

Article

Open Problems in Universal Induction & Intelligence

Marcus Hutter

Research School of Information Sciences and Engineering (RSISE), Australian National University,
and Statistical Machine Learning (SML), NICTA, Canberra, ACT, 0200, Australia
E-mail: marcus@hutter1.net; <http://www.hutter1.net>

Received: 8 April 2009; in revised form: 15 June 2009 / Accepted: 16 June 2009 /

Published: 2 July 2009

Abstract: Specialized intelligent systems can be found everywhere: finger print, handwriting, speech, and face recognition, spam filtering, chess and other game programs, robots, et al. This decade the first presumably complete *mathematical* theory of artificial intelligence based on universal induction-prediction-decision-action has been proposed. This information-theoretic approach solidifies the foundations of inductive inference and artificial intelligence. Getting the foundations right usually marks a significant progress and maturing of a field. The theory provides a gold standard and guidance for researchers working on intelligent algorithms. The roots of universal induction have been laid exactly half-a-century ago and the roots of universal intelligence exactly one decade ago. So it is timely to take stock of what has been achieved and what remains to be done. Since there are already good recent surveys, I describe the state-of-the-art only in passing and refer the reader to the literature. This article concentrates on the open problems in universal induction and its extension to universal intelligence.

Keywords: Kolmogorov complexity; information theory; sequential decision theory; reinforcement learning; artificial intelligence; universal Solomonoff induction; rational agents

“The mathematician is by now accustomed to intractable equations, and even to unsolved problems, in many parts of his discipline. However, it is still a matter of some fascination to realize that there are parts of mathematics where the very construction of a precise mathematical statement of a verbal problem is itself a problem of major difficulty.”

— Richard Bellman, *Adaptive Control Processes* (1961) p.194

1. Introduction

What is a good model of the weather changes? Are there useful models of the world economy? What is the true regularity behind the number sequence 1,4,9,16,...? What is the correct relationship between mass, force, and acceleration of a physical object? Is there a causal relation between interest rates and inflation? Are models of the stock market purely descriptive or do they have any predictive power?

Induction. The questions above look like a set of unrelated inquiries. What they have in common is that they seem to be amenable to scientific investigation. They all ask about a model for or relation between observations. The purpose seems to be to explain or understand the data. Generalizing from data to general rules is called *inductive inference*, a core problem in philosophy [1–3] and a key task of science [4–6].

But why do or should we care about modeling the world? Because this is what science is about [7]? As indicated above, models should be good, useful, true, correct, causal, predictive, or descriptive [8]. Digging deeper, we see that models are mostly used for prediction in related but new situations, especially for predicting future events [9].

Predictions. Consider the apparently only slight variation of the questions above: What is the correct answer in an IQ test asking to continue the sequence 1,4,9,16,...? Given historic stock-charts, can one predict the quotes of tomorrow? Or questions like: Assuming the sun rose every day for 5000 years, how likely is doomsday (that the sun will not rise) tomorrow? What is my risk of dying from cancer next year?

These questions are instances of the important problem of time-series forecasting, also called sequence prediction [10, 11]. While inductive inference is about finding models or hypotheses that *explain* the data (whatever *explain* actually shall mean), *prediction* is concerned about forecasting the future. Finding models is interesting and useful, since they usually help us to (partially) answer such *predictive* questions [12, 13]. While the usefulness of predictions is clearer to the layman than the purpose of the scientific inquiry for models, one may again ask, why we do or should we care about making predictions?

Decisions. Consider the following questions: Shall I take my umbrella or wear sunglasses today? Shall I invest my assets in stocks or bonds? Shall I skip work today because it might be my last day on earth? Shall I irradiate or remove the tumor of my patient? These questions ask for decisions that have some (minor to drastic) consequences. We usually want to make “good” decisions, where the quality is measured in terms of some reward (money, life expectancy) or loss [14–16]. In order to compute this reward as a function of our decision, we need to predict the environment: whether there will be rain or sunshine today, whether the market will go up or down, whether doomsday is tomorrow, or which type of cancer the patient has. Often forecasts are uncertain [17], but this is still better than no prediction. Once we arrived at a (hopefully good) decision, what do we do next?

Actions. The obvious thing is to execute the decision, i.e. to perform some action consistent with the decision arrived at. The action may not influence the environment, like taking umbrella versus sunglasses does not influence the future weather (ignoring the butterfly effect) or small stock trades. These settings are called passive [18], and the action part is of marginal importance and usually not discussed. On the other hand, a patient might die from a wrong treatment, or a chess player loses a figure and possibly

the whole game by making one mistake. These settings are called (re)active [19], and their analysis is immensely more involved than the passive case [20].

And now? There are many theories and algorithms and whole research fields and communities dealing with some aspects of induction, prediction, decision, or action. Some of them will be detailed below. Finding solutions for every particular (new) problem is possible and useful for many specific applications. Trouble is that this approach is cumbersome and prone to disagreement or contradiction [21]. Some researchers feel that this is the nature of their discipline and one can do little about it [22]. But in science (in particular math, physics, and computer science) previously separate approaches are constantly being unified towards more and more powerful theories and algorithms [23, 24]. There is at least one field, where we *must* put everything (induction+prediction+decision+action) together in a completely formal (preferably elegant) way, namely Artificial Intelligence [25]. Such a general and formal theory of AI has been invented about a decade ago [26].

Contents. In *Section 2*. I will give a brief introduction into this universal theory of AI. It is based on an unexpected unification of algorithmic information theory and sequential decision theory. The corresponding AIXI agent is the first sound, complete, general, rational agent in any relevant but unknown environment with reinforcement feedback [27, 28]. It is likely the best possible such agent in a sense to be explained below.

Section 3. describes the historic origin of the AIXI model. One root is Solomonoff's theory [29] of *universal induction*, which is closely connected to algorithmic complexity. The other root is Bellman's *adaptive control theory* [30] for optimal *sequential decision* making. Both theories are now half-a-century old. From an algorithmic information theory perspective, AIXI generalizes optimal passive universal induction to the case of active agents. From a decision-theoretic perspective, AIXI is a universal Bayes-optimal learning algorithm.

Sections 4.–7. constitute the core of this article describing the open problems around universal induction & intelligence. Most of them are taken from the book [27] and paper [31]. I focus on questions whose solution has a realistic chance of advancing the field. I avoid technical open problems whose global significance is questionable.

Solomonoff's half-a-century-old theory of universal induction is already well developed. Naturally, most remaining open problems are either philosophically or technically deep.

Its generalization to Universal Artificial Intelligence seems to be quite intricate. While the AIXI model itself is very elegant, its analysis is much more cumbersome. Although AIXI has been shown to be optimal in some senses, a convincing notion of optimality is still lacking. Convergence results also exist, but are much weaker than in the passive case.

Its construction makes it plausible that AIXI is the optimal rational general learning agent, but unlike the induction case, victory cannot be claimed yet. It would be natural, hence, to compare AIXI to alternatives, if there were any. Since there are no competitors yet, one could try to create some. Finally, AIXI is only “essentially” unique, which gives rise to some more open questions.

Given that AI is about designing intelligent systems, a serious attempt should be made to *formally* define intelligence in the first place. Astonishingly there have been not too many attempts. There is one definition that is closely related to AIXI, but its properties have yet to be explored.

The final Section 8. briefly discusses the flavor, feasibility, difficulty, and interestingness of the raised questions, and takes a step back and briefly compares the information-theoretic approach to AI discussed in this article to others.

2. Universal Artificial Intelligence

Artificial Intelligence. The science of artificial intelligence (AI) may be defined as the construction of intelligent systems (*artificial agents*) and their analysis [25]. A natural definition of a *system* is anything that has an input and an output stream, or equivalently an agent that acts and observes. *Intelligence* is more complicated. It can have many faces like *creativity, solving problems, pattern recognition, classification, learning, induction, deduction, building analogies, optimization, surviving in an environment, language processing, planning, and knowledge acquisition and processing*. Informally, *AI is concerned with developing agents that perform well in a large range of environments* [32]. A formal definition incorporating every aspect of intelligence, however, seems difficult. In order to solve this problem we need to solve the induction, prediction, decision, and action problem, which seems like a daunting (some even claim impossible) task: Intelligent *actions* are based on informed *decisions*. Attaining good decisions requires *predictions* which are typically based on models of the environments. Models are constructed or learned from past observations via *induction*. Fortunately, based on the *deep philosophical insights* and *powerful mathematical developments* listed in Section 3., these problems have been overcome, at least in theory.

Universal Artificial Intelligence (UAI). Most, if not all, known facets of intelligence can be formulated as goal driven or, more precisely, as maximizing some reward or utility function. It is, therefore, sufficient to study goal-driven AI; e.g. the (biological) goal of animals and humans is to survive and spread. The goal of AI systems should be to be useful to humans. The problem is that, except for special cases, we know neither the utility function nor the environment in which the agent will operate in advance. What do we need (from a mathematical point of view) to construct a universal optimal learning agent interacting with an arbitrary unknown environment? The theory, coined *AIXI*, developed in this decade and explained in [27] says: *All you need is Occam* [33], *Epicurus* [34], *Turing* [35], *Bayes* [36], *Solomonoff* [37], *Kolmogorov* [38], and *Bellman* [30]: Sequential decision theory [20] (*Bellman's* equation) formally solves the problem of rational agents in uncertain worlds if the true environmental probability distribution is known. If the environment is unknown, *Bayesians* [39] replace the true distribution by a weighted mixture of distributions from some (hypothesis) class. Using the large class of all (semi)measures that are (semi)computable on a *Turing* machine bears in mind *Epicurus*, who teaches not to discard any (consistent) hypothesis. In order not to ignore *Occam*, who would select the simplest hypothesis, *Solomonoff* defined a universal prior that assigns high/low prior weight to simple/complex environments, where *Kolmogorov* quantifies complexity [40, 41]. All other concepts and phenomena attributed to intelligence are emergent. All together, this solves all conceptual problems [27], and “only” computational problems remain.

Kolmogorov complexity. Kolmogorov [38] defined the complexity of a string $x \in \mathcal{X}^*$ over some finite alphabet $\mathcal{X}$ as the length ℓ of a shortest description $p \in \{0,1\}^*$ on a universal Turing machine U :

$$\text{Kolmogorov complexity:} \quad K(x) := \min_p \{\ell(p) : U(p) = x\} \quad (1)$$

A string is simple if it can be described by a short program, like “the string of one million ones”, and is complex if there is no such short description, like for a random string whose shortest description is specifying it bit-by-bit. For non-string objects o one defines $K(o) := K(\langle o \rangle)$, where $\langle o \rangle \in \mathcal{X}^*$ is some standard code for o . Kolmogorov complexity [38, 42] is a key concept in (algorithmic) information theory [41]. An important property of K is that it is nearly independent of the choice of U , i.e. different choices of U change K “only” by an additive constant (see Section 4.h). Furthermore it leads to shorter codes than any other effective code. K shares many properties with Shannon’s entropy (information measure) S [43, 44], but K is superior to S in many respects. Foremost, K measures the information of individual outcomes, while S can only measure expected information of random variables. To be brief, K is an excellent universal complexity measure, suitable for quantifying Occam’s razor. The major drawback of K as complexity measure is its incomputability. So in practical applications it has always to be approximated, e.g. by Lempel-Ziv compression [45, 46], or by CTW [47] compression, or by using two-part codes like in MDL and MML, or by others.

Solomonoff induction. Solomonoff [37] defined (earlier) the closely related universal a priori probability $M(x)$ as the probability that the output of a universal (monotone) Turing machine U starts with x when provided with fair coin flips on the input tape [48]. Formally,

$$\text{Solomonoff prior:} \quad M(x) := \sum_{p: U(p)=x*} 2^{-\ell(p)}, \quad (2)$$

where the sum is over all (possibly non-halting) so-called minimal programs p which output a string starting with x . Since the sum is dominated by short programs, we have $M(x) \approx 2^{-K(x)}$ (formally $-\log M(x) = K(x) + O(\log \ell(x))$), i.e. simple/complex strings are assigned a high/low a-priori probability. A different representation is as follows [49]: Let $\mathcal{M} = \{\nu\}$ be a countable class of probability measures ν (environments) on infinite sequences $\mathcal{X}^\infty$, $\mu \in \mathcal{M}$ be the true sampling distribution, i.e. $\mu(x)$ is the true probability that an infinite sequences starts with x , and $\xi_{\mathcal{M}}(x) := \sum_{\nu \in \mathcal{M}} w_\nu \nu(x)$ be the w -weighted average called Bayesian mixture distribution. One can show that $M(x) = \xi_{\mathcal{M}_U}(x)$, where $\mathcal{M}_U$ includes all computable probability measures and $w_\nu = 2^{-K(\nu)}$. More precisely, $\mathcal{M}_U := \{\nu_1, \nu_2, \dots\}$ consists of an effective enumeration of all so-called lower semi-computable semi-measures ν_i , and $K(\nu_i) := K(i) := K(\langle i \rangle)$ [41].

M can be used as a universal sequence predictor, which outperforms in a strong sense all other predictors. Consider the classical online sequence prediction task: Given $x_{<t} \equiv x_{1:t-1} := x_1 \dots x_{t-1}$, predict x_t ; then observe the true x_t ; $t \rightsquigarrow t+1$; repeat. For $x_{1:\infty}$ generated by the unknown “true” distribution $\mu \in \mathcal{M}_U$, one can show [50] that the universal predictor $M(x_t | x_{<t}) := M(x_{1:t}) / M(x_{<t})$ rapidly converges to the true probability $\mu(x_t | x_{<t}) = \mu(x_{1:t}) / \mu(x_{<t})$ of the next observation $x_t \in \mathcal{X}$ given history $x_{<t}$. That is, M serves as an excellent predictor of any sequence sampled from any computable probability distribution.

The AIXI model. It is possible to write down the AIXI model explicitly in one line [19], although one should not expect to be able to grasp the full meaning and power from this compact representation.

AIXI is an agent that interacts with an environment in cycles $k = 1, 2, \dots, m$. In cycle k , AIXI takes action a_k (e.g. a limb movement) based on past perceptions $o_1 r_1 \dots o_{k-1} r_{k-1}$ as defined below. Thereafter, the environment provides a (regular) observation o_k (e.g. a camera image) to AIXI and a real-valued

reward r_k . The reward can be very scarce, e.g. just +1 (-1) for winning (losing) a chess game, and 0 at all other times. Then the next cycle $k+1$ starts. Given the above, AIXI is defined by:

$$\text{AIXI:} \quad a_k := \arg \max_{a_k} \sum_{o_k r_k} \dots \max_{a_m} \sum_{o_m r_m} [r_k + \dots + r_m] \sum_{q: U(q, a_1 \dots a_m) = o_1 r_1 \dots o_m r_m} 2^{-\ell(q)} \quad (3)$$

The expression shows that AIXI tries to maximize its total future reward $r_k + \dots + r_m$. If the environment is modeled by a deterministic program q , then the future perceptions $\dots o_k r_k \dots o_m r_m = U(q, a_1 \dots a_m)$ can be computed, where U is a universal (monotone Turing) machine executing q given $a_1 \dots a_m$. Since q is unknown, AIXI has to maximize its expected reward, i.e. average $r_k + \dots + r_m$ over all possible perceptions created by all possible environments q . The simpler an environment, the higher is its a-priori contribution $2^{-\ell(q)}$, where simplicity is measured by the length ℓ of program q . The inner sum $\sum_{q \dots} 2^{-\ell(q)}$ generalizes Solomonoff's a-priori distribution M by including actions. Since noisy environments are just mixtures of deterministic environments, they are automatically included. The sums in the formula constitute the averaging process. Averaging and maximization have to be performed in chronological order, hence the interleaving of max and Σ (similarly to minimax for games). The *value* V of AIXI (or any other agent) is its expected reward sum.

One can fix any finite action and perception space, any reasonable U , and any large finite lifetime m . This completely and uniquely defines AIXI's actions a_k , which are limit-computable via the expression above (all quantities are known).

That's it! Ok, not really. It takes a whole book and more to explain why AIXI likely is the most intelligent general-purpose agent and incorporates all aspects of rational intelligence. In practice, AIXI needs to be approximated. AIXI can also be regarded as the gold standard which other practical general purpose AI programs should aim at (analogue to minimax approximations/heuristics).

The role of AIXI for AI. The AIXI model can be regarded as the first *complete theory of AI*. Most if not all AI problems can easily be formulated within this theory, which reduces the conceptual problems to pure computational questions. Solving the conceptual part of a problem often causes a quantum leap forward in a field. Two analogies may help: QED is a *complete* theory of *all* chemical processes. ZFC solved the *conceptual* problems of sets (e.g. Russell's paradox).

From an algorithmic information theory (AIT) perspective, the AIXI model generalizes optimal passive universal induction to the case of active agents. From a decision-theoretic perspective, AIXI is a suggestion of a new (implicit) "learning" algorithm, which may overcome all (except computational) problems of previous reinforcement learning algorithms. If the optimality theorems of universal induction and decision theory generalize to the unified AIXI model, we would have, for the first time, a universal (parameterless) model of an optimal rational agent in any computable but unknown environment with reinforcement feedback.

Although deeply rooted in algorithm theory, AIT mainly neglects computation time and so does AIXI. It is important to note that this does not make the AI problem trivial. Playing chess optimally or solving NP-complete problems become trivial, but driving a car or surviving in nature do not. This is because it is a challenge itself to well-define the latter problems, not to mention presenting an algorithm. In other words: The AI problem has not yet been well defined (cf. the quote after the abstract). One may view AIXI as a suggestion of such a mathematical definition.

Although Kolmogorov complexity is uncomputable in general, Solomonoff's theory triggered an entire field of research on computable approximations. This led to numerous practical applications [51]. If the AIXI model should lead to a universal "active" decision maker with properties analogous to those of universal "passive" predictors, then we could expect a similar stimulation of research on resource-bounded, practically feasible variants. First attempt have been made to test the power and limitations of AIXI and downscaled versions like AIXI_{tl} and AI ξ [52, 53], as well as related models derived from basic concepts of algorithmic information theory.

So far, some remarkable and surprising results have already been obtained (see Section 3.). A 2, 12, 60, 300 page introduction to the AIXI model can be found in [19, 27, 54, 55], respectively, and a gentle introduction to UAI in [56].

3. History and State-of-the-Art

The theory of UAI and AIXI build on the theories of universal induction, universal prediction, universal decision making, and universal agents. From a historical and research-field perspective, the *AIXI model* is based on two otherwise unconnected fundamental theories:

- (1) The major basis is *Algorithmic information theory* [41], initiated by [37, 38, 57], which builds the foundation of complexity and randomness of individual objects. It can be used to quantify Occam's razor principle (use the simplest theory consistent with the data). This in turn allowed *Solomonoff* to come up with a *universal theory of induction* [37, 50].
- (2) The other basis is the theory of optimal *sequential decisions*, initiated by Von Neumann [58] and Bellman [30]. This theory builds the basis of modern *reinforcement learning* [59].

This section outlines the history and state-of-the-art of the theories and research fields involved in the AIXI model.

Algorithmic information theory (AIT). In the 1960's [37, 38, 57] introduced a new machine independent complexity measure for arbitrary computable data. The Kolmogorov complexity $K(x)$ is defined as the length of the shortest program on a universal Turing machine that computes x . It is closely related to Solomonoff's universal a-priori probability $M(x) \approx 2^{-K(x)}$ (see above), Martin-Löf randomness of individual sequences [60], time-bounded complexity [61], universal optimal search [62], the speed prior [63], the halting probability Ω [64], strong mathematical undecidability [65], generalized probability and complexity [66], algorithmic statistics [67–69], and others.

Despite its uncomputability, AIT found many applications in philosophy, practice, and science: The minimum message/description length (MML/MDL) principles [70–72] can be regarded as a practical approximation of Kolmogorov complexity. MML&MDL are widely used in machine learning applications [6, 73–77]. The latest, most direct and impressive applications are via the universal similarity metric [46, 78]. Schmidhuber produced another range of impressive applications to neural networks [79, 80], in search problems [81], and even in the fine arts [82]. By carefully approximating Kolmogorov complexity, AIT sometimes lead to results unmatched by other approaches. Besides these practical applications, AIT is used to simplify proofs via the incompressibility method, improves Shannon information, is used

in reversible computing, physical entropy and Maxwell daemon issues, artificial intelligence, and the asymptotically fastest algorithm for all well-defined problems [27, 40, 41, 83, 84].

Universal Solomonoff induction. How and in which sense induction is possible at all has been subject to long philosophical controversies [1, 27, 85]. Highlights are Epicurus' principle of multiple explanations [34], Occam's razor (simplicity) principle [33], and Bayes' rule for conditional probabilities [5, 36]. Solomonoff [37] elegantly unified these aspects with the concept of universal Turing machines [35] to one formal theory of inductive inference based on a universal probability distribution M , which is closely related to Kolmogorov complexity K ($M(x) \approx 2^{-K(x)}$). The theory allows for optimally predicting sequences without knowing their true generating distribution μ [50], and presumably solves the induction problem. The theory remained for more than 20 years at this stage, till the work on AIXI started, which resulted in a beautiful elaboration and extension of Solomonoff's theory.

Meanwhile, the (non)existence of universal priors for several generalized computability concepts [66, 86, 87] has been classified, rapid convergence of M to the unknown true environmental distribution μ [88] and tight error [89] and loss bounds for arbitrary bounded loss functions and finite alphabet [90, 91] have been proven, and (Pareto) optimality of M [18, 86] has been shown, exemplified on games of chance and compared to predictions with expert advice [18, 92]. The bounds have been further improved by introducing a version of Kolmogorov complexity that is monotone in the condition [93, 94]. Similar but necessarily weaker non-asymptotic bounds for universal deterministic/one-part MDL [95, 96] and discrete two-part MDL [97–100] have also been proven. Quite unexpectedly [101] M does *not* converge on all Martin-Löf random sequences [102], but there is a sophisticated remedy [103].

All together this shows that Solomonoff's induction scheme represents a universal (formal, but in-computable) solution to all *passive* prediction problems. The most recent studies [104] suggest that this theory could solve the induction problem at whole, or at least constitute a significant progress in this fundamental problem [31].

Sequential decision theory. Sequential decision theory provides a framework for finding optimal reward-maximizing strategies in *reactive* environments (e.g. chess playing as opposed to weather forecasting), assuming the environmental probability distribution μ is known. The Bellman equations [30] are at the heart of sequential decision theory [25, 58, 105]. The book [20] summarizes open problems and progress in infinite horizon problems. Sequential decision theory can deal with actions and observations depending on arbitrary past events. This general setup has been called $AI\mu$ model in [19, 27]. Optimality of $AI\mu$ is obvious by construction. This model reduces in special cases to a range of known models.

Reinforcement learning. If the true environmental probability distribution μ or the reward function are unknown, they need to be learned [59]. This dramatically complicates the problem due to the exploration $\leftrightarrow$ exploitation dilemma [27, 106–108]. In order to attack this intrinsically difficult problem, control theorists typically confine themselves to linear systems with quadratic loss function, relevant in the control of (simple) machines, but irrelevant for AI. There are notable exceptions to this confinement, e.g. the book [109] on stochastic adaptive control and [110, 111], and an increasing number of more recent work. Reinforcement learning (RL) (sometimes associated with temporal difference learning or neural nets) is the instantiation of stochastic adaptive control theory [109] in the machine learning com-

munity. Current research on RL is vast; the most important conferences are ICML, COLT, ECML, ALT, and NIPS; the most important journals are JMLR and MLJ. Some highlights and surveys are [108, 112–130] and [20, 59, 131–133] respectively. RL has been applied to a variety of real-world problems, occasionally with stunning success: Backgammon and Checkers [59, Chp.11], helicopter control [134], and others. Nevertheless, existing learning algorithms are very limited (typically to Markov domains), and non-optimal — from the very outset they are approximate or asymptotic only. Indeed, AIXI is currently the only general and rigorous mathematical formulation of the addressed problems.

The universal algorithmic agent AIXI. Reinforcement learning algorithms [59, 131, 135] are usually used in the case of unknown μ . They can succeed if the state space is either small or has effectively been made small by generalization techniques. The algorithms work only in restricted, (e.g. Markov) domains, have problems with optimally trading off exploration versus exploitation, have non-optimal learning rate, are prone to diverge, or are otherwise ad hoc.

The formal solution proposed in [27, 55] is to generalize the universal probability M to include actions as conditions and replace μ by M in the $AI\mu$ model, resulting in the AIXI model, which is presumably universally optimal. It is quite non-trivial what can be expected from a universally optimal agent and to properly interpret or define *universal*, *optimal*, etc [19]. It is known that M converges to μ also in case of multi-step lookahead as occurs in the AIXI model [136], and that a variant of AIXI is asymptotically self-optimizing and Pareto optimal [137, 138].

The book [27] gives a comprehensive introduction and discussion of previous achievements on or related to AIXI, including a critical review, more open problems, comparison to other approaches to AI, and philosophical issues.

Important environmental classes. In practice, one is often interested in specific classes of problems rather than the fully universal setting; for example we might be interested in evaluating the performance of an algorithm designed solely for function maximization. A taxonomy of abstract environmental classes from the mathematical perspective of interacting chronological systems [56, 139] has been established. The relationships between Bandit problems, MDP problems, ergodic MDPs, higher order MDPs, sequence prediction problems, function optimization problems, strategic games, classification, and many others are formally defined and explored therein. The work also suggests new abstract environmental classes that could be useful from an analytic perspective. In [27], each problem class is formulated in its natural way for known μ , and then a formulation within the $AI\mu$ model is constructed and their equivalence is shown. Then, the consequences of replacing μ by M are considered, and in which sense the problems are formally solved by AIXI.

Computational aspects. The major drawback of AIXI is that it is uncomputable, or more precisely, only asymptotically computable, which makes a direct implementation impossible. To overcome this problem, the AIXI model can be scaled down to a model coined $AIXI_{t,l}$, which is still superior to any other time t and length l bounded agent [27, 55]. The computation time of $AIXI_{t,l}$ is of the order $t \cdot 2^l$. A way of overcoming the large multiplicative constant 2^l is possible at the expense of an (unfortunately even larger) additive constant. The constructed algorithm builds upon Levin search [62, 140]. The algorithm is capable of solving all well-defined problems p as quickly as the fastest algorithm computing a solution to p , save for a factor of $1 + \varepsilon$ and lower-order additive terms [84]. The solution requires an implementa-

tion of first-order logic, the definition of a universal Turing machine within it and a proof theory system. The algorithm as it is, is only of theoretical interest, but there are more practical variations [81, 141]. A different, more limited but more practical scaled-down version (coined AI ξ) has been implemented and applied successfully to 2×2 matrix games like the notoriously difficult repeated prisoner problem and generalized variants thereof [52].

4. Open Problems in Universal Induction

The induction problem is a fundamental problem in philosophy [1, 5] and science [142]. Solomonoff's model is a promising universal solution of the induction problem. In [31], an attempt has been made to collect the most important fundamental philosophical and statistical problems, regarded as open, and to present arguments and proofs that Solomonoff's theory overcomes them. Despite the force of the arguments, they are likely not yet sufficient to convince the (scientific) world that the induction problem is solved. The discussion needs to be rolled out much further, say, at least one generally accessible article per one allegedly open problem. Indeed, this endeavor might even discover some catch in Solomonoff's theory. Some problems identified and outlined in [31] worth to investigate in more detail are:

- a) **The zero prior problem.** The problem is how to *confirm universal hypotheses* like $H :=$ “all balls in some urn (or all ravens) are black”. A natural model is to assume that balls (or ravens) are drawn randomly from an infinite population with fraction θ of black balls (or ravens) and to assume *some* prior *density* over $\theta \in [0;1]$ (a uniform density gives the Bayes-Laplace model). Now we draw n objects and observe that they are all black. The problem is that the posterior probability $P[H|\text{black}_1 \dots \text{black}_n] \equiv 0$, since the prior probability $P[H] = P[\theta = 1] \equiv 0$. Maher's [143] approach does not solve the problem [31].
- b) **The black raven paradox** by Carl Gustav Hempel goes as follows [144, Ch.11.4]: Observing *Black Ravens* confirms the hypothesis H that all ravens are black. In general, (i) hypothesis $R \rightarrow B$ is confirmed by R -instances with property B . Formally substituting $R \rightsquigarrow \neg B$ and $B \rightsquigarrow \neg R$ leads to (ii) hypothesis $\neg B \rightarrow \neg R$ is confirmed by $\neg B$ -instances with property $\neg R$. But (iii) since $R \rightarrow B$ and $\neg B \rightarrow \neg R$ are logically equivalent, $R \rightarrow B$ must also be confirmed by $\neg B$ -instance with property $\neg R$. Hence by (i), observing *Black Ravens* confirms Hypothesis H , so by (iii), observing *White Socks* also confirms that all Ravens are Black, since *White Socks* are non-Ravens which are non-Black. But this conclusion is absurd. Again, neither Maher's nor any other approach solves this problem.
- c) **The Grue problem** [145]. Consider the following two hypotheses: $H1 :=$ “All emeralds are green”, and $H2 :=$ “All emeralds found till year 2020 are green, thereafter all emeralds are blue”. Both hypotheses are equally well supported by empirical evidence. Occam's razor seems to favor the more plausible hypothesis $H1$, but by using new predicates *grue*:=“green till y2020 and blue thereafter” and *bleen*:=“blue till y2020 and green thereafter”, $H2$ gets simpler than $H1$.
- d) **Reparametrization invariance** [146]. The question is how to extend the symmetry principle from finite hypothesis classes (all hypotheses are equally likely) to infinite hypothesis classes. For

“compact” classes, Jeffrey’s prior [147] is a solution, but for non-compact spaces like $\mathbb{N}$ or $\mathbb{R}$, classical statistical principles lead to improper distributions, which are often not acceptable.

- e) **Old-evidence/updating problem** and *ad-hoc hypotheses* [148]. How shall a Bayesian treat the case when some evidence $E \triangleq x$ (e.g. Mercury’s perihelion advance) is known well-before the correct hypothesis/theory/model $H \triangleq \mu$ (Einstein’s general relativity theory) is found? How shall H be added to the Bayesian machinery a posteriori? What is the prior of H ? Should it be the belief in H in a hypothetical counterfactual world in which E is not known? Can old evidence E confirm H ? After all, H could simply be constructed/biased/fitted towards “explaining” E . Strictly speaking, a Bayesian needs to choose the hypothesis/model class before seeing the data, which seldom reflects scientific practice [5].
- f) **Other issues/problems.** Comparison to Carnap’s confirmation theory [149] and Laplace rule [150], allowing for continuous model classes, how to incorporate prior knowledge [151, 152], and others.

Solomonoff’s theory has already been intensively studied in the predictive setting [18, 50, 89, 91, 94] mostly confirming its power, with the occasional unexpected exception [103]. Important open questions are:

- g) **Prediction of selected bits.** Consider a very simple and special case of problem 5.i, a binary sequence that coincides at even times with the preceding (odd) bit, but is otherwise uncomputable. Every child will quickly realize that the even bits coincide with the preceding odd bit, and after a while perfectly predict the even bits, given the past bits. The uncomputability of the sequence is no hindrance. It is unknown whether Solomonoff works or fails in this situation. I expect that a solution of this special case will lead to general useful insights and advance this theory (cf. problem 5.i).
- h) **Identification of “natural” Turing machines.** In order to pin down the additive/multiplicative constants that plague most results in AIT, it would be highly desirable to identify a class of “natural” UTMs/USMs which have a variety of favorable properties. A more moderate approach may be to consider classes $\mathcal{C}_i$ of universal Turing machines (UTM) or universal semimeasures (USM) satisfying certain properties $\mathcal{P}_i$ and showing that the intersection $\cap_i \mathcal{C}_i$ is not empty. Indeed, very occasionally results in AIT only hold for particular (subclasses of) UTMs [153]. A grander vision is to find the single “best” UTM or USM [154] (a remarkable approach).
- i) **Martin-Löf convergence.** Quite unexpectedly, a loophole in the proof of Martin-Löf (M.L.) convergence of M to μ in the literature has been found [101]. In [102] it has been shown that this loophole cannot be fixed, since M.L.-convergence actually can fail. The construction of non-universal (semi)measures D and W that M.L. converge to μ [103] partially rescued the situation. The major problem left open is the convergence rate for $W \xrightarrow{M.L.} \mu$. The current bound for $D \xrightarrow{M.L.} \mu$ is double exponentially worse than for $M \xrightarrow{w.p.1} \mu$. It is also unknown whether convergence in ratio holds. Finally, there could still exist *universal* semimeasures M (dominating all enumerable semimeasures) for which M.L.-convergence holds. In case they exist, they probably have particularly interesting additional structure and properties.

- j) **Generalized mixtures and convergence concepts.** Another interesting and potentially fruitful approach to the above convergence problem is to consider other classes of semimeasures $\mathcal{M}$ [63, 66, 86], define mixtures ξ over $\mathcal{M}$, and (possibly) generalized randomness concepts by using this ξ to define a generalized notion of randomness. Using this approach, in [87] it has been shown that convergence holds for a subclass of Bernoulli distributions if the class is dense, but fails if the class is gappy, showing that a denseness characterization of $\mathcal{M}$ could be promising in general. See also [155, 156].
- k) **Lower convergence bounds and defect of M .** One can show that $M(\bar{x}_t | x_{<t}) \geq 2^{-K(t)}$, i.e. the probability of making a wrong prediction $\bar{x}_t$ converges to zero slower than any computable summable function. This shows that, although M converges rapidly to μ in a cumulative sense, occasionally, namely for simply describable t , the prediction quality is poor. An easy way to show the lower bound is to exploit the semimeasure defect of M . Do similar lower bounds hold for a proper (Solomonoff) normalized measure M_{norm} ? I conjecture the answer is yes, i.e. the lower bound is not a semimeasure artifact, but “real”.
- l) **Using AIXI for prediction.** Since AIXI is a unification of sequential decision theory with the idea of universal probability one may think that the AIXI model for a sequence prediction problem exactly reduces to Solomonoff’s universal sequence prediction scheme. Unfortunately this is not the case. For one reason, M is only a probability distribution on the inputs but not on the outputs. This is also one of the origins of the difficulty of proving general value bounds for AIXI. The question is whether, nevertheless, AIXI predicts sequences as well as Solomonoff’s scheme. A first weak bound in a very restricted setting has been proven in [27, Sec.6.2], showing that progress in this question is possible.

The most important open, but unfortunately likely also the hardest, problem is the formal identification of natural universal (Turing) machines (h). A proper solution would eliminate one of the two most important critiques of the whole field of AIT. Item (l) is an important question for universal AI.

5. Open Problems regarding Optimality of AIXI

AIXI has been shown to be Pareto-optimal and a variant of AIXI to be self-optimizing [137]. These are important results supporting the claim that AIXI is universally optimal. More results can be found in [27]. Unlike the induction case, the results are not strong enough to allay all doubts. Indeed, the major problem is not to prove optimality but to come up with a sufficiently strong but still satisfiable optimality notion in the reinforcement learning case. The following items list four potential approaches towards a solution:

- a) **What is meant by universal optimality?** A “learner” (like AIXI) may converge to the optimal informed decision maker (like AI_μ) in several senses. Possibly relevant concepts from statistics are, *consistency*, *self-tuningness*, *self-optimizingness*, *efficiency*, *unbiasedness*, *asymptotically or finite convergence* [109], *Pareto-optimality*, and some more defined in [27]. Some concepts are stronger than necessary, others are weaker than desirable but suitable to start with. It is necessary to investigate in more breadth which properties the AIXI model satisfies.

- b) **Limited environmental classes.** The problem of defining and proving general value bounds becomes more feasible by considering, in a first step, restricted concept classes. One could analyze AIXI for known classes (like Markov or factorizable environments) and especially for the new classes (*forgetful*, *relevant*, *asymptotically learnable*, *farsighted*, *uniform*, and (*pseudo*-)*passive*) defined in [27].
- c) **Generalization of AIXI to general Bayes mixtures.** Alternatively one can generalize AIXI to $AI\xi$, where $\xi(\cdot) = \sum_{\nu \in \mathcal{M}} w_{\nu} \nu(\cdot)$ is a general Bayes-mixture of distributions ν in some class $\mathcal{M}$ and prior w_{ν} . If $\mathcal{M}$ is the multi-set of all enumerable semi-measures, then $AI\xi$ coincides with AIXI. If $\mathcal{M}$ is the (multi)set of passive semi-computable environments, then AIXI reduces to Solomonoff's optimal predictor [18]. The key is not to prove absolute results for specific problem classes, but to prove *relative results* of the form "if there exists a policy with certain desirable properties, then $AI\xi$ also possesses these desirable properties". If there are tasks which cannot be solved by any policy, $AI\xi$ should not be blamed for failing.
- d) **Intelligence Aspects of AIXI.** Intelligence can have many faces. As argued in [27], it is plausible that AIXI possesses all or at least most properties an intelligent rational agent should possess. Some of the following properties could and should be investigated mathematically: *creativity*, *problem solving*, *pattern recognition*, *classification*, *learning*, *induction*, *deduction*, *building analogies*, *optimization*, *surviving in an environment*, *language processing*, *planning*.

Sources of inspiration can be previously proven loss bounds for Solomonoff sequence prediction generalized to unbounded horizon, optimality results from the adaptive control literature, and the asymptotic self-optimizingness results for the related $AI\xi$ model. Value bounds for AIXI are expected to be, in a sense, weaker than the loss bounds for Solomonoff induction because the problem class covered by AIXI is much larger than the class of sequence prediction problems.

In the same sense as Gittins' solution to the bandit problem and Laplace' rule for Bernoulli sequences, AIXI may simply be regarded as (Bayes-)optimal by construction. Even when accepting this "easy way out", the above questions remain significant: Theorems relating AIXI to $AI\mu$ would no longer be regarded as optimality proofs of AIXI, but just as how much harder it becomes to operate when μ is unknown, i.e. progress on the items above is simply reinterpreted.

A weaker goal than to prove optimality of AIXI is to ask for reasonable convergence properties:

- f) **Posterior convergence for unbounded horizon.** Convergence of M to μ holds somewhat surprisingly even for unbounded horizon, which is good news for AIXI. Unfortunately convergence can be slow, but I expect that convergence is "reasonably" fast for "slowly" growing horizon, which is important in AIXI. It would be useful to quantify and prove such a result.
- g) **Reinforcement learning.** Although there is no explicit learning algorithm built into the AIXI model, AIXI is a reinforcement learning system capable of receiving and exploiting rewards. The system learns by eliminating Turing machines q in the definition of M once they become inconsistent with the progressing history. This is similar to Gold-style learning [157]. For Markov environments (but not for partially observable environments) there are efficient general reinforcement

learning algorithms, like $TD(\lambda)$ and Q learning. One could compare the performance (learning speed and quality) of $AI\xi$ to e.g. $TD(\lambda)$ and Q learning, extending [52].

- h) **Posterization.** Many properties of Kolmogorov complexity, Solomonoff's prior, and reinforcement learning algorithms remain valid after "posterization". With posterization I mean replacing the total value V_{1m} , the weights w_ν , the complexity $K(\nu)$, the environment $\nu(or_{1:m}|a_{1:m})$, etc. by their "posteriors" V_{km} , $w_\nu(aor_{<k})$, $K(\nu|aor_{<k})$, $\nu(or_{k:m}|or_{<k}a_{1:m})$, etc, where k is the current cycle and m the lifespan of AIXI. Strangely enough for w_ν chosen as $2^{-K(\nu)}$ it is not true that $w_\nu(aor_{<k}) \sim 2^{-K(\nu|aor_{<k})}$. If this property were true, weak bounds as the one proven in [27, Sec.6.2] (which is too weak to be of practical importance) could be boosted to practical bounds of order 1. Hence, it is highly important to rescue the posterization property in some way. It may be valid when grouping together essentially equal distributions ν .
- i) **Relevant and non-computable environments μ .** Assume that the observations of AIXI contain irrelevant information, like noise. Irrelevance can formally be defined as being statistically independent of future observations and rewards, i.e. neither affecting rewards, nor containing information about future observations. It is easy to see that Solomonoff prediction does not decline under such noise if it is sampled from a computable distribution. This likely transfers to AIXI. More interesting is the case, where the irrelevant input is complex. If it is easily separable from the useful input it should not affect AIXI. On the other hand, even in prediction this problem is non-trivial, see problem 4.g. How robustly does AIXI deal with complex but irrelevant inputs? A model that explicitly deals with this situation has been developed in [129, 130].
- j) **Grain of truth problem [158].** Assume AIXI is used in a multi-agent setup [159] interacting with other agents. For simplicity I only discuss the case of a single other agent in a competitive setup, i.e. a two-person zero-sum game situation. We can entangle agents A and B by letting A observe B 's actions and vice versa. The rewards are provided externally by the rules of the game. The situation where A is AIXI and B is a perfect minimax player was analyzed in [27, Sec.6.3]. In multi-agent systems one is mostly interested in a symmetric setup, i.e. B is also an AIXI. Whereas both AIXIs may be able to learn the game and improve their strategies (to optimal minimax or more generally *Nash equilibrium*), this setup violates one of the basic assumptions. Since AIXI is incomputable, $AIXI(B)$ does not constitute a computable environment for $AIXI(A)$. More generally, starting with any class of environments $\mathcal{M}$, then $\mu \triangleq AI\xi_{\mathcal{M}}$ seems not to belong to class $\mathcal{M}$ for most (all?) choices of $\mathcal{M}$. Various results can no longer be applied, since $\mu \notin \mathcal{M}$ when coupling two $AI\xi$ s. Many questions arise: Are there interesting environmental classes for which $AI\xi_{\mathcal{M}} \in \mathcal{M}$ or $AI\xi_{tl_{\mathcal{M}}} \in \mathcal{M}$? Do $AIXI(A/B)$ converge to optimal minimax players? Do AIXIs perform well in general multi-agent setups?

From the optimality questions above, the first one (a) is the most important, least defined, and likely hardest one: In which sense can a rational agent in general and AIXI in particular be optimal? The multi-agent setting adds another layer of difficulty: The grain of truth problem (j) is in my opinion the most important fundamental problem in game theory and multi-agent systems. Its satisfactory solution should be worth a Nobel prize or Turing award.

6. Open Problems regarding Uniqueness of AIXI

As a unification of two optimal theories, it is plausible that AIXI is optimal in the “union” of their domains, which has been affirmed but not finally settled by the positive results derived so far. In the absence of a definite answer, one should be open to alternative models, but no convincing competitor exists to date. Most of the following items describe ideas which, if worked out, might result in alternative models:

- a) **Action with expert advice.** Expected performance bounds for predictions based on Solomonoff’s prior exist. Inspired by Solomonoff induction, a dual, currently very popular approach, is “prediction with expert advice” (PEA) [11, 160, 161]. Whereas PEA performs well in any environment, but only relative to a given set of experts, Solomonoff’s predictor competes with *any* other predictor, but only in expectation for environments with computable distribution. It seems philosophically less compromising to make assumptions on prediction strategies than on the environment, however weak. PEA has been generalized to active learning [11, 162], but the full reinforcement learning case is still open [52]. If successful, it could result in a model dual to AIXI, but I expect the answer to be negative, which on the positive side would show the distinguishedness of AIXI. Other ad-hoc approaches like [126, 163] are also unlikely to be competitive.
- b) **Actions as random variables.** There may be more than one way for the choice of the generalized M in the AIXI model. For instance, instead of defining M as in [27] one could treat the agent’s actions a also as universally distributed random variables and then conditionalize M on a .
- c) **Structure of AIXI.** The algebraic properties and the structure of AIXI has barely been investigated. It is known that the value of $\text{AI}\mu$ is a linear function in μ and the value of AIXI is a convex function in μ , but this is neither very deep nor very specific to AIXI. It should be possible to extract all essentials from AIXI which finally should lead to an axiomatic characterization of AIXI. The benefit is as in any *axiomatic approach*: It would clearly exhibit the assumptions, separate the essentials from technicalities, simplify understanding and, most importantly, guide in finding proofs.
- d) **Parameter dependence.** The AIXI model depends on a few parameters: the choice of observation and action spaces $\mathcal{O}$ and $\mathcal{A}$, the horizon m , and the universal machine U . So strictly speaking, AIXI is only (essentially) unique, if it is (essentially) independent of the parameters. I expect this to be true, but it has not been proven yet. The U -dependence has been discussed in problem 4.h. Countably infinite $\mathcal{O}$ and $\mathcal{A}$ would provide a rich enough interface for all problems, but even binary $\mathcal{O}$ and $\mathcal{A}$ are sufficient by sequentializing complex observations and actions. For special classes one could choose $m \rightarrow \infty$ [20]; unfortunately, the universal environment M does not belong to any of these special classes. See [27, 32, 164] for some preliminary considerations.

7. Open Problems in Defining Intelligence

A fundamental and long standing difficulty in the field of artificial intelligence is that (generic) intelligence itself is not well defined. It is an anomaly that nowadays most AI researchers avoid discussing

intelligence, which is caused by several factors: It is a difficult old subject, it is politically charged, it is not necessary for narrow AI which focusses on specific applications, AI research is done mainly by computer scientists who mainly care about algorithms rather than philosophical foundations, and the popular belief that general intelligence is principally unamenable to a mathematical definition. These reasons explain but only partially justify the low effort in trying to define intelligence.

Assume we had a definition, ideally a formal, objective, non-anthropocentric, and direct method of measuring intelligence, or at least a very general intelligence-like performance measure that could serve as an adequate substitute. This would bring the higher goals of the field into tight focus and allow us to objectively compare different approaches and judge the overall progress. Indeed, formalizing and rigorously defining a previously vague concept usually constitutes a quantum leap forward in the field: Cf. set theory, logical reasoning, infinitesimal calculus, energy, temperature, etc. Of course there is (some) work on defining [165] and testing [166] intelligence (see [32] for a comprehensive list of references):

The famous Turing test [167–169] involves human interaction, so is unfortunately informal and anthropocentric, others are large “messy” collections of existing intelligence tests [170, 171] (“shotgun” approaches), which are subjective and lack a clear theoretical grounding, and are potentially too narrow.

There are some more elegant solutions based on classical [172] and algorithmic [173] information theory (“C-Test” [174–176]), the latter closely related to Solomonoff’s [37] “perfect” inductive inference model. The simple program in [177] reached good IQ scores on some of the more mathematical tests.

One limitation of the C-Test however is that it only deals with compression and (passive) sequence prediction, while humans or machines face reactive environments where they are able to change the state of the environment through their actions. AIXI generalizes Solomonoff to reactive environments, which suggested an extremely *general, objective, fundamental, and formal* performance measure [56, 178]. This so-called Intelligence Order Relation (IOR) [27] even attracted the popular scientific press [179, 180], but the theory surrounding it has not yet been adequately explored. Here I only describe three non-technical open problems in defining intelligence.

- a) **General and specific performance measures.** Currently it is only partially understood how the IOR theoretically compares to the myriad of other tests of intelligence such as conventional IQ tests or even other performance tests proposed by AI other researchers. Another open question is whether the IOR might in some sense be too general. One may narrow the IOR to specific classes of problems [139] and compare how the resulting IOR measures compare to standard performance measures for each problem class. This could shed light on aspects of the IOR and possibly also establish connections between seemingly unrelated performance metrics for different classes of problems.
- b) **Practical performance measures.** A more practically orientated line of investigation would be to produce a resource bounded version of the IOR like the one in [27, Sec.7], or perhaps some of its special cases. This would allow one to define a practically implementable performance test, similar to the way in which the C-Test has been derived from incomputable definitions of compression using Kt complexity [175]. As there are many subtle kinds of resource bounded complexity [41], the advantages and disadvantages of each in this context would need to be carefully examined.

Another possibility is the recent Speed Prior [63] or variants of this approach.

- c) **Experimental evaluation.** Once a computable version of the IOR had been defined, one could write a computer program that implements it. One could then experimentally explore its characteristics in a range of different problem spaces. For example, it might be possible to find correlations with IQ test scores when applied to humans, like has been done with the C-Test [174]. Another possibility would be to consider more limited domains like classification problems or sequence prediction problems and to see whether the relative performance of algorithms according to the IOR agrees with standard performance measures and real world performance.

A comprehensive collection, discussion and comparison of verbal and formal intelligence tests, definitions, and measures can be found in [32].

8. Conclusions

The flavor of the open questions. While most of the key questions about universal sequence prediction have been solved, many key questions about universal AI remain open to date. The questions in Sections 4.-7. are centered around the AIT approach to induction and AI, but many require interdisciplinary working. A more detailed account with technical details can be found in the book [27] and paper [31]. Most questions are amenable to a *rigorous mathematical* treatment, including the more philosophically or vaguely sounding ones. Progress on the latter can be achieved in the usual way by cycling through (i) craft or improve mathematical definitions that resemble the intuitive concepts to be studied (e.g. “natural”, “generalization”, “optimal”), (ii) formulate or adapt a mathematical conjecture resembling the informal question, (iii) (dis)prove the conjecture. Some questions are about approximating, implementing, and testing various ideas and concepts. Technically, many questions are on (the *interface* between) and exploit *techniques* used in (*algorithmic*) *information theory*, *machine learning*, *Bayesian statistics*, (*adaptive*) *control theory*, and *reinforcement learning*.

Feasibility, difficulty, and interestingness of the open questions. I concentrated on questions whose answers probably help to develop the foundations of universal induction and UAI. Some problems are very hard, and their satisfactory solution worth a Nobel prize or Turing award, e.g. problem 5.j. I included those questions that looked promising and interesting at the time of writing this article. In the following I try to estimate their relative feasibility, difficulty, and interestingness:

- Problems roughly sorted from most important or interesting to least:
5.j,4.h,7.b,5.a,7.c,4.b,4.f,4.l,5.d,5.f,5.i,4.a,4.c,4.d,4.e,4.g,5.c,6.a,6.b,5.g,5.h,4.j,5.b,6.c,7.a,4.i,4.k,6.d.
- Problems roughly sorted from most to least time consuming:
4.b,4.h,5.d,5.j,7.c,4.c,5.b,5.c,6.c,6.a,5.g,5.i,7.b,4.i,4.j,4.l,5.a,6.d,6.b,5.f,5.h,4.d,4.f,4.g,4.k,7.a,4.a,4.e.
- Problems roughly sorted from hard to easy:
4.h,6.c,5.j,4.b,4.i,4.l,5.b,6.a,7.b,4.c,4.g,5.a,5.c,6.b,5.f,5.i,4.j,5.h,7.a,4.d,4.f,4.k,5.d,6.d,7.c,4.a,4.e,5.g.

These rankings hopefully do not mislead but give the interested reader some guidance where (not) to start. The final paragraphs of this article are devoted to the role UAI plays in the grand goal of AI.

Other approaches to AI. There are many fields that try to understand the phenomenon of intelligence and whose insights help in creating intelligent systems: *Cognitive psychology* and *behaviorism* [181], *philosophy of mind* [182, 183], *neuroscience* [184], *linguistics* [185, 186], *anthropology* [187], *machine learning* [59, 188], *logic* [189, 190], *computer science* [25, 191], *biological evolution* [192], and others. In computer science, most AI research is *bottom-up*; extending and improving existing or developing new *algorithms* and increasing their range of applicability; an interplay between experimentation on toy problems and theory, with occasional real-world applications. The agent perspective of AI [25] brings some order and unification in the large variety of problems the fields wants to address, but it is only a framework rather than a complete theory. In the absence of a perfect (stochastic) model of the environment, machine learning techniques are needed and employed. Apart from AIXI, there is no general theory for learning agents. This resulted in an ever increasing number of *limited models and algorithms* in the past.

The information-theoretic approach to AI. Solomonoff induction and AIXI are mathematical *top-down* approaches. The price for this generality is that the full models are computationally intractable, and investigations have to be mostly theoretical at this stage. From a different perspective, UAI strictly *separates the conceptual and algorithmic AI questions*. Two analogies may help: Von Neumann's optimal minimax strategy [58] is a conceptual solution of zero-sum games, but is infeasible for most interesting zero-sum games. Nevertheless most algorithms are based on approximations of this ideal. In physics, the quest for a "theory of everything" (TOE) lead to extremely successful unified theories, despite their computational intractability [23, 24]. The role of UAI in AI should be understood as analogous to the role of minimax in zero-sum games or of the TOE in physics.

Epilogue. As we have seen, algorithmic information theory offers answers to the following two *key scientific questions*: (1) The problem of *induction*, which is what science itself is mostly about: Induction $\approx$ finding regularities in data $\approx$ understanding the world $\approx$ science. (2) Understanding *intelligence*, the key property that distinguishes humans from animals and inanimate things.

This modern *mathematical* approach to both questions (1) and (2) is quite different to the more traditional philosophical, logic-based, engineering, psychological, or neurological approaches. Among the few other mathematical approaches, none captures rational intelligence as completely as the AIXI model does. Still, a lot of questions remain open. Raising and discussing them was the primary focus of this article.

Imagine a *complete practical* solution of the AI problem (by the next generation or so), i.e. systems that surpass human intelligence. This would *transform society* more than the industrial revolution two centuries ago, the computer last century, and the internet this century. Although individually, some questions I raised seem quite technical and narrow, they derive their significance from their role in a truly outstanding scientific endeavor. As with most innovations, the social benefit of course depends on its benevolent use.

References and Notes

1. Hume, D. *A Treatise of Human Nature, Book I*. [Edited version by L. A. Selby-Bigge and P. H. Nidditch, Oxford University Press, 1978; 1739.

2. Popper, K.R. *Logik der Forschung*; Springer: Berlin, Germany, 1934. [English translation: *The Logic of Scientific Discovery* Basic Books, New York, NY, USA, 1959, and Hutchinson, London, UK, revised edition, 1968.
3. Howson, C. *Hume's Problem: Induction and the Justification of Belief*, 2nd Ed.; Oxford University Press: Oxford, UK, 2003.
4. Levi, I. *Gambling with Truth: An Essay on Induction and the Aims of Science*; MIT Press: Cambridge, MA, USA, 1974.
5. Earman, J. *Bayes or Bust? A Critical Examination of Bayesian Confirmation Theory*; MIT Press: Cambridge, MA, USA, 1993.
6. Wallace, C.S. *Statistical and Inductive Inference by Minimum Message Length*; Springer: Berlin, Germany, 2005.
7. Salmon, W.C. *Four Decades of Scientific Explanation*; University of Pittsburgh Press: Pittsburgh, PA, USA, 2006.
8. Frigg, R.; Hartmann, S. Models in science. *Stanford Encyclopedia of Philosophy* **2006**. <http://plato.stanford.edu/entries/models-science/>.
9. Wikipedia. Predictive modelling, 2008.
10. Brockwell, P.J.; Davis, R.A. *Introduction to Time Series and Forecasting*, 2nd Ed.; Springer: New York, NY, USA, 2002.
11. Cesa-Bianchi, N.; Lugosi, G. *Prediction, Learning, and Games*; Cambridge University Press: Cambridge, UK, 2006.
12. Geisser, S. *Predictive Inference*; Chapman & Hall/CRC: New York, NY, USA, 1993.
13. Chatfield, C. *The Analysis of Time Series: An Introduction*, 6th Ed.; Chapman & Hall / CRC: New York, NY, USA, 2003.
14. Ferguson, T.S. *Mathematical Statistics: A Decision Theoretic Approach*, 3rd Ed.; Academic Press: New York, NY, USA, 1967.
15. DeGroot, M.H. *Optimal Statistical Decisions*; McGraw-Hill: New York, NY, USA, 1970.
16. Jeffrey, R.C. *The Logic of Decision*, 2nd Ed.; University of Chicago Press: Chicago, IL, USA, 1983.
17. Paris, J.B. *The Uncertain Reasoner's Companion: A Mathematical Perspective*; Cambridge University Press: Cambridge, UK, 1995.
18. Hutter, M. Optimality of universal Bayesian prediction for general loss and alphabet. *Journal of Machine Learning Research* **2003**, *4*, 971–1000.
19. Hutter, M. Universal algorithmic intelligence: A mathematical top→down approach, In *Artificial General Intelligence*; Springer: Berlin, Germany, 2007; pp. 227–290.
20. Bertsekas, D.P. *Dynamic Programming and Optimal Control, volume I and II*, 3rd Ed.; Athena Scientific: Belmont, MA, USA, 2006; Volumes 1 and 2.
21. Kemp, S. Toward a monistic theory of science: The 'strong programme' reconsidered. *Philosophy of the Social Sciences* **2003**, *33*, 311–338.
22. Kellert, S.H.; Longino, H.E.; Waters, C.K., Eds. *Scientific Pluralism*; Univ. of Minnesota Press: Minneapolis, MN, USA, 2006.
23. Green, M.B.; Schwarz, J.H.; Witten, E. *Superstring Theory: Volumes I and 2*; Cambridge Univer-

- sity Press: Cambridge, UK, 2000.
24. Greene, B. *The Elegant Universe: Superstrings, Hidden Dimensions, and the Quest for the Ultimate Theory*; Vintage Press: London, UK, 2000.
 25. Russell, S.J.; Norvig, P. *Artificial Intelligence. A Modern Approach*, 2nd Ed.; Prentice-Hall: Englewood Cliffs, NJ, USA, 2003.
 26. Hutter, M. A theory of universal artificial intelligence based on algorithmic complexity. Technical Report cs.AI/0004001, München, Germany, 62 pages, 2000. <http://arxiv.org/abs/cs.AI/0004001>.
 27. Hutter, M. *Universal Artificial Intelligence: Sequential Decisions based on Algorithmic Probability*; Springer: Berlin, Germany, 2005. 300 pages, <http://www.hutter1.net/ai/uaibook.htm>.
 28. Oates, T.; Chong, W. Book review: Marcus Hutter, universal artificial intelligence, Springer (2004). *Artificial Intelligence* **2006**, 170, 1222–1226.
 29. Solomonoff, R.J. A preliminary report on a general theory of inductive inference. Technical Report V-131, Zator Co.: Cambridge, MA, USA, 1960. Distributed at the *Conference on Cerebral Systems and Computers*, 8–11 Feb. 1960.
 30. Bellman, R.E. *Dynamic Programming*; Princeton University Press: Princeton, NJ, USA, 1957.
 31. Hutter, M. On universal prediction and Bayesian confirmation. *Theoretical Computer Science* **2007**, 384, 33–48.
 32. Legg, S.; Hutter, M. Universal intelligence: A definition of machine intelligence. *Minds & Machines* **2007**, 17, 391–444.
 33. Franklin, J. *The Science of Conjecture: Evidence and Probability before Pascal*; Johns Hopkins University Press: Baltimore, MD, USA, 2002.
 34. Asmis, E. *Epicurus' Scientific Method*; Cornell Univ. Press: Ithaca, NY, USA, 1984.
 35. Turing, A.M. On computable numbers, with an application to the Entscheidungsproblem. *Proc. London Mathematical Society* **1937**, 2, 230–265.
 36. Bayes, T. An essay towards solving a problem in the doctrine of chances. *Philosophical Transactions of the Royal Society* **1763**, 53, 376–398. [Reprinted in *Biometrika*, 45, 296–315, 1958].
 37. Solomonoff, R.J. A formal theory of inductive inference: Parts 1 and 2. *Information and Control* **1964**, 7, 1–22 and 224–254.
 38. Kolmogorov, A.N. Three approaches to the quantitative definition of information. *Problems of Information and Transmission* **1965**, 1, 1–7.
 39. Berger, J. *Statistical Decision Theory and Bayesian Analysis*, 3rd Ed.; Springer: Berlin, Germany, 1993.
 40. Hutter, M. Algorithmic information theory: a brief non-technical guide to the field. *Scholarpedia* **2007**, 2, 2519.
 41. Li, M.; Vitányi, P.M.B. *An Introduction to Kolmogorov Complexity and its Applications*, 3rd Ed.; Springer: Berlin, Germany, 2008.
 42. Hutter, M. Algorithmic complexity. *Scholarpedia* **2008**, 3, 2573.
 43. MacKay, D.J.C. *Information theory, inference and learning algorithms*; Cambridge University Press: Cambridge, MA, USA, 2003.
 44. Cover, T.M.; Thomas, J.A. *Elements of Information Theory*, 2nd Ed.; Wiley-Interscience: New York, NY, USA, 2006.

45. Lempel, A.; Ziv, J. On the complexity of finite sequences. *IEEE Transactions on Information Theory* **1976**, *22*, 75–81.
46. Cilibrasi, R.; Vitányi, P.M.B. Clustering by compression. *IEEE Trans. Information Theory* **2005**, *51*, 1523–1545.
47. Willems, F.M.J.; Shtarkov, Y.M.; Tjalkens, T.J. Reflections on the prize paper: The context-tree weighting method: Basic properties. *IEEE Information Theory Society Newsletter* **1997**, pp. 20–27.
48. Hutter, M.; Legg, S.; Vitányi, P.M.B. Algorithmic probability. *Scholarpedia* **2007**, *2*, 2572.
49. Zvonkin, A.K.; Levin, L.A. The complexity of finite objects and the development of the concepts of information and randomness by means of the theory of algorithms. *Russian Mathematical Surveys* **1970**, *25*, 83–124.
50. Solomonoff, R.J. Complexity-based induction systems: Comparisons and convergence theorems. *IEEE Transactions on Information Theory* **1978**, *IT-24*, 422–432.
51. Li, M.; Vitányi, P.M.B. Applications of algorithmic information theory. *Scholarpedia* **2007**, *2*, 2658.
52. Poland, J.; Hutter, M. Universal learning of repeated matrix games. In *Proc. 15th Annual Machine Learning Conf. of Belgium and The Netherlands (Benelearn'06)*. Ghent, Belgium, 2006; pp. 7–14.
53. Pankov, S. A computational approximation to the AIXI model. In *Proc. 1st Conference on Artificial General Intelligence*, 2008; Vol. 171, pp. 256–267.
54. Hutter, M. Universal sequential decisions in unknown environments. In *Proc. 5th European Workshop on Reinforcement Learning (EWRL-5)*. Onderwijsinstituut CKI, Utrecht Univ., Netherlands, 2001; Vol. 27, pp. 25–26.
55. Hutter, M. Towards a universal theory of artificial intelligence based on algorithmic probability and sequential decisions. In *Proc. 12th European Conf. on Machine Learning (ECML'01)*. Freiburg, Germany, 2001; Springer: Berlin, Germany; Vol. 2167, *LNAI*, pp. 226–238.
56. Legg, S. *Machine Super Intelligence*. PhD thesis, IDSIA, Lugano, Switzerland, 2008.
57. Chaitin, G.J. On the length of programs for computing finite binary sequences. *Journal of the ACM* **1966**, *13*, 547–569.
58. Neumann, J.V.; Morgenstern, O. *Theory of Games and Economic Behavior*; Princeton University Press: Princeton, NJ, USA, 1944.
59. Sutton, R.S.; Barto, A.G. *Reinforcement Learning: An Introduction*; MIT Press: Cambridge, MA, USA, 1998.
60. Martin-Löf, P. The definition of random sequences. *Information and Control* **1966**, *9*, 602–619.
61. Levin, L.A. Randomness conservation inequalities: Information and independence in mathematical theories. *Information and Control* **1984**, *61*, 15–37.
62. Levin, L.A. Universal sequential search problems. *Problems of Information Transmission* **1973**, *9*, 265–266.
63. Schmidhuber, J. The speed prior: A new simplicity measure yielding near-optimal computable predictions. In *Proc. 15th Conf. on Computational Learning Theory (COLT'02)*. Sydney, Australia, 2002; Springer: Berlin, Germany; Vol. 2375, *LNAI*, pp. 216–228.
64. Chaitin, G.J. *Algorithmic Information Theory*; Cambridge University Press: Cambridge, UK,

- 1987.
65. Chaitin, G.J. *The Limits of Mathematics: A Course on Information Theory and the Limits of Formal Reasoning*; Springer: Berlin, Germany, 2003.
 66. Schmidhuber, J. Hierarchies of generalized Kolmogorov complexities and nonenumerable universal measures computable in the limit. *International Journal of Foundations of Computer Science* **2002**, *13*, 587–612.
 67. Gács, P.; Tromp, J.; Vitányi, P.M.B. Algorithmic statistics. *IEEE Transactions on Information Theory* **2001**, *47*, 2443–2463.
 68. Vereshchagin, N.; Vitányi, P.M.B. Kolmogorov's structure functions with an application to the foundations of model selection. In *Proc. 43rd Symposium on Foundations of Computer Science*. Vancouver, Canada, 2002; pp. 751–760.
 69. Vitányi, P.M.B. Meaningful information. *Proc. 13th International Symposium on Algorithms and Computation (ISAAC'02)* **2002**, *2518*, 588–599.
 70. Wallace, C.S.; Boulton, D.M. An information measure for classification. *Computer Journal* **1968**, *11*, 185–194.
 71. Rissanen, J.J. Modeling by shortest data description. *Automatica* **1978**, *14*, 465–471.
 72. Rissanen, J.J. *Stochastic Complexity in Statistical Inquiry*; World Scientific: Singapore, Singapore, 1989.
 73. Quinlan, J.R.; Rivest, R.L. Inferring decision trees using the minimum description length principle. *Information and Computation* **1989**, *80*, 227–248.
 74. Gao, Q.; Li, M. The minimum description length principle and its application to online learning of handprinted characters. In *Proc. 11th International Joint Conf. on Artificial Intelligence*. Detroit, MI, USA, 1989; pp. 843–848.
 75. Milosavljević, A.; Jurka, J. Discovery by minimal length encoding: A case study in molecular evolution. *Machine Learning* **1993**, *12*, 96–87.
 76. Pednault, E.P.D. Some experiments in applying inductive inference principles to surface reconstruction. In *Proc. 11th International Joint Conf. on Artificial Intelligence*. San Mateo, CA, USA, 1989; Morgan Kaufmann: San Francisco, CA, USA; pp. 1603–1609.
 77. Grünwald, P.D. *The Minimum Description Length Principle*; The MIT Press: Cambridge, MA, USA 2007.
 78. Cilibrasi, R.; Vitányi, P.M.B. Similarity of objects and the meaning of words. In *Proc. 3rd Annual Conferene on Theory and Applications of Models of Computation (TAMC'06)*. Beijing, China, 2006; Springer: Berlin, Germany; Vol. 3959, LNCS, pp. 21–45.
 79. Schmidhuber, J. Discovering neural nets with low Kolmogorov complexity and high generalization capability. *Neural Networks* **1997**, *10*, 857–873.
 80. Schmidhuber, J.; Zhao, J.; Wiering, M.A. Shifting inductive bias with success-story algorithm, adaptive Levin search, and incremental self-improvement. *Machine Learning* **1997**, *28*, 105–130.
 81. Schmidhuber, J. Optimal ordered problem solver. *Machine Learning* **2004**, *54*, 211–254.
 82. Schmidhuber, J. Low-complexity art. *Leonardo, Journal of the International Society for the Arts, Sciences, and Technology* **1997**, *30*, 97–103.
 83. Calude, C.S. *Information and Randomness: An Algorithmic Perspective*, 2nd Ed.; Springer: Berlin,

Germany, 2002.

84. Hutter, M. The fastest and shortest algorithm for all well-defined problems. *International Journal of Foundations of Computer Science* **2002**, *13*, 431–443.
85. Stork, D. Foundations of Occam's razor and parsimony in learning. *NIPS 2001 Workshop* **2001**. <http://www.rii.ricoh.com/~stork/OccamWorkshop.html>.
86. Hutter, M. On the existence and convergence of computable universal priors. In *Proc. 14th International Conf. on Algorithmic Learning Theory (ALT'03)*. Sapporo, Japan, 2003; Springer: Berlin, Germany; Vol. 2842, *LNAI*, pp. 298–312.
87. Hutter, M. On generalized computable universal priors and their convergence. *Theoretical Computer Science* **2006**, *364*, 27–41.
88. Hutter, M. Convergence and error bounds for universal prediction of nonbinary sequences. In *Proc. 12th European Conf. on Machine Learning (ECML'01)*. Freiburg, Germany, 2001; Springer, Berlin, Germany, 2001; Vol. 2167, *LNAI*, pp. 239–250.
89. Hutter, M. New error bounds for Solomonoff prediction. *Journal of Computer and System Sciences* **2001**, *62*, 653–667.
90. Hutter, M. General loss bounds for universal sequence prediction. In *Proc. 18th International Conf. on Machine Learning (ICML'01)*. Williams College, Williamstown, MA, USA, 2001; Morgan Kaufmann Publishers Inc.: San Francisco, CA, USA; pp. 210–217.
91. Hutter, M. Convergence and loss bounds for Bayesian sequence prediction. *IEEE Transactions on Information Theory* **2003**, *49*, 2061–2067.
92. Hutter, M. Online prediction – Bayes versus experts. Technical report, <http://www.hutter1.net/ai/bayespea.htm>, 2004. Presented at the *EU PASCAL Workshop on Learning Theoretic and Bayesian Inductive Principles (LTBIP'04)*.
93. Chernov, A.; Hutter, M. Monotone conditional complexity bounds on future prediction errors. In *Proc. 16th International Conf. on Algorithmic Learning Theory (ALT'05)*. Singapore, 2005; Springer: Berlin, Germany; Vol. 3734, *LNAI*, pp. 414–428.
94. Chernov, A.; Hutter, M.; Schmidhuber, J. Algorithmic complexity bounds on future prediction errors. *Information and Computation* **2007**, *205*, 242–261.
95. Hutter, M. Sequence prediction based on monotone complexity. In *Proc. 16th Annual Conf. on Learning Theory (COLT'03)*. Washington, DC, USA, 2003; Springer: Berlin, Germany; Vol. 2777, *LNAI*, pp. 506–521.
96. Hutter, M. Sequential predictions based on algorithmic complexity. *Journal of Computer and System Sciences* **2006**, *72*, 95–117.
97. Poland, J.; Hutter, M. Convergence of discrete MDL for sequential prediction. In *Proc. 17th Annual Conf. on Learning Theory (COLT'04)*. Banff, Canada, 2004; Springer: Berlin, Germany; Vol. 3120, *LNAI*, pp. 300–314.
98. Poland, J.; Hutter, M. Asymptotics of discrete MDL for online prediction. *IEEE Transactions on Information Theory* **2005**, *51*, 3780–3795.
99. Poland, J.; Hutter, M. On the convergence speed of MDL predictions for Bernoulli sequences. In *Proc. 15th International Conf. on Algorithmic Learning Theory (ALT'04)*. Padova, Italy, 2004; Springer: Berlin, Germany; Vol. 3244, *LNAI*, pp. 294–308.

100. Poland, J.; Hutter, M. MDL convergence speed for Bernoulli sequences. *Statistics and Computing* **2006**, *16*, 161–175.
101. Hutter, M. An open problem regarding the convergence of universal a priori probability. In *Proc. 16th Annual Conf. on Learning Theory (COLT'03)*. Washington, DC, USA, 2003; Springer: Berlin, Germany; Vol. 2777, *LNAI*, pp. 738–740.
102. Hutter, M.; Muchnik, A.A. Universal convergence of semimeasures on individual random sequences. In *Proc. 15th International Conf. on Algorithmic Learning Theory (ALT'04)*. Padova, Italy, 2004; Springer, Berlin, Germany; Vol. 3244, *LNAI*, pp. 234–248.
103. Hutter, M.; Muchnik, A.A. On semimeasures predicting Martin-Löf random sequences. *Theoretical Computer Science* **2007**, *382*, 247–261.
104. Hutter, M. On the foundations of universal sequence prediction. In *Proc. 3rd Annual Conference on Theory and Applications of Models of Computation (TAMC'06)*. Beijing, China, 2006; Springer: Berlin, Germany; Vol. 3959, *LNCS*, pp. 408–420.
105. Michie, D. Game-playing and game-learning automata, In *Advances in Programming and Non-Numerical Computation*; Pergamon: New York, NY, USA, 1966; pp. 183–200.
106. Berry, D.A.; Fristedt, B. *Bandit Problems: Sequential Allocation of Experiments*; Chapman and Hall: London, UK, 1985.
107. Duff, M. *Optimal Learning: Computational procedures for Bayes-adaptive Markov decision processes*. PhD thesis, Department of Computer Science, University of Massachusetts Amherst, MA, USA, 2002.
108. Szita, I.; Lörincz, A. The many faces of optimism: a unifying approach. In *Proc. 12th International Conference (ICML 2008)*. Helsinki, Finland, 2008; Vol. 307.
109. Kumar, P.R.; Varaiya, P.P. *Stochastic Systems: Estimation, Identification, and Adaptive Control*; Prentice Hall: Englewood Cliffs, NJ, USA, 1986.
110. Agrawal, R.; Teneketzis, D.; Anantharam, V. Asymptotically efficient adaptive allocation schemes for controlled i.i.d. processes: Finite parameter space. *IEEE Trans. Automatic Control* **1989**, *34*, 258–266.
111. Agrawal, R.; Teneketzis, D.; Anantharam, V. Asymptotically efficient adaptive allocation schemes for controlled Markov chains: Finite parameter space. *IEEE Trans. Automatic Control* **1989**, *34*, 1249–1259.
112. Samuel, A.L. Some studies in machine learning using the game of checkers. *IBM Journal on Research and Development* **1959**, *3*, 210–229.
113. Barto, A.G.; Sutton, R.S.; Anderson, C.W. Neuronlike adaptive elements that can solve difficult learning control problems. *IEEE Transactions on Systems, Man, and Cybernetics* **1983**, *SMC-13*, 834–846.
114. Sutton, R.S. Learning to predict by the methods of temporal differences. *Machine Learning* **1988**, *3*, 9–44.
115. Watkins, C. *Learning from Delayed Rewards*. PhD thesis, King's College, Oxford, UK, 1989.
116. Watkins, C.; Dayan, P. Q-learning. *Machine Learning* **1992**, *8*, 279–292.
117. Moore, A.W.; Atkeson, C.G. Prioritized sweeping: Reinforcement learning with less data and less time. *Machine Learning* **1993**, *13*, 103–130.

118. Tesauro, G. “TD”-Gammon, a self-teaching backgammon program, achieves master-level play. *Neural Computation* **1994**, *6*, 215–219.
119. Wiering, M.A.; Schmidhuber, J. Fast online “Q”(λ). *Machine Learning* **1998**, *33*, 105–116.
120. Kearns, M.; Koller, D. Efficient reinforcement learning in factored MDPs. In *Proc. 16th International Joint Conference on Artificial Intelligence (IJCAI-99)*. Stockholm, Sweden, 1999; Morgan Kaufmann, San Francisco, CA, USA; pp. 740–747.
121. Wiering, M.A.; Salustowicz, R.P.; Schmidhuber, J. Reinforcement learning soccer teams with incomplete world models. *Artificial Neural Networks for Robot Learning. Special issue of Autonomous Robots* **1999**, *7*, 77–88.
122. Baum, E.B. Toward a model of intelligence as an economy of agents. *Machine Learning* **1999**, *35*, 155–185.
123. Koller, D.; Parr, R. Policy iteration for factored MDPs. In *Proc. 16th Conference on Uncertainty in Artificial Intelligence (UAI-00)*. Stanford University, Stanford, CA, USA, 2000; Morgan Kaufmann, San Francisco, CA, USA; pp. 326–334.
124. Singh, S.; Littman, M.; Jong, N.; Pardoe, D.; Stone, P. Learning predictive state representations. In *Proc. 20th International Conference on Machine Learning (ICML’03)*, Washington, DC, USA, 2003; pp. 712–719.
125. Guestrin, C.; Koller, D.; Parr, R.; Venkataraman, S. Efficient solution algorithms for factored MDPs. *Journal of Artificial Intelligence Research (JAIR)* **2003**, *19*, 399–468.
126. Ryabko, D.; Hutter, M. On the possibility of learning in reactive environments with arbitrary dependence. *Theoretical Computer Science* **2008**, *405*, 274–284.
127. Strehl, A.L.; Diuk, C.; Littman, M.L. Efficient structure learning in factored-state MDPs. In *Proc. 27th AAAI Conference on Artificial Intelligence*. Vancouver, BC, Canada, 2007; AAAI Press, pp. 645–650.
128. Ross, S.; Pineau, J.; Paquet, S.; Chaib-draa, B. Online planning algorithms for POMDPs. *Journal of Artificial Intelligence Research* **2008**, *2008*, 663–704.
129. Hutter, M. Feature Markov decision processes. In *Proc. 2nd Conf. on Artificial General Intelligence (AGI’09)*. Arlington, VA, USA, 2009; Atlantis Press, Vol. 8, pp. 61–66.
130. Hutter, M. Feature dynamic Bayesian networks. In *Proc. 2nd Conf. on Artificial General Intelligence (AGI’09)*. Arlington, VA, USA, 2009; Atlantis Press, Vol. 8, pp. 67–73.
131. Kaelbling, L.P.; Littman, M.L.; Moore, A.W. Reinforcement learning: a survey. *Journal of Artificial Intelligence Research* **1996**, *4*, 237–285.
132. Kaelbling, L.P.; Littman, M.L.; Cassandra, A.R. Planning and acting in partially observable stochastic domains. *Artificial Intelligence* **1998**, *101*, 99–134.
133. Boutilier, C.; Dean, T.; Hanks, S. Decision-theoretic planning: Structural assumptions and computational leverage. *Journal of Artificial Intelligence Research* **1999**, *11*, 1–94.
134. Ng, A.Y.; Coates, A.; Diel, M.; Ganapathi, V.; Schulte, J.; Tse, B.; Berger, E.; Liang, E. Autonomous inverted helicopter flight via reinforcement learning. In *ISER*. Springer: New York, NY, USA, 2004; Vol. 21, *Springer Tracts in Advanced Robotics*, pp. 363–372.
135. Bertsekas, D.P.; Tsitsiklis, J.N. *Neuro-Dynamic Programming*; Athena Scientific: Belmont, MA, USA, 1996.

136. Hutter, M. Bayes optimal agents in general environments. Technical report, Istituto Dalle Molle di Studi sull'Intelligenza Artificiale (IDSIA), 2004. unpublished manuscript.
137. Hutter, M. Self-optimizing and Pareto-optimal policies in general environments based on Bayes-mixtures. In *Proc. 15th Annual Conf. on Computational Learning Theory (COLT'02)*, Sydney, Australia, 2002; Springer: Berlin, Germany; Vol. 2375, *LNAI*, pp. 364–379.
138. Legg, S.; Hutter, M. Ergodic MDPs admit self-optimising policies. Technical Report IDSIA-21-04, IDSIA, 2004.
139. Legg, S.; Hutter, M. A taxonomy for abstract environments. Technical Report IDSIA-20-04, IDSIA, 2004.
140. Gaglio, M. Universal search. *Scholarpedia* **2007**, 2, 2575.
141. Schmidhuber, J. Gödel machines: Self-referential universal problem solvers making provably optimal self-improvements, In *Artificial General Intelligence*; Springer, in press, 2005.
142. Jaynes, E.T. *Probability Theory: The Logic of Science*; Cambridge University Press: Cambridge, MA, USA, 2003.
143. Maher, P. Probability captures the logic of scientific confirmation, In *Contemporary Debates in Philosophy of Science*; Hitchcock, C., Ed.; Blackwell Publishing: Malden, MA, USA, 2004; chapter 3, pp. 69–93.
144. Rescher, N. *Paradoxes: Their Roots, Range, and Resolution*; Open Court: Lanham, MD, USA, 2001.
145. Goodman, N. *Fact, Fiction, and Forecast*, 4th Ed.; Harvard University Press: Cambridge, MA, USA, 1983.
146. Kass, R.E.; Wasserman, L. The selection of prior distributions by formal rules. *Journal of the American Statistical Association* **1996**, 91, 1343–1370.
147. Jeffreys, H. An invariant form for the prior probability in estimation problems. In *Proc. Royal Society London*, 1946, Vol. Series A 186, pp. 453–461.
148. Glymour, C. *Theory and Evidence*; Princeton Univ. Press: Princeton, NJ, USA, 1980.
149. Carnap, R. *The Continuum of Inductive Methods*; University of Chicago Press: Chicago, IL, USA, 1952.
150. Laplace, P. *Théorie analytique des probabilités*; Courcier, Paris, France, 1812. [English translation by Truscott, F.W. and Emory, F.L.: *A Philosophical Essay on Probabilities*. Dover, 1952].
151. Press, S.J. *Subjective and Objective Bayesian Statistics: Principles, Models, and Applications*, 2nd Ed.; Wiley: New York, NY, USA, 2002.
152. Goldstein, M. Subjective bayesian analysis: Principles and practice. *Bayesian Analysis* **2006**, 1, 403–420.
153. Muchnik, A.A.; Positselsky, S.Y. Kolmogorov entropy in the context of computability theory. *Theoretical Computer Science* **2002**, 271, 15–35.
154. Müller, M. Stationary algorithmic probability. Technical Report <http://arXiv.org/abs/cs/0608095>, TU Berlin, Berlin, 2006.
155. Ryabko, D.; Hutter, M. On sequence prediction for arbitrary measures. In *Proc. IEEE International Symposium on Information Theory (ISIT'07)*. IEEE, Nice, France, 2007, pp. 2346–2350.
156. Ryabko, D.; Hutter, M. Predicting non-stationary processes. *Applied Mathematics Letters* **2008**,

- 21, 477–482.
157. Gold, E.M. Language identification in the limit. *Information and Control* **1967**, *10*, 447–474.
158. Kalai, E.; Lehrer, E. Rational learning leads to Nash equilibrium. *Econometrica* **1993**, *61*, 1019–1045.
159. Weiss, G., Ed. *Multiagent Systems: A Modern Approach to Distributed Artificial Intelligence*; MIT Press: Cambridge, MA, USA, 2000.
160. Littlestone, N.; Warmuth, M.K. The weighted majority algorithm. In *30th Annual Symposium on Foundations of Computer Science*; IEEE, Research Triangle Park, NC, USA, 1989; pp. 256–261.
161. Vovk, V.G. Universal forecasting algorithms. *Information and Computation* **1992**, *96*, 245–277.
162. Poland, J.; Hutter, M. Defensive universal learning with experts. In *Proc. 16th International Conf. on Algorithmic Learning Theory (ALT'05)*, Singapore, 2005; Springer: Berlin, Germany; Vol. 3734, *LNAI*, pp. 356–370.
163. Ryabko, D.; Hutter, M. Asymptotic learnability of reinforcement problems with arbitrary dependence. In *Proc. 17th International Conf. on Algorithmic Learning Theory (ALT'06)*, Barcelona, Spain, 2006; Springer: Berlin, Germany; Vol. 4264, *LNAI*, pp. 334–347.
164. Hutter, M. General discounting versus average reward. In *Proc. 17th International Conf. on Algorithmic Learning Theory (ALT'06)*, Barcelona, Spain, 2006; Springer: Berlin, Germany; Vol. 4264, *LNAI*, pp. 244–258.
165. Legg, S.; Hutter, M. A collection of definitions of intelligence. In *Advances in Artificial General Intelligence: Concepts, Architectures and Algorithms*; Goertzel, B.; Wang, P., Eds.; IOS Press: Amsterdam, Netherlands, 2007; Vol. 157, *Frontiers in Artificial Intelligence and Applications*, pp. 17–24.
166. Legg, S.; Hutter, M. Tests of machine intelligence. In *50 Years of Artificial Intelligence*. Monte Verita, Switzerland, 2007; Vol. 4850, *LNAI*, pp. 232–242.
167. Turing, A.M. Computing machinery and intelligence. *Mind* **1950**.
168. Saygin, A.; Cicekli, I.; Akman, V. Turing test: 50 years later. *Minds and Machines* **2000**, *10*.
169. Loebner, H. The loebner prize – the first turing test. <http://www.loebner.net/Prize/loebner-prize.html> **1990**.
170. Bringsjord, S.; Schimanski, B. What is artificial intelligence? psychometric ai as an answer. *Proc. 18th International Joint Conf. on Artificial Intelligence* 2003, *18*, 887–893.
171. Alvarado, N.; Adams, S.; Burbeck, S.; Latta, C. Beyond the turing test: Performance metrics for evaluating a computer simulation of the human mind. In *Performance Metrics for Intelligent Systems Workshop*, Gaithersburg, MD, USA, 2002.
172. Horst, J. A native intelligence metric for artificial systems. In *Performance Metrics for Intelligent Systems Workshop*. Gaithersburg, MD, USA, 2002.
173. Chaitin, G.J. Gödel's theorem and information. *International Journal of Theoretical Physics* **1982**, *22*, 941–954.
174. Hernández-Orallo, J.; Minaya-Collado, N. A formal definition of intelligence based on an intentional variant of kolmogorov complexity. In *International Symposium of Engineering of Intelligent Systems*, 1998, pp. 146–163.
175. Hernández-Orallo, J. Beyond the turing test. *Journal of Logic, Language and Information* **2000**,

- 9, 447–466.
176. Hernández-Orallo, J. On the computational measurement of intelligence factors. In *Performance Metrics for Intelligent Systems Workshop*. Gaithersburg, MD, USA, 2000; pp. 1–8.
 177. Sanghi, P.; Dowe, D.L. A computer program capable of passing i.q. tests. In *Proc. 4th ICCS International Conf. on Cognitive Science (ICCS'03)*. Sydney, NSW, Australia, 2003; pp. 570–575.
 178. Legg, S.; Hutter, M. A formal measure of machine intelligence. In *Proc. 15th Annual Machine Learning Conference of Belgium and The Netherlands (Benelearn'06)*. Ghent, Belgium, 2006; pp. 73–80.
 179. Graham-Rowe, D. Spotting the bots with brains. In *New Scientist magazine*, (13 August 2005), Vol. 2512, p. 27.
 180. Fiévet, C. Mesurer l'intelligence d'une machine. In *Le Monde de l'intelligence*. Mondeo publishing, Paris, France, (2005), Vol. 1, pp. 42–45.
 181. Solso, R.L.; MacLin, O.H.; MacLin, M.K. *Cognitive Psychology*, 8th Ed.; Allyn & Bacon, 2007.
 182. Chalmers, D.J., Ed. *Philosophy of Mind: Classical and Contemporary Readings*; Oxford University Press: New York, NY, USA, 2002.
 183. Searle, J.R. *Mind: A Brief Introduction*; Oxford University Press: New York, NY, USA, 2005.
 184. Hawkins, J.; Blakeslee, S. *On Intelligence*; Times Books: New York, NY, USA, 2004.
 185. Hausser, R. *Foundations of Computational Linguistics: Human-Computer Communication in Natural Language*, 2nd Ed.; Springer: New York, NY, USA, 2001.
 186. Chomsky, N. *Language and Mind*, 3rd Ed; Cambridge University Press: Cambridge, UK., 2006.
 187. Park, M.A. *Introducing Anthropology: An Integrated Approach*, 4th Ed.; McGraw-Hill: New York, NY, USA, 2007.
 188. Bishop, C.M. *Pattern Recognition and Machine Learning*; Springer: New York, NY, USA, 2006.
 189. Turner, R. *Logics for Artificial Intelligence*; Ellis Horwood Series in Artificial Intelligence, 1984.
 190. Lloyd, J.W. *Foundations of Logic Programming*, 2nd Ed.; Springer: New York, NY, USA, 1987.
 191. Tettamanzi, A.; Tomassini, M.; Janßen, J. *Soft Computing: Integrating Evolutionary, Neural, and Fuzzy Systems*; Springer: Berlin, Germany, 2001.
 192. Kardong, K.V. *An Introduction to Biological Evolution*, 2nd Ed.; McGraw-Hill Science/Engineering/Math: New York, NY, USA, 2007.