and David Poeppel . 2008. A cortical network for semantics: (de)constructing the n400. *Nature Reviews Neuroscience* 9: 920–933.

LeCun, Yann. 2015. Facebook AI director Yann Lecun on his quest to unleash deep learning and make machines smarter. <https://spectrum.ieee.org/automaton/artificial-intelligence/machinelearning/facebook-ai-director-yann-lecun-on-deep-learning..>

Leff, Harvey , and Andrew F. Rex . 2002. *Maxwell's demon 2—entropy, classical and quantum information, computing*. London: IOP Publishing.

Legg, Shane , and Marcus Hutter . 2007. Universal intelligence: a definition of machine intelligence. *Minds and Machines* 17: 391–444.

Leike, Jan , David Krueger , 2018. Scalable agent alignment via reward modeling: a research direction. *arXiv abs/1811.07871*.

Lepore, Ernie , and Matthew Stone . 2018. Pejorative tone. In *Bad words: philosophical perspectives on slurs*, edited by David Sosa , 134–153. Oxford: Oxford University Press.

Levesque, Henri J. 2014. On our best behaviour. *Artificial Intelligence* 213: 27–35.

Levinson, Stephen C. 1983. Deixis in pragmatics. In *Pragmatics*, 54–96. Cambridge, MA: Cambridge University Press.

Lewis, David. 1979. Scorekeeping in a language game. In *Semantics from different points of view*, 172–187. Berlin: Springer.

Li, Bian , Michaela Fooksa , 2018. Finding the needle in the haystack: towards solving the protein-folding problem computationally. *Critical Reviews in Biochemistry and Molecular Biology* 53: 1–28.

Li, Jiwei , Will Monroe , 2016. Deep reinforcement learning for dialogue generation. *CoRR abs/1606.01541*.

Li, Michael. 2018. *An introduction to mathematical modeling of infectious diseases*. Berlin: Springer.

Li, Zongyi , Nikola Kovachki , 2020. Fourier neural operator for parametric partial differential equations. *arXiv, eprint: 2010.08895*.

Lin, Stephanie , Jacob Hilton , and Owain Evans . 2021. TruthfulQA: measuring how models mimic human falsehoods. *arXiv: 2109.07958*.

Loebner, Sebastian. 2013. *Understanding semantics*. New York: Routledge.

Lorini, Giuseppe. 2018. *Animal norms: an investigation of normativity in the non-human social world*. Law, Culture and the Humanities. Los Angeles: SAGE Publications, <https://doi.org/10.1177/1743872118800008>.

Lotze, Hermann. 1841. *Metaphysik*. Leipzig: Weidmann'sche Buchhandlung.

Lovelace, Ada , and Luigi Menabrea . 1843. Sketch of the Analytical Engine invented by Charles Babbage esq. In *Scientific memoirs*. Selected from the Transactions of Foreign Academies of Science and Learned Societies and from Foreign Journals, 3, edited by Richard Taylor . London, 666–690.

Lowe, E. Jonathan. 2005. How are ordinary objects possible? *The Monist* 88 (4): 510–533.

Lucas, John R. 1961. Minds, machines and Gödel. *Philosophy*: 112–127.

Lucas, John R. 2003. The Gödelian argument: turn over the page. *Etica e Politica* 5 (1): 1.

Lucas, Peter , and Linda Van Der Gaag . 1991. Principles of expert systems. Wokingham: Addison-Wesley.

Luo, Huaishao , Lei Ji , 2020a. UniVL: a unified video and language pre-training model for multimodal understanding and generation. arXiv: 2002.06353.

Luo, Yuan , Alal Eran , 2020b. A multidimensional precision medicine approach identifies an autism subtype characterized by dyslipidemia. *Nature Medicine* 26 (9): 1375–1379.

Lyons, John. 1977. Deixis, space and time. In *Semantics*, 636–724. Cambridge: Cambridge University Press.

Mackie, John L. 1977. *Inventing right and wrong*. London: Penguin.

Maier, John. 2020. Abilities. In *The Stanford encyclopedia of philosophy*, Winter 2020, <https://plato.stanford.edu/entries/abilities/>.

Manning, Christopher D. , and Hinrich Schütze . 1999. *Foundations of statistical natural language processing*. Cambridge, MA: MIT Press.

Manolio, Teri A. , Francis S. Collins , 2009. Finding the missing heritability of complex diseases. *Nature* 461: 747–753.

Marcus, Gary. 2018. Deep learning: a critical appraisal. arXiv 1801.00631.

Marcus, Gary , and Ernest Davis . 2019. *Rebooting AI: building artificial intelligence we can trust*. New York: Vintage.

Marcus, Gary , and Ernest Davis . 2020. GPT-3, Bloviator: openAI's language generator has no idea what it's talking about. <https://www.technologyreview.com/2020/08/22/1007539/gpt3-openai-language-generator-artificial-intelligence-ai-opinion/>.

Marcus, Gary , and Ernest Davis . 2021. Has AI found a new foundation? *The Gradient*. <https://thegradient.pub/has-ai-found-a-new-foundation>.

Margolin, Emmanuel A. , Richard Strasser , 2020. Engineering the plant secretory pathway for the production of next-generation pharmaceuticals. *Trends in Biotechnology* 28: 1034–1044.

Marshak, Alexander , and Anthony Davis . 2005. *3D radiative transfer in cloudy atmospheres*. Berlin: Springer.

Mathew, Sarah , and Robert Boyd . 2011. Punishment sustains large-scale cooperation in prestate warfare. *PNAS* 108 (28): 11375–11380.

Matthiessen, Christian M. , and Abhishek K. Kashyap . 2014. The construal of space in different registers: an exploratory study. *Language Sciences* 45: 1–27.

McCann, Bryan , James Bradbury , 2017. Learned in translation: contextualized word vectors. arXiv abs/1708.00107.

McCauley, Joseph L. 1993. *Chaos, dynamics, and fractals: an algorithmic approach to deterministic chaos*. Cambridge: Cambridge University Press.

McCauley, Joseph L. 2009. *Dynamics of markets: the new financial economics*. 2nd ed. Cambridge: Cambridge University Press.

McCulloch, Warren S. , and Walter Pitts . 1943. A logical calculus of the ideas immanent in nervous activity. *The Bulletin of Mathematical Biophysics* 5 (4): 115–133.

McDermott, Drew. 1976. Artificial intelligence meets natural stupidity. *ACM Sigart Bulletin*, 57: 4–9.

McFadden, Daniel. 1963. Constant elasticity of substitution production functions. *The Review of Economic Studies* 30 (2): 73–83.

McLaughlin, Brian , and Karen Bennett . 2018. Supervenience. In *The Stanford encyclopedia of philosophy*, Winter 2018, <https://plato.stanford.edu/entries/supervenience/>.

McShane, Marjorie , and Sergei Nirenburg . 2021. *Linguistics for the age of AI*. Cambridge, MA: MIT Press.

Medin, Doug , and Brian H. Ross . January 1989. The specific character of abstract thought: categorization, problem solving, and induction. *Advances in the Psychology of Human Intelligence* 5: 189–223.

Meibauer, Jörg. 2001. *Pragmatik*. Tübingen: Stauffenburg.

Merkle, Ralph C. 2013. Uploading. In *The transhumanist reader*, edited by Max More and Natasha Vita-More , 157–164. Oxford: Wiley-Blackwell.

Merleau-Ponty, Maurice. 2012. *Phenomenology of perception* [1945]. London: Routledge.

Merrell, Eric , David Limbaugh , 2021. Capabilities. <https://philpapers.org/rec/MERC-14>.

Milkowski, Marcin. 2013. Explaining the computational mind. Cambridge, MA: MIT Press.

Mill, John S. 1863. Utilitarianism. London: Parker, Son/Bourn.

Millikan, Ruth Garrett. 2018. Biosemantics and words that don't represent. *Theoria* 84 (3): 229–241.

Mills, John A. 2012. Behaviorism. In *Encyclopedia of the history of psychological theories*, edited by Robert W. Rieber , 98–110. New York: Springer.

Milton, Katharine. 2000. Quo vadis? Tactics of food search and group movement in primates and other animals. In *On the move: how and why animals travel in groups*, edited by Sue Boinski and Paul A. Graber , 375–418. Chicago: Chicago University Press.

Min, Erxue , Xifeng Guo , 2018. A survey of clustering with deep learning: from the perspective of network architecture. *IEEE Access* 6: 39501–39514.

Minsky, Marvin L. 1991. Logical versus analogical or symbolic versus connectionist or neat versus scruffy. *AI Magazine* 12 (2): 34–34.

Mises, Ludwig von. 1936. *Socialism: an economic and sociological analysis*. London: Jonathan Cape.

Mitchell, Melanie. 2009. *Complexity: a guided tour*. Oxford: Oxford University Press.

Moor, James . August 2006. The nature, importance, and difficulty of machine ethics. *IEEE Intelligent Systems* 21: 18–21.

Moor, James . 2009. Four kinds of ethical robots. *Philosophy Now* 72.

Moortgat, Michael. 1997. Categorical type logics. In *Handbook of logic and language*, edited by J. van Benthem and A. ter Meulen . London: Elsevier.

Moosavi-Dezfooli, Seyed-Mohsen , Alhussein Fawzi , 2016. Universal adversarial perturbations. *CoRR abs/1610.08401*.

Mora, Peter T. 1963. Urge and molecular biology. *Nature (Berlin)* 199 (4890): 212–219.

Morgenstern, Martin. 1997. Nicolai Hartmann zur Einführung. Hamburg: Junius-Verlag.

Mouritsen, Henrik , Dominik Heyers , and Onur Güntürkün . 2016. The neural basis of long-distance navigation in birds. *Annual Review of Physiology* 78: 133–154.

Muehlhauser, Luke. 2013. What is AGI? <https://intelligence.org/2013/08/11/what-isagi/>.

Muehlhauser, Luke , and Louie Helm . 2012a. The singularity and machine ethics. In *Singularity hypotheses: a scientific and philosophical assessment*, edited by Amnon H. Eden and James H. Moor , 101–125. Dordrecht: Springer.

Muehlhauser, Luke , and Anna Salamon . 2012b. Intelligence explosion: evidence and import. In *Singularity hypotheses: a scientific and philosophical assessment*, edited by Amnon H. Eden and James H. Moor , 15–40. Dordrecht: Springer.

Müller, Stefan. 2016. *Grammatical theory: from transformational grammar to constraint-based approaches*. Berlin: Language Science Press.

Mulligan, Kevin. 1987. Promisings and other social acts: their constituents and structure. In *Speech act and Sachverhalt. Reinach and the foundations of realist phenomenology*, edited by Kevin Mulligan . New York: Springer.

Mulligan, Kevin. 1995. Perception. In *The Cambridge companion to Husserl*, edited by David Woodruff Smith and Barry Smith . Cambridge: Cambridge University Press.

Mulligan, Kevin. 1998. From appropriate emotions to values. *The Monist* 81 (1): 161–188.

Mulligan, Kevin. 2018. How to marry phenomenology and pragmatism. Scheler's proposal. In *Pragmatism and the European traditions*, edited by Maria Baghramian and Sarin Marchetti , 37–64. London: Routledge.

Mulligan, Kevin , Peter Simons , and Barry Smith . 2006. What's wrong with contemporary philosophy? *Topoi* 25 (1–2): 63–67.

Nagel, Thomas. 1975. What is it like to be a bat? *The Philosophical Review* 83 (4): 435–450.

Nandakumaran, A. K. , and P. S. Datti . 2020. *Partial differential equations: classical theory with a modern touch*. Cambridge, MA: Cambridge University Press.

Neil, Daniel , Michael Pfeiffer , and Shih-Chii Liu . 2016. Phased LSTM: accelerating recurrent network training for long or event-based sequences. *arXiv abs/1610.09513*.

Newell, Allen , and Herbert A. Simon . 1976. Computer science as empirical inquiry: symbols and search. *Communications of the ACM* 19 (3): 113–126.

Newman, Mark E. J. 2005. Power laws, Pareto distributions and Zipf's law. *Contemporary Physics* 46 (5): 323–351.

Nielsen, Michael , and Isaac Chuang . 2010. Quantum computation and quantum information. New York: Cambridge University Press.

Nienhuys-Cheng, Shan-Hwei , and Ronald de Wolf . 2008. Foundations of inductive logic programming. Berlin: Springer.

Nietzsche, Friedrich. 1980. Über Wahrheit und Lüge im aussermoralischen Sinne. In Friedrich Nietzsche: sämtliche Werke, edited by Giorgio Colli and Mazzino Montinari . Vol. 1. Berlin: de Gruyter.

Nozick, Robert. 1985. Interpersonal utility theory. *Social Choice and Welfare* (Berlin) 3 (2): 161–179.

Nyíri, Kristóf. 2014. Image and time in the theory of gestures. In *Meaning and motoricity: essays on image and time*. Frankfurt a. M.: Peter Lang.

Oare, Steve . June 17, 2018. Practical brass physics to improve your teaching and playing. *Kansas Music Review*.

Ochando, Jordi , Farideh Ordikhani , 2020. Tolerogenic dendritic cells in organ transplantation. *Transplant International* 33 (2): 113–127.

O'Regan, J. Kevin , and Alva Noë . 2001. A sensorimotor account of vision and visual consciousness. *Behavioral and Brain Sciences* 24: 939–1031.

O'Shaughnessy, Brian. 1980. The will: a dual aspect theory. 2 vols. Cambridge: Cambridge University Press.

Osoba, Osonde A. , Benjamin Boudreaux , and Douglas Yeung . 2020. Steps towards value-aligned systems. In *Proceedings of the AAAI/ACM conference on AI, ethics, and society*, 332–336, <https://dl.acm.org/doi/10.1145/3375627.3375872>.

Ouimet, Kirk. 2020. The universe function. <https://kirkouimet.medium.com/the-universe-function-92012c0c67c5>.

Parsons, Talcott. 1949. The structure of social action [1937]. New York: Free Press.

Parsons, Talcott. 1951. The social system. London: Routledge.

Parsons, Talcott. 1968. Social interaction. In *International encyclopedia of the social sciences*, edited by David L. Sills , 429–440. New York: Macmillan.

Parsons, Talcott. 1971. The system of modern societies. Englewood Cliffs, NJ: Prentice-Hall.

Parzen, Emanuel. 2015. Stochastic processes. Mineola, NY: Courier Dover Publications.

Pautz, Adam , and Daniel Stoljar , editors. 2019. Poise, dispositions, and access consciousness: reply to Daniel Stoljar. In *Blockheads! Essays on Ned Block's philosophy of mind and consciousness*, 537–544. New York: MIT Press.

Pavone, Arianna , and Alessio Plebe . 2021. How neurons in deep models relate with neurons in the brain. *Algorithms* 14 (9): 272.

Pearl, Judea. 2020. The limitations of opaque learning machines. In *Possible minds: twenty-five ways of looking at AI*, edited by John Brockman , 13–19. New York: Penguin Books.

Pearl, Leavitt J. 2018. Popitz's imaginative variation on power as model for critical phenomenology. *Human Studies* 41 (3): 475–483.

Peirce, Charles Sanders. 1935. Some amazing mazes. In *Collected papers of Charles Sanders Peirce*, edited by Charles Hartshorne and Paul Weiss . Cambridge, MA: Harvard University Press.

Pennachin, Cassio , and Ben Goertzel . 2007. Contemporary approaches to artificial general intelligence. In *Artificial general intelligence*, edited by Ben Goertzel and Cassio Pennachin , 1–30. Berlin and Heidelberg: Springer-Verlag.

Penrose, Roger. 1994a. Mathematical intelligence. In *What is intelligence?*, edited by Jean Khalfa , 107–136. Cambridge: Cambridge University Press.

Penrose, Roger. 1994b. *Shadows of the mind*. Oxford: Oxford University Press.

Percival, Steven L. , Sladjana Malic , 2011. Introduction to biofilms. In *Biofilms and veterinary medicine*, 41–68. Berlin: Springer.

Perko, Lawrence. 2013. Differential equations and dynamical systems. 3rd ed. Berlin: Springer.

Petitot, Jean , and Barry Smith . 1990. New foundations for qualitative physics. In *Evolving knowledge in natural science and artificial intelligence*, edited by J. E. Tiles , G. T. McKee ,

and C. G. Dean , 231–249. London: Pitman Publishing.

Piccinini, Gualtiero. 2003. Alan Turing and the mathematical objection. *Minds and Machines* 13: 23–48.

Piccinini, Gualtiero. 2015. *Physical computation: a mechanistic account*. Oxford: Oxford University Press.

Piccinini, Gualtiero. 2017. Computation in physical systems. In *Stanford encyclopedia of philosophy*, Summer 2017 Edition, Edward N. Zalta (ed.), <https://plato.stanford.edu/archives/sum2017/entries/computation-physicalsystems/>.

Pierson, Hugh O. 2012. *Handbook of carbon, graphite, diamonds and fullerenes: processing, properties and applications*. Norwich, NY: William Andrew.

Pinker, Steven. 2020. Tech prophecy and the underappreciated causal power of ideas. In *Possible minds: twenty-five ways of looking at AI*, edited by John Brockman . New York: Penguin Books.

Popitz, Heinrich. 2017a. *Phenomena of power*. New York: Columbia University Press.

Popitz, Heinrich. 2017b. Social norms. *Genocide Studies and Prevention: An International Journal* 11: 3–12.

Prigogine, Ilya. 1955. *Introduction to thermodynamics of irreversible processes*. New York: Interscience Publishers.

Prigogine, Ilya , and René Lefever . 1973. Theory of dissipative structures. In *Synergetics*, edited by H. Haken , 124–135. Wiesbaden: Vieweg+Teubner Verlag.

Purcell, Edward M. 1977. Life at low Reynolds number. *American Journal of Physics* 45: 3–11.

Putnam, Hillary. 1975. The meaning of “meaning”. *Minnesota Studies in the Philosophy of Science* 7: 131–193.

Quine, Willard Van Orman. 1969. Ontological relativity. In *Ontological relativity and other essays*, 26–68. New York: Columbia University Press.

Rabinowitz, Neil C. , Frank Perbet , 2018. Machine theory of mind. CoRR abs/1802.07740.

Radford, Alec , Jeffrey Wu , 2018. Language models are unsupervised multitask learners. Technical report. <https://paperswithcode.com/paper/language-models-are-unsupervised-multitask>.

Rajpurkar, Pranav , Jian Zhang , November 2016. SQuAD: 100,000+ questions for machine comprehension of text. In *Proceedings of the 2016 conference on empirical methods in natural language processing*, 2383–2392. Austin, TX: Association for Computational Linguistics.

Ramesh, Aditya , Mikhail Pavlov , 2021. Zero-shot text-to-image generation. arXiv: 2102.12092.

Ramsey, Frank Plumpton. 1931. *The foundations of mathematics and other logical essays*. London: K. Paul, Trench, Trubner & Company, Limited.

Ramsey, William. 2019. Eliminative materialism. In *Stanford encyclopedia of philosophy*, <https://plato.stanford.edu/entries/materialism-eliminative/>.

Rao, Rajesh P. N. , and Dana H. Ballard . 1999. Predictive coding in the visual cortex: a functional interpretation of some extra-classical receptive-field effects. *Nature Neuroscience* 2 (1): 79–87.

Reich, Wendelin. 2010. Three problems of intersubjectivity—and one solution. *Sociological Theory* 28 (1): 40–63.

Reinach, Adolf. 1969. Concerning phenomenology. *The Personalist* 50 (2): 194–221.

Reinach, Adolf. 1989. *Sämtliche Werke: textkritische Ausgabe in 2 Bänden*. Edited by Karl Schuhmann and Barry Smith . Munich: Philosophia.

Reinach, Adolf. 2012. *The a priori foundations of the civil law [1913]*. Edited by John F. Crosby . Frankfurt am Main: ontos Verlag.

Rescorla, Michael. 2020. The computational theory of mind. In *The Stanford encyclopedia of philosophy*, Winter 2020, <https://plato.stanford.edu/entries/computational-mind/>.

Richerson, Peter , and Robert Boyd . 2005. *Not by genes alone: how culture transformed human evolution*. Chicago, IL: Chicago University Press.

Roberts, David , Drake Morgan , and Yu Liu . 2007. How to make a rat addicted to cocaine. *Progress in Neuro-Psychopharmacology and Biological Psychiatry* 31: 1614–1624.

Roberts, Larry S. , John Janovy , and Steve Nadler . 2013. *Foundations of parasitology*. New York: McGraw-Hill.

Robinson, A. , and A. Voronkov . 2001. *Handbook of automated reasoning*. Cambridge, MA: Elsevier Science.

Robinson, William. 2019. Epiphenomenalism. In *Stanford encyclopedia of philosophy*, <https://stanford.library.sydney.edu.au/entries/epiphenomenalism/>.

Rojas, Raúl. 1997. Konrad Zuse's legacy: the architecture of the Z1 and Z3. *IEEE Annals of the History of Computing* 19 (2): 5–16.

Rosch, Eleanor. 1975. Cognitive representation of semantic categories. *Journal of Experimental Psychology* 104: 192–233.

Rose, Michael C. 2013. Immortalist fictions and strategies. In *The transhumanist reader*, edited by Max More and Natasha Vita-More , 196–204. Oxford: Wiley-Blackwell.

Rothblatt, Martine. 2013. Mind is deeper than matter. In *The transhumanist reader*, edited by Max More and Natasha Vita-More , 317–326. Oxford: Wiley-Blackwell.

Rothschild, Daniel , and Seth Yalcin . 2017. On the dynamics of conversation. *Noûs* 51 (1): 24–48.

Russell, Bertrand. 1938. *Power: a new social analysis*. New York: Routledge.

Sacks, Harvey , Emanuel Schegloff , and Gail Jefferson . 1974. A simplest systematics for the organization of turn-taking for conversation. *Language* 50: 696–735.

Sandberg, Anders. 2013. Feasibility of whole brain emulation. In *Theory and philosophy of artificial intelligence*, edited by Vincent C. Müller , 251–264. Berlin: Springer.

Schaeffer, Jonathan , Neil Burch , 2007. Checkers is solved. *Science* 317 (5844): 1518–1522.

Schaffer, Jonathan. 2015. What not to multiply without necessity. *Australasian Journal of Philosophy* 93 (4): 644–664.

Schaffer, Jonathan. 2021. Ground functionalism. In *Oxford Studies in the Philosophy of Mind*, 171–207. Oxford: Oxford University Press.

Schegloff, Emanuel. 2000. Overlapping talk and the organization of turn-taking for conversation. *Language in Society* 29 (1): 1–63.

Schegloff, Emanuel. 2017. Conversation analysis. In *The Oxford handbook of pragmatics*, edited by Yan Huang , 435–449. London: Oxford University Press.

Scheler, Max. 1961. *Man's place in nature*. New York: The Noonday Press.

Scheler, Max. 1973. *Formalism in ethics and non-formal ethics of values*. Evanston: Northwestern University Press.

Scheler, Max. 2008. *The nature of sympathy*. Piscataway, NJ: Transaction Publishers.

Schlossberger, Matthias. 2016. The varieties of togetherness: Scheler on collective affective intentionality. In *The phenomenological approach to social reality. History, concepts, problems*, edited by Alessandro Salice and Hans B. Schmid , 173–195. Berlin: Springer.

Schmidhuber, Jürgen. 1990. Making the world differentiable: on using self-supervised fully recurrent neural networks for dynamic reinforcement learning and planning in non-stationary environments. Technical report. TU Munich.

Schmidhuber, Jürgen. 2007. Gödel machines: fully self-referential optimal universal self-improvers. In *Artificial general intelligence*, edited by Ben Goertzel and Cassio Pennachin , 199–226. Berlin and Heidelberg: Springer-Verlag.

Schmidhuber, Jürgen. 2012. Philosophers and futurists, catch up. *Journal of Consciousness Studies* 19 (1–2): 173–182.

Schmidt, Karen L. , and Jeffrey Cohn . 2001. Human facial expressions as adaptations: evolutionary questions in facial expression research. *American Journal of Physical Anthropology Suppl* 33: 3–24.

Schoggen, Phil. 1989. *Behavior settings: a revision and extension of Roger G. Barker's ecological psychology*. Stanford, CA: Stanford University Press.

Schopenhauer, Artur. 1986. *Die Welt als Wille und Vorstellung*. Frankfurt am Main: Suhrkamp.

Schuhmann, Karl , and Barry Smith . 1987. Questions: an essay in Daubertian phenomenology. *Philosophy and Phenomenological Research* 47: 353–384.

Schuhmann, Karl , and Barry Smith . 1990. Elements of speech act theory in the work of Thomas Reid. *History of Philosophy Quarterly* 7: 47–66.

Schuster, Heinz G. , and Wolfram Just . 2005. *Deterministic chaos*. 4th ed. New York: Wiley VCH.

Schwab, Klaus. 2017. *The fourth industrial revolution*. New York: Penguin Books.

Searle, John R. 1969. *Speech acts. An essay in the philosophy of language*. Cambridge, MA: Cambridge University Press.

Searle, John R. 1975. A taxonomy of illocutionary acts. In *Language, mind and knowledge*, edited by K. Gunderson , 344–369. Minneapolis: University of Minnesota Press.

Searle, John R. 1978. Literal meaning. *Erkenntnis* 13 (1): 207–228.

Searle, John R. 1980. Minds, brains, and programs. *Behavioral and Brain Sciences* 3: 417–457.

Searle, John R. 1983. *Intentionality: an essay in the philosophy of mind*. Cambridge, MA: Cambridge University Press.

Searle, John R. 1990. Collective intentions and actions. In *Intentions in communication*, edited by Philip R. Cohen , Jerry Morgan , and Martha Pollack , 401–415. Cambridge, MA: MIT Press.

Searle, John R. 1992. *The rediscovery of the mind*. Cambridge, MA: MIT Press.

Searle, John R. 1995. *The construction of social reality*. New York: Simon/Schuster.

Searle, John R. 1998. How to study consciousness scientifically. *Brain Research Reviews* 26: 379–387.

Searle, John R. 2002. Why I am not a property dualist. *Journal of Consciousness Studies* 9 (12): 57–64.

Searle, John R. October 2014. What your computer can't know. *New York Review of Books*.

Seirin-Lee, Sungrim , T. Sukekawa , 2020. Transitions to slow or fast diffusions provide a general property for in-phase or anti-phase polarity in a cell. *Journal of Mathematical Biology* 80 (6): 1885.

Sellars, Wilfrid. 1963. Philosophy and the scientific image of man. In *Science, perception and reality*, 1–40. New York: Humanities Press.

Shapiro, Scott J. 2014. Massively shared agency. In *Rational and social agency: the philosophy of Michael Bratman*, 257–293. Oxford: Oxford University Press.

Sidnell, J. , and N. J. Enfield . 2017. Deixis and the interactional foundations of reference. In *The Oxford handbook of pragmatics*, edited by Yan Huang , 217–239. London: Oxford University Press.

Siegelmann, Hava T. , and Eduardo D. Sontag . 1994. Analog computation via neural networks. *Theoretical Computer Science* 131: 331–360.

Silberstein, Michael , and Anthony Chemero . 2012. Complexity and extended phenomenological-cognitive systems. *Topics in Cognitive Science* 4 (1): 35–50.

Silver, David , Demis Hassabis , 2016. Mastering the game of Go with deep neural networks and tree search. *Nature* 529 (7587): 484–489.

Silver, David , Thomas Hubert , 2018. A general reinforcement learning algorithm that masters chess, shogi, and Go through self-play. *Science* 362 (6419): 1140–1144.

Smaili, Fatima Zohra , Xin Gao , and Robert Hoehndorf . 2019. Formal axioms in biomedical ontologies improve analysis and interpretation of associated data. *Bioinformatics* 36 (7): 2229–2236.

Smith, Adam. 1790. *Theory of moral sentiments, or an essay towards an analysis of the principles by which men naturally judge concerning the conduct and character*. 6th ed. The Strand, London: A. Strahan/T. Cadell.

Smith, Barry. 1990. Towards a history of speech act theory. In *Speech acts, meanings and intentions. Critical approaches to the philosophy of John R. Searle*, edited by A. Burkhardt , 29–61. Berlin: de Gruyter.

Smith, Barry. 1992. An essay on material necessity. *Canadian Journal of Philosophy*: 301–322.

Smith, Barry. 1994. Fiat objects. In *Parts and wholes: conceptual part-whole relations and formal mereology*, 11th European conference on artificial intelligence, edited by Nicola Guarino , Laure Vieu , and Simone Pribbenow , 14–22. Amsterdam: European Coordinating

Committee for Artificial Intelligence.

Smith, Barry. 1995. Formal ontology, common sense, and cognitive science. *International Journal of Human-Computer Studies* 43 (5–6): 641–667.

Smith, Barry. 1997. Realistic phenomenology. In *Encyclopedia of phenomenology*, edited by Lester Embree, 586–590. Dordrecht/Boston/London: Kluwer Academic Publishers.

Smith, Barry. 1999. Truth and the visual field. In *Naturalizing phenomenology: issues in contemporary phenomenology and cognitive science*, edited by Jean Petitot, F. J. Varela, 317–329. Stanford: Stanford University Press.

Smith, Barry. 2000. Objects and their environments: from Aristotle to ecological ontology. In *Life and motion of socio-economic units*, 84–102. London/New York: Taylor and Francis.

Smith, Barry. 2003. John Searle: from speech acts to social reality. In John Searle, edited by Barry Smith, 1–33. Cambridge: Cambridge University Press.

Smith, Barry. 2005. Against fantology. In *Experience and analysis*, edited by J. Marek and E. M. Reicher, 153–170. Vienna: ÖBV & HPT Verlag.

Smith, Barry. 2008. Searle and de Soto: the new ontology of the social world. In *The mystery of capital and the construction of social reality*, edited by Barry Smith, David Mark, and Isaac Ehrlich, 35–51. Chicago: Open Court.

Smith, Barry. 2013. Diagrams, documents, and the meshing of plans. In *How to do things with pictures: skill, practice, performance*, edited by András Benedek and Kristóf Nyíri, 165–179. Frankfurt a. M.: Peter Lang.

Smith, Barry. 2021. Making space: the natural, cultural, cognitive and social niches of human activity. *Cognitive Processing* 22 (1): 77–87.

Smith, Barry, Michael Ashburner, 2007. The OBO foundry: coordinated evolution of ontologies to support biomedical data integration. *Nature Biotechnology* 25 (11): 1251–1255.

Smith, Barry, and Berit Brogaard. 2000. A unified theory of truth and reference. *Logique et Analyse*: 169/170, 49–93.

Smith, Barry. 2002. Quantum mereotopology. *Annals of Mathematics and Artificial Intelligence* 36 (1): 153–175.

Smith, Barry, and Roberto Casati. 1994. Naïve physics: an essay in ontology. *Philosophical Psychology* 7 (2): 227–247.

Smith, Barry, and Werner Ceusters. 2010. Ontological realism: a methodology for coordinated evolution of scientific ontologies. *Applied Ontology* 5 (3–4): 139–188.

Smith, Barry, Olimpia Giuliana Loddo, and Giuseppe Lorini. 2020. On credentials. *Journal of Social Ontology* 6 (1): 47–67.

Smith, Barry, and Achille C. Varzi. 1999. The niche. *Noûs* 33 (2): 214–238.

Smith, Barry, and Achille C. Varzi. 2002. Surrounding space. *Theory in Biosciences* 121 (2): 139–162.

Smolensky, Paul. 1990. Tensor product variable binding and the representation of symbolic structures in connectionist systems. *Artificial Intelligence* 46 (1): 159–216.

Solomon, Karen O., Doug Medin, and Elizabeth Lynch. 1999. Concepts do more than categorize. *Trends in Cognitive Sciences* 3: 99–105.

Spaulding, Shannon. 2013. Mirror neurons and social cognition. *Mind and Language* 28 (2): 233–257.

Spear, Andrew, Werner Ceusters, and Barry Smith. 2016. Functions in Basic Formal Ontology. *Applied Ontology* 11 (2): 103–128.

Spelke, Elizabeth S. 1990. Principles of object perception. *Cognitive Science* 14 (1): 29–56.

Spelke, Elizabeth S. 2000. Core knowledge. *American Psychologist* 55 (11): 1233–1243.

Spelke, Elizabeth S., and Linda Hermer. 1996. Early cognitive development: objects and space. In *Perceptual and Cognitive Development*, 71–114. San Diego and London: Academic Press.

Spelke, Elizabeth S., and Katherine D. Kinzler. 2007. Core knowledge. *Developmental Science* 10 (1): 89–96.

Sperber, Dan, and Nicolas Claidière. 2008. Defining and explaining culture (comments on Richerson and Boyd, *Not by Genes Alone*). *Biology and Philosophy* 23 (2): 283–292.

Stalnaker, Robert. 2012. *Mere possibilities: metaphysical foundations of modal semantics*. Princeton, NJ: Princeton University Press.

Steele, David Ramsay. 2013. *From Marx to Mises: post-capitalist society and the challenge of economic calculation*. Chicago: Open Court.

Stern, William. 1920. *Die Intelligenz der Kinder und Jugendlichen und Methoden ihrer Untersuchung*. Leipzig: Barth.

Stevenson, Charles Leslie. 1938. Persuasive definitions. *Mind* 47 (187): 331–350.

Stivers, Tanya , N. J. Enfield , 2009. Universals and cultural variation in turn-taking in conversation. *Proceedings of the National Academy of Sciences* 106 (26): 10587–10592.

Stoljar, Daniel. 2017. Physicalism. In *Stanford encyclopedia of philosophy*, <https://plato.stanford.edu/entries/physicalism/>.

Stooke, Adam , Anuj Mahajan , 2021. Open-ended learning leads to generally capable agents. *CoRR abs/2107.12808*. <https://arxiv.org/abs/2107.12808>.

Strachan, Tom , and Andrew Read . 2018. *Human molecular genetics*. London: Chapman & Hall/CRC.

Strecker, Jonathan , Sara Jones , 2019. Engineering of CRISPR-Cas12b for human genome editing. *Nature Communications* 10: 212.

Strickland, Eliza. 2019. How IBM Watson overpromised and underdelivered on AI health care. <https://spectrum.ieee.org/biomedical/diagnostics/how-ibm-watson-overpromised-and-underdelivered-on-ai-health-care>.

Strychalski, Elizabeth A. , Clyde A. Hutchinson III , 2016. Design and synthesis of a minimal bacterial genome. *Science* 351 (6280): aad6253.

Sutton, Richard S. , and Andrew G. Barto . 2018. *Reinforcement learning: an introduction*. Cambridge, MA: The MIT Press.

Swat, Maciej J. , Pierre Grenon , and Sarala Wimalaratne . 2016. ProbOnto: ontology and knowledge base of probability distributions. *Bioinformatics* 32 (17): 2719–2721.

Syropoulos, Apostolos . 2008. *Hypercomputation. Computing beyond the Church-Turing barrier*. Boston, MA: Springer.

Talmy, Leonard. 2018. *The targeting system of language*. Cambridge, MA: MIT Press.

Talmy, Leonard. 2021. Structure within morphemic meaning. *Cognitive Semantics* 7 (2).

Tam, Vivian , Nikunj Patel , 2019. Benefits and limitations of genome-wide association studies. *Nature Reviews Genetics* 20: 467–486.

Taylor, Charles. 1985. The concept of a person. In *Philosophical papers, Volume 1: human agency and language*, Cambridge: Cambridge University Press. 97–114.

Thomsen, Erik , and Barry Smith . 2018. Ontology-based fusion of sensor data and natural language. *Applied Ontology* 13: 295–333.

Thurner, Stefan , Peter Klimek , and Rudolf Hanel . 2018. *Introduction to the theory of complex systems*. Oxford: Oxford University Press.

Turing, Alan. 1937. On computable numbers, with an application to the Entscheidungsproblem. *Proceedings of the London Mathematical Society* 42 (1): 230–265.

Turing, Alan. 1950. Computing machinery and intelligence. *Mind* LIX: 433–460.

Turing, Alan. 1951. Can digital computers think? <http://www.turingarchive.org/browse.php/B/5>.

Turing, Alan. 1952. The chemical basis of morphogenesis. *Philosophical Transactions of the Royal Society of London Series B* 237 (641): 37–72.

Valiant, Leslie G. 1984. A theory of the learnable. *Communications of the ACM* 27 (11): 1134–1142.

Vallor, Shannon. 2016. *Technology and the virtues: a philosophical guide to a future worth wanting*. Oxford: Oxford University Press.

Vanderstraeten, Raf. 2002. Parsons, Luhmann and the theorem of double contingency. *Journal of Classical Sociology* 2 (1): 77–92.

Vaswani, Ashish , Noam Shazeer , 2017. Attention is all you need. *arXiv abs/1706.03762*.

Verschueren, Jef. 1999. *Understanding pragmatics*. London and New York: Arnold.

Vetter, Barbara. 2020. Perceiving potentiality: a metaphysics for affordances. *Topoi* 39 (5): 1177–1191.

Vinge, Vernor. 1993. The coming technological singularity. <https://edoras.sdsu.edu/~vinge/misc/singularity.html>.

Vinge, Vernor. 2013. Technological singularity. In *The transhumanist reader*, edited by Max More and Natasha Vita-More, 365–375. Oxford: Wiley-Blackwell.

Vogt, Lars, Peter Grobe. 2012. Fiat or bona fide boundary—a matter of granular perspective. *PLoS ONE* 7 (12): e48603.

Wallach, Wendell, and Colin Allen. 2008. *Moral machines: teaching robots right from wrong*. Oxford: Oxford University Press.

Walsh, Toby. 2017. The singularity may never be near. *AI Magazine* 38 (3): 58–62.

Wang, Chaoyue, Chang Xu. 2019. Evolutionary generative adversarial networks. *IEEE Transactions on Evolutionary Computation* 23 (6): 921–934.

Warwick, Kevin, and Huma Shah. 2016. Can machines think? A report on Turing test experiments at the Royal Society. *Journal of Experimental & Theoretical Artificial Intelligence* 28 (6): 989–1007.

Weber, Max. 1976. *Wirtschaft und Gesellschaft*. Tübingen: Mohr Siebeck.

Weber, Max. 1988. *Gesammelte Aufsätze zur Wissenschaftslehre*. Tübingen: J.C.B. Mohr.

Weinstein, Michael, and Jürgen Schmidhuber. 2017. Von Menschen, Stachelschweinen und Robotern. *Schweizer Monat Dossier: Big Data und KI*. <https://schweizermonat.ch/von-menschen-stachelschweinen-und-robotern/>

Werner, Konrad. 2020. Enactment and construction of the cognitive niche: toward an ontology of the mind-world connection. *Synthese* 197 (3): 1313–1341.

Werner, Sven. 2017. International review of district heating and cooling. *Energy* 137: 617–631.

Whitehead, Alfred North. 1911. *An introduction to mathematics*. London: Williams/Norgate.

Wiles, Andrew John. 1995. Modular elliptic curves and Fermat's last theorem. *Annals of Mathematics* 141 (3): 443–551.

Williams, Bernard. 1985. *Ethics and the limits of philosophy*. Cambridge, MA: Harvard University Press.

Williamson, Timothy. 2005. Contextualism, subject-sensitive invariantism and knowledge of knowledge. *The Philosophical Quarterly* 55 (219): 213–235.

Wilson, Edward O. 2000. *Sociobiology: the new synthesis*. Cambridge, MA: Harvard University Press.

Wimsatt, William C. 1994. The ontology of complex systems: levels of organization, perspectives, and causal thickets. *Canadian Journal of Philosophy* 24 (Sup 1): 207–274.

Wisniewski, Jeremy. 2019. Affordances, embodiment, and moral perception: a sketch of a moral theory. *Philosophy in the Contemporary World* 25 (1): 35–48.

Wittgenstein, Ludwig. 2003. *Philosophical investigations: the German text, with a revised English translation*. Malden, MA: Blackwell.

Wojtyła, Karol. 1979. *Primat des Geistes: Philosophische Schriften*. Degerloch, Stuttgart: Seewald Verlag.

Wolpert, David H., and William G. Macready. 1997. No free lunch theorems for optimization. *IEEE Transactions on Evolutionary Computation* 1 (1): 67–82.

Wu, Jiayan, Jingfa Xiao. 2014. Ribogenomics: the science and knowledge of RNA. *Genomics, Proteomics & Bioinformatics* 12 (2): 57–63.

Yudkowsky, Eliezer. 2001a. *Coherent extrapolated volition*. San Francisco: The Singularity Institute.

Yudkowsky, Eliezer. 2001b. *Creating friendly AI 1.0: the analysis and design of benevolent goal architectures*. San Francisco: The Singularity Institute.

Yudkowsky, Eliezer. 2008a. Artificial intelligence as a positive and negative factor in global risk. In *Global catastrophic risks*, edited by Nick Bostrom and Milan M. Cirkovic, 308–345. New York: Oxford University Press.

Yudkowsky, Eliezer. 2008b. *Efficient cross-domain optimization*. <https://www.lesswrong.com/posts/yLeEPFnB9wE7KLx2/efficient-cross-%20domain-optimization>.

Zaibert, Leo. 2018. *Rethinking punishment*. Cambridge: Cambridge University Press.

Zhou, Li, Jianfeng Gao. 2020. The design and implementation of Xiaolce, an empathetic social chatbot. *Computational Linguistics* 46 (1): 53–93.

Ziv Konzelmann, Anita. 2013. Introduction. In *Institutions, emotions, and group agents: contributions to social ontology*, edited by Anita Ziv Konzelmann and Hans Bernhard Schmid, 1–15. Berlin: Springer.