

Towards Formalizing Non-Monotonic Reasoning in Physics: Logical Approach Based on Physical Induction and Its Relation to Kolmogorov Complexity

Vladik Kreinovich

University of Texas at El Paso, El Paso, TX 79968, USA,
vladik@utep.edu,
WWW home page: <http://www.cs.utep.edu/vladik>

Abstract. To formalize some types of non-monotonic reasoning in physics, researchers have proposed an approach based on Kolmogorov complexity. Inspired by Vladimir Lifschitz's belief that many features of reasoning can be described on a purely logical level, we show that an equivalent formalization can be described in purely logical terms: namely, in terms of physical induction.

One of the consequences of this formalization is that the set of not-abnormal states is (pre-)compact. We can therefore use Lifschitz's result that when there is only one state that satisfies a given equation (or system of equations), then we can algorithmically find this state. In this paper, we show that this result can be extended to the case of approximate uniqueness.

Keywords: non-monotonic reasoning, physical induction, uniqueness implies computability, approximate uniqueness

1 Non-Monotonic Features of Physics Reasoning and their Formalization Based on Kolmogorov Complexity

Non-monotonic features of physics reasoning. Many areas of physics – ranging from quantum physics (the physics of microscopic objects) to cosmology (the physics of very large-scale objects) – have well-defined well-studied mathematical equations and models. At first glance, one may get an impression that these equations are all we need to make conclusions about the physical world. In practice, however, in addition to equations and precise logical conclusions, physicists also use intuitive informal reasoning – some of which is non-monotonic. Specifically, they believe that not all solutions to the corresponding equations are physically meaningful – only those solutions which are, in some reasonable sense, “typical” (“not abnormal”). Let us give a few examples of such reasoning; for details, see, e.g., [7].

First example: statistical physics. The first example comes from the study of micro-objects, namely, from statistical physics. According to modern physics, all

the molecules that form a gas are constantly in random motion. It is, in principle, possible that due to this motion, all the molecules of a gas will concentrate in one half of a vessel. The probability of this event is very low, but still positive; so, from the purely mathematical viewpoint, this event may occur – we just have to wait a very long time. Physicists, however, believe that such an event is simply not possible at all [7]. Their argument is that while mathematically, such an event is possible, this event is abnormal (atypical), and we should only consider not-abnormal (typical) situations.

This physicists’ belief may sound unusual, but actually it is in good accordance with common sense. Indeed, if we toss a fair coin many times, then from the purely mathematical viewpoint, it is possible to have heads a hundred or even a million times in a row: the probability of this event is small, but if we wait long enough, it will happen. Similarly, in a state lottery, it is mathematically possible that the same person wins several times in a row. However, in practice, if the same individual wins a state lottery several times in a row, then every person using common sense will conclude that the lottery is rigged.

Second example: cosmology. Our second example comes from cosmology, the study of very large objects. According to modern physics, the large-scale state of the Universe is described by the equations of General Relativity. In principle, these equations allow many different types of solutions. Some of these solutions correspond to “generic” initial conditions, some to specific “degenerate” situations. It turns out that all solutions corresponding to the generic initial conditions have the same asymptotic. Because of this, physicists conclude that the actual space-time has this same asymptotic – this is the usual picture of the expansion following the Big Bang. The physicists’ argument is that degenerate solutions are abnormal, and the actual solution should be not-abnormal; see, e.g., [15].

Third example: general physical reasoning. One of the most productive way of making conclusions in physics is to use linearized versions of different equations. In general, the dependence $y = f(x_1, \dots, x_n)$ of different physical quantities on each other is non-linear, but when the values x_i are close to some values $x_i^{(0)}$, we can expand the dependence

$$f(x_1, \dots, x_n) = f(x_1^{(0)}, \dots, x_n^{(0)}) + \Delta x_1 \frac{\partial f}{\partial x_1} + \dots + \Delta x_n \frac{\partial f}{\partial x_n} + \dots$$

into Taylor series in terms of the differences $\Delta x_i = x_i - x_i^{(0)}$ and retain only linear terms in this expansion. The physicists’ usual argument (see, e.g., [7]) is that quadratic terms are proportional to $\Delta x_i \cdot \Delta x_j$ and, since the differences Δx_i are small, these quadratic terms can be safely ignored. Of course, the Taylor series contain each quadratic term with a numerical factor. From the purely mathematical viewpoint, this factor can be huge, in which case we can no longer ignore the corresponding quadratic term. The physicists’ argument is that such situations are abnormal, and in not-abnormal situations, each factor is reasonably small.

How to formalize such reasoning: let us start with the simplified version of the statistical case. Let us start our description with the simplified version of the above statistical case. Crudely speaking, the above case means that if an event has a very small probability, then it cannot happen. Of course, we cannot take this statement literally: for example, we believe that it is not possible to have 1000 heads in a row when tossing a coin, but every other sequence of 1000 heads and tails has the same probability 2^{-1000} , and surely one of these sequences will appear if we toss a coin 1000 times.

The simplified statement is that if an event has probability 0, then it cannot happen. This statement may also sound unusual, but it is an implicit basis of all real-life conclusions about random events. For example, we usually believe that for a fair coin, in the limit, the frequency of heads tends to $1/2$. How do we justify this belief? From the purely mathematical viewpoint, the only conclusion that we can make is that the frequency of heads converges to $1/2$ with probability 1, i.e., that the probability that the frequency *does not* converge to $1/2$ is 0. So, when we transition from this mathematically justified conclusion to a belief that for the fair coin, the frequency tends to $1/2$, we implicitly use the statement that events with probability 0 cannot happen.

Similarly, we believe that deviations of the frequency from $1/2$ are (asymptotically) normally distributed – based on the mathematical result that this asymptotical behavior occurs with probability 1.

The above implicit statement – that events with probability 0 cannot occur – is the basis of Kolmogorov-Martin-Löf formalization of the notion of a random sequence (and, more generally, a random object); see, e.g., [13]. The need for such a definition comes from the fact that in traditional statistics, there is no definition of a random sequence, while from the physics viewpoint, some sequences are random and some are not. What we want from this definition is the ability to conclude that the random sequence satisfies all the laws of probability: that the frequency tends to $1/2$, that deviations from the frequency are asymptotically normally distributed, etc. It is therefore reasonable to define a random sequence as a sequence that satisfies all the corresponding probability laws.

A probability law can be defined as a statement which is true with probability 1 – and whose negation is true with probability 0. Thus, crudely speaking, a sequence is random if it belongs to every set of probability measure 1 – or, equivalently, does not belong to every set of probability measure 0. Of course, this cannot be literally true since for every infinite sequence x , the set $\{x\}$ consisting of this very sequence has probability measure 0. To make the above definition consistent, we must therefore restrict ourselves to sets which are *definable* in some reasonable sets. Crudely speaking, a definable set is a set which is uniquely determined by a statement in some language. Every language has only countably many words, so if we require that an element does not belong to any definable set of measure 0, then we dismiss countably many set of measure 0 – i.e., a set of total measure 0. As a result, almost all sequences are random in this sense.

There are different versions of Kolmogorov-Martin-Löf complexity, depending on how we define definable sets. For example, if we consider sets which are

computable (in some reasonable sense), then we get an equivalent definition in terms of Kolmogorov complexity $K(s)$ – the shortest length of a program that generates a string s . Intuitively, a sequence which is not random – such as $0101\dots01$ – can be generated by a simple for-loop, while to generate a truly random sequence, we have to print the corresponding sequence of symbols one by one – and no shorter program can produce the given random sequence. Thus, for random strings, the Kolmogorov complexity is close to their length, while for non-random strings s , the Kolmogorov complexity is much smaller than the length $\text{len}(s)$: $K(s) \ll \text{len}(s)$. It turns out that an infinite sequence $x = x_1x_2\dots$ is random if and only if, in some reasonable sense, $K(x_1\dots x_n) \approx \text{len}(x_1\dots x_n) = n$ for all n ; see [13] for details.

From simplified version to a full statistical case. How can we formalize the physicists’ idea that an event with a very small probability cannot occur? We have already mentioned that we cannot describe this idea by simply fixing some threshold p_0 and requiring that all events with probability $< p_0$ cannot occur. Instead, it is natural to use the following idea: for every definable decreasing sequence of events $A_1 \supseteq A_2 \supseteq A_3 \supseteq \dots$, if the probability $P(A_n)$ tends to 0, then there exists an index N for which the probability is so small that this event cannot occur.

For example, for coin tosses, A_n is the set of all of the sequences for which the first n tosses resulted in all heads. Here, clearly, $A_n \supseteq A_{n+1}$ and $P(A_n) = 2^{-n} \rightarrow 0$, so there exists an N for which having N heads in a row is not possible.

In general, if we have a set X with a probability measure P , then a set $\mathcal{R} \subseteq X$ is called a *set of random elements* if for every definable sequence of sets $A_n \subseteq X$ for which $A_n \supseteq A_{n+1}$ and $P(A_n) \rightarrow 0$, there exists an integer N for which $A_N \cap \mathcal{R} = \emptyset$ – i.e., for which no element from the atypical set A_N can be viewed as truly random.

From statistical case to the general description. In our other two examples, we do not have probabilities. However, we can raise a similar argument. For example, in our third example, we do not know beforehand how large the factors need to be for the situation to become abnormal, but we are confident that some values are too large to be typical. Similarly, we may not know which human heights are abnormal, but we know that some heights are too large to be normal.

In all these cases, we can consider the set A_n of all situations in which (the absolute values of) some factors exceed n . Here, $A_n \supseteq A_{n+1}$ and $\bigcap A_n = \emptyset$, and we conclude that there exists an integer N for which none of the elements of the set A_N are typical. Thus, we arrive at the following definition; see, e.g., [8–10]. For this definition, we need to select a theory $\mathcal{L}$ which is rich enough to contain all physicists’ arguments and at the same time weak enough so that we will be able to formally talk about definability in $\mathcal{L}$; for a detailed discussion, see Appendix.

Definition 1. Let $\mathcal{L}$ be a theory, and let $P(x)$ be a formula from the language of the theory $\mathcal{L}$, with one free variable x for which, in the theory $\mathcal{L}$, there exists a set $\{x \mid P(x)\}$. We will then call the set $\{x \mid P(x)\}$ $\mathcal{L}$ -definable.

Comment. In the following text, we will assume that the language $\mathcal{L}$ is fixed, so we will simply talk about definability.

Definition 2. Let X be a set. We say that a subset $\mathcal{T}$ is a set of typical (not-abnormal) elements if for every definable sequence A_n for which $A_n \supseteq A_{n+1}$ and $\cap A_n = \emptyset$, there exists an integer N for which $A_N \cap \mathcal{T} = \emptyset$.

Physical induction: an important consequence of the above definition.

As a consequence of the above definition, we get an explanation of *physical induction*: the principle that when we have observed some property A sufficiently many times, then this property must be always true. This is how physical laws are confirmed: we perform a large number of experiments and/or observations, and if the hypothetic law is confirmed in all these experiments and observations, we consider it valid.

The principle of physical induction becomes a theorem if we assume that the state of the world s is not abnormal. Let $A(s, k)$ mean that the property A was confirmed during the k -th measurement. Then, physical induction means that there exists a natural number N_A (depending on A) such that if we have $A(s, 1)$, $A(s, 2)$, $\dots$, $A(s, N_A)$, then $A(s, n)$ holds for every natural number n .

Proposition 1. Let S be a set; its elements are called states of the world. Let $\mathcal{T} \subseteq S$ be a set of not-abnormal elements. Then, for every definable property A , there exists an integer N_A such that, if the state s is not abnormal (i.e., $s \in \mathcal{T}$) and the property $A(s, n)$ holds for all $n \leq N_A$, then the property $A(s, n)$ holds for all natural numbers n .

Comment. Physical induction can be described in purely logical terms, as the following deduction rule:

$$\frac{A(s, 1), A(s, 2), \dots, A(s, N_A), \neg ab(s)}{A(s, n)},$$

where $ab(s)$ means that $s \notin \mathcal{T}$.

Proof of Proposition 1. To prove this proposition, let us consider the the following sequence of sets

$$A_n \stackrel{\text{def}}{=} \{s : A(s, 1) \& \dots \& A(s, n) \& \neg \forall m A(s, m)\}.$$

One can easily see that this sequence is definable, that $A_n \supseteq A_{n+1}$, and that $\cap A_n = \emptyset$. Thus, by definition of a set of not-abnormal elements, there exists an integer N for which $A_N \cap \mathcal{T} = \emptyset$. This means that if $s \in \mathcal{T}$, then $s \notin A_N$. So, if $s \in \mathcal{T}$ and we have $A(s, 1)$, $\dots$, $A(s, N)$, then we cannot have $\neg \forall m A(s, m)$. Therefore, when $s \in \mathcal{T}$, we have the desired property $\forall m A(s, m)$.

Another important property: the set of typical elements is pre-compact and so, inverse problems become well-defined. Another important consequence of the above definition is related to the fact that usually, we do not directly observe the state of the world $s \in S$, we observe the result $r = f(s)$ of applying some transformation to this state. We would like to reconstruct the state s from this observation, as the state s for which $f(s) = r$, i.e., as the value $f^{-1}(r)$, where f^{-1} denotes an inverse function. This reconstruction problem is known in physics as an *inverse problem*.

One of the main challenges related to the inverse problem is that measurements are never absolutely accurate. As a result, instead of observing the exact combination of values $r = f(s)$, we observe a combination of values $\tilde{r}$ which is *close* to r . It would be nice to be able to conclude that the corresponding reconstructed state $f^{-1}(\tilde{r})$ is close to the actual state s – but for that, we need the inverse function f^{-1} to be continuous.

Most physical functions are continuous, so it is reasonable to assume that the function f is continuous. However, the inverse to a continuous function is, in general, not continuous. As a result, small changes in the measurement results can, in principle, lead to drastic changes in the reconstructed state. This discontinuity is described by saying that the inverse problem is *ill-defined*; see, e.g., [17].

To make a definite state reconstruction, physicists often make additional assumptions about the state: e.g., if we are reconstructing a signal $x(t)$, we assume certain bounds on the value of the signal and bounds on its derivative. The set of all the functions that satisfy these bounds form a compact set, and it is known that for a continuous function f from a compact set, its inverse f^{-1} is also continuous. Thus, once we impose such a restriction, the inverse problem becomes well-defined.

We will show that, in principle, there is no need to come up with artificial compactness restrictions: the mere suggestion that the state s is not abnormal (in the above precise sense) is sufficient to conclude that the corresponding set is compact. Let us describe this in precise terms.

Definition 3. *By a definable separable metric space, we mean a set X with a definable metric $d(x, y)$ and a definable sequence $\{x_n\}$ which is everywhere dense in the set X .*

Proposition 2. *Let X be a definable separable metric space, and let $\mathcal{T} \subseteq X$ be a set of typical elements. Then, the closure $\overline{\mathcal{T}}$ of this set is a compact set.*

Proof of Proposition 2. In a separable metric space, a set C is compact if and only if it is closed and for each $\varepsilon > 0$, it has a finite ε -net, i.e., a finite set $c_1, \dots, c_N$ for which every point $c \in C$ is $\leq \varepsilon$ -close to one of these points c_i : $\forall c \in C \exists i (d(c, c_i) \leq \varepsilon)$. The property that the points c_i form an ε -net is equivalent to the condition that the set C is covered by the union of the corresponding balls: $C \subseteq \bigcup B_\varepsilon(c_i)$, where $B_\varepsilon(c) \stackrel{\text{def}}{=} \{x : d(x, c) \leq \varepsilon\}$.

It is sufficient to prove the existence of an ε -net for rational values $\varepsilon > 0$ (actually, it is sufficient to prove it, e.g., for $\varepsilon = 2^{-k}$).

So, to prove that the closure $\overline{\mathcal{T}}$ is a compact set, it is sufficient to prove that for every rational number $\varepsilon > 0$, the set $\mathcal{T}$ has a finite ε -net. To prove this, let us consider the following sequence of sets: $A_n = X - \bigcup_{i=1}^n B_\varepsilon(x_i)$. This sequence is definable – we have just given a definition, and it is easy to prove that $A_n \supseteq A_{n+1}$ and that $\bigcap A_n = \emptyset$. Thus, there exists an integer N for which $A_N \cap \mathcal{T} = \emptyset$, i.e., for which $\mathcal{T} \subseteq \bigcup_{i=1}^N B_\varepsilon(x_i)$. Therefore, the elements $x_1, \dots, x_N$ form a finite ε -net for the set $\mathcal{T}$. The proposition is proven.

2 First Result: Reformulating the Above Definition of Typical Elements in Purely Logical Terms

Need for a logical formalization. Our objective is to formalize an important feature of the physicists' *reasoning*. Since logic is what describes reasoning, it is therefore natural to expect a formalization in terms of *logic*. Instead, we have a formalization in terms of sets. It is thus desirable to provide an equivalent formulation of the above definition in terms of logic.

Possibility of a logical formalization? In searching for such a logical reformulation, I was inspired by the experience of Vladimir Lifschitz who, via his numerous papers, showed that many important things related to human reasoning can be reformulated in logical terms. First, he worked in constructive mathematics, the analysis of algorithmic computability of different mathematical objects; in his research, among other things, he analyzed what can be expressed in the corresponding (intuitionistic) logic. Then, he started working in logic programming and in the formalization of commonsense reasoning; here, he also showed that many complex formalisms can be equivalently reformulated in terms of the corresponding logics.

Towards our result. In our case, there is already a logical consequence: physical induction. What we will prove here is that physical induction is not just a *consequence* of the above non-logical definition, it is actually *equivalent* to this definition.

Definition 4. Let X be a set. We say that a property $ab(x)$ describes abnormality if and only if for every definable property A , the following rule is valid for an some integer N_A (depending on A):

$$\frac{A(x, 1), A(x, 2), \dots, A(x, N_A), \neg ab(x)}{A(x, n)}.$$

Theorem 1. For every set X and for every property $ab(x)$, the following two conditions are equivalent to each other:

- the property $ab(x)$ describes abnormality (in the sense of Definition 4), and

– the set $\{x : \neg ab(x)\}$ is a set of typical elements (in the sense of Definition 2).

Proof of Theorem 1. We have already proven that if the set $\mathcal{T}$ is a set of typical elements in the sense of Definition 2, then the corresponding property $ab(x) \Leftrightarrow x \notin \mathcal{T}$ describes abnormality in the sense of Definition 4. So, to complete our proof, we need to show that vice versa, if the property $ab(x)$ describes abnormality, then the set $\mathcal{T} \stackrel{\text{def}}{=} \{x : \neg ab(x)\}$ is a set of typical elements.

Indeed, let A_n be a definable sequence of sets for which $A_n \supseteq A_{n+1}$ and $\cap A_n = \emptyset$. Let us take $A(x, k) \stackrel{\text{def}}{=} x \in A_k$. The general physical induction rules means that if we have $A(x, 1), \dots, A(x, N_A)$, and $\neg ab(x)$, then we have $\forall n A(x, n)$. In our case, this means that if $x \in A_1, \dots, x \in A_{N_A}$, and $x \in \mathcal{T}$, then for every n , we have $x \in A_n$, i.e., we have $x \in \cap A_n$. Since $\cap A_n = \emptyset$, this means that it is not possible to have $x \in A_1, \dots, x \in A_{N_A}$, and $x \in \mathcal{T}$. Thus, if we already know that $x \in \mathcal{T}$, then we cannot have $x \in A_1, \dots$, and $x \in A_{N_A}$, i.e., we must have $x \notin A_k$ for some $k \leq N_A$. For all such k , we have $A_k \supseteq A_{N_A}$, so $x \notin A_k$ implies $x \notin A_{N_A}$. Thus, $x \in \mathcal{T}$ implies that $x \notin A_{N_A}$, i.e., $\mathcal{T} \cap A_{N_A} = \emptyset$. The theorem is proven.

Comment. It is important to emphasize that physical induction is a *meta-rule*, a sequence of rules corresponding to different *definable* properties A . In general, it cannot be equivalently reformulated as a rule of second-order logic – which would mean that this implication holds for *all* properties A . Indeed, as we will show, the corresponding second-order logical statement

$$\forall A \exists N \forall x ((A(x, 1) \& \dots \& A(x, N) \& \neg ab(x)) \Rightarrow \forall n A(x, n))$$

implies that only finitely many elements are not-abnormal.

Indeed, let us assume that there are infinitely many not-abnormal elements. Then, we can find countably many among them. Let us denote these not-abnormal elements by $x_1, \dots, x_n, \dots$. Let us select the following property $A(x, k)$: $A(x, k)$ holds if and only if $x = x_i$ and $k \leq i$. According to the above second-order formula, for this property A , there exists an integer N for which, for every not-abnormal element x , the condition $A(x, 1) \& \dots \& A(x, N)$ implies that $A(x, n)$ holds for every integer n . In particular, this implication is true for a not-abnormal element x_N . For this element, by definition of the property A , we have $A(x_N, 1), \dots, A(x_N, N)$, and $\neg ab(x_N)$. Thus, we should be able to conclude that $A(x_N, n)$ holds for every integer n , but by definition of the property A , the property $A(x_N, n)$ does not hold already for $n = N + 1$.

This contradiction proves that under the second-order reformulation of physical induction, there are indeed only finitely many not-abnormal elements – and thus, that this reformulation is not adequate for describing physicists' intuition.

3 Second Result: Computability from Uniqueness to Approximate Uniqueness

Uniqueness implies computability: reminder. One of the advantages of compactness is that in a compact set, if we know that there is only one el-

ement with a certain property – e.g., the property that $F(x) = 0$ for some computable function f – then we can algorithmically find this element x . For example, if we are reconstructing the state s from measurement results $f(s) = (f_1(s), \dots, f_m(s)) = (r_1, \dots, r_m) = r$, then as the desired function $F(x)$ we can take the sum of the squares $F(x) = \sum_{i=1}^m (f_i(x) - r_i)^2$.

To describe this result – originally proven by V. Lifschitz [14] – in precise terms, let us recall the definitions of computable numbers, computable functions, and computable compact sets; see, e.g., [16, 18] (see also [1–6, 11, 12]).

Definition 5. *A real number x is called computable if there exists an algorithm (program) that transforms an arbitrary natural number k into a rational number r_k which is 2^{-k} -close to x . It is said that this algorithm computes the real number x .*

When we say that a computable real number is given, we mean that we are given an algorithm that computes this real number.

Definition 6. *A sequence of real numbers $x_1, x_2, \dots, x_n, \dots$ is called computable if there exists an algorithm (program) that transforms arbitrary natural numbers n and k into a rational number r_{nk} which is 2^{-k} -close to x_n . It is said that this algorithm computes the sequence x_n .*

When we say that a computable sequence of real numbers is given, we mean that we are given an algorithm that computes this sequence.

Definition 7. *By a computable metric space, we mean a triple $(X, d, \{x_n\})$, where (X, d) is a metric space, $\{x_1, x_2, \dots, x_n, \dots\}$ is a dense subset of X , and there exists an algorithm that, given two natural numbers i and j , computes the distance $d(x_i, x_j)$.*

In other words, we have an algorithm that, given i, j , and an accuracy k , computes the 2^{-k} -rational approximation to $d(x_i, x_j)$.

Definition 8. *A point $x \in X$ of a computable metric space $(X, d, \{x_n\})$ is called computable if there exists an algorithm that transforms an arbitrary natural number k into a natural number i for which $d(x, x_i) \leq 2^{-k}$. It is said that this algorithm computes the point x .*

A space is a compact set if there is an algorithm that, given $\varepsilon = 2^{-k}$, computes the ε -net:

Definition 9. *A computable metric space $(X, d, \{x_n\})$ is called a computable compact space if there exists an algorithm that, given an arbitrary natural number k , returns a finite set of indices $F_k \subset \{1, 2, \dots, n, \dots\}$ such that for every i there is a $f \in F_k$ for which $d(x_i, x_f) \leq 2^{-k}$.*

Many real-life quantities x, y are related by an (efficiently computable) functional relation $y = F(x)$. For example, the volume V of a cube is equal to the cube of its linear size s : $V = F(s) = s^3$. This means that, once we know the linear size, we can compute the volume.

At every moment of time, we can only know an approximate value of the actual quality $x \in X$. Thus, to be able to compute $F(x)$ with a given accuracy 2^{-k} , we must:

- be able to tell with what accuracy we need to know x , and then
- be able to use the corresponding approximation to compute $F(x)$.

We thus arrive at the following definition.

Definition 10. A function $F : X \rightarrow X'$ from a computable metric space $(X, d, \{x_n\})$ to a computable metric space $(X', d', \{x'_n\})$ is called computable if there exist two algorithms U_F and φ with the following properties:

- the algorithm φ takes a natural number k and produces a natural number $\ell = \varphi(k)$ such that $d(x, y) \leq 2^{-\ell}$ implies that $d'(F(x), F(y)) \leq 2^{-k}$;
- U_F takes two natural numbers n and k and produces a 2^{-k} -approximation to $F(x_n)$, i.e., a point x'_ℓ for which $d'(x'_\ell, F(x_n)) \leq 2^{-k}$.

Several computability results are known for computable functions on computable compact spaces.

Proposition 3. There exists an algorithm that, given a computable compact spaces X and a computable function $F : X \rightarrow \mathbb{R}$ from X to real numbers, compute its maximum and its minimum on X .

Proof. Indeed, to compute $M \stackrel{\text{def}}{=} \max F(x)$ with the accuracy 2^{-k} , we must first use the fact that F is computable and find with what accuracy $2^{-\ell}$ we must compute x to be able to estimate $F(x)$ with the accuracy $2^{-(k+1)}$. Then, we use the fact that X is a computable compact space to find a finite $2^{-\ell}$ -net. For each point x_i from this $2^{-\ell}$ -net, we compute the $2^{-(k+1)}$ -approximation $\tilde{F}(x_i)$ to the value $F(x_i)$. Then, $\tilde{M} \stackrel{\text{def}}{=} \max \tilde{F}(x_i)$ is the desired 2^{-k} -approximation to $M = \max f(x)$. Indeed, since $f(x_i) \geq \tilde{F}(x_i) - 2^{-(k+1)}$, we have

$$M = \max F(x) \geq \max F(x_i) \geq \max \tilde{F}(x_i) - 2^{-(k+1)} = \tilde{M} - 2^{-(k+1)}.$$

On the other hand, since the values x_i form a $2^{-\ell}$ -net, for every value x , there is an x_i for which $d(x, x_i) \leq 2^{-\ell}$ and hence $|F(x) - F(x_i)| \leq 2^{-(k+1)}$; therefore, $F(x) \leq \max F(x_i) + 2^{-(k+1)}$ for all x and $M = \max F(x) \leq \max F(x_i) + 2^{-(k+1)}$. Here, $F(x_i) \leq \tilde{F}(x_i) + 2^{-(k+1)}$ so

$$M \leq \max \tilde{F}(x_i) + 2^{-(k+1)} + 2^{-(k+1)} \leq \tilde{M} + 2^{-k}.$$

The proposition is proven.

Proposition 4. [3, 4] If $G : X \rightarrow \mathbb{R}$ is a computable mapping from a computable compact space X into real numbers, then, for every two rational numbers r and r' for which $r < r' \leq \max G(x)$, we can algorithmically produce a computable number $\alpha \in [r, r']$ for which the pre-image $\{x : G(x) \geq \alpha\}$ is also constructively compact (and the corresponding 2^{-k} -nets are also algorithmically produced).

Now, we are ready to reproduce (and prove) Lifschitz's result that uniqueness implies algorithmic computability:

Proposition 5. [14] *There exists an algorithm that, given a computable function $F : X \rightarrow \mathbb{R}$ that has exactly one root x_0 (for which $F(x_0) = 0$) on a computable compact space X , computes this root x_0 .*

Comment. While the result was first proven in [14], we will provide a different proof of this result, a proof that will be easy to modify to cover our new result as well.

Proof of Proposition 5. Let us show how to compute the root x_0 with a given accuracy $\delta > 0$. Let us take $\eta = \frac{\delta}{8}$, and build an η -net $\{p_1, \dots, p_k\}$ for the computable compact space X . Let us compute the distances $d(p_i, p_j)$ between the points p_i with accuracy η . As a result, we get the values $\tilde{d}(p_i, p_j)$ for which $|\tilde{d}(p_i, p_j) - d(p_i, p_j)| \leq \eta$.

According to Proposition 4, for each $i = 1, \dots, k$, there exists a value $\eta_i \in [\eta, 2\eta]$ for which the ball $B_i \stackrel{\text{def}}{=} B_{\eta_i}(p_i) = \{x : d(x, p_i) \leq \eta_i\}$ is a computable compact. Due to Proposition 3, we can compute each minimum $m_i = \min_{x \in B_i} |F(x)|$ with an arbitrary accuracy 2^{-k} . In other words, given an integer k , we can compute a rational value $\tilde{m}_{ik}$ for which $|\tilde{m}_{ik} - m_i| \leq 2^{-k}$.

For each $k = 0, 1, 2, \dots$ we compute these values $\tilde{m}_{ik}$ until for all points p_i and p_j for which $\tilde{m}_{ik} \leq 2^{-k}$ and $\tilde{m}_{jk} \leq 2^{-k}$, we get $\tilde{d}(p_i, p_j) \leq 5\eta$. Once such a k is reached, we return one of the points p_i for which $\tilde{m}_{ik} \leq 2^{-k}$ as the desired δ -approximation to the desired root x_0 .

Let us prove that this algorithm always converges, and that once it converges, the produced point p_i is indeed a δ -approximation to x_0 . Let us start with the second statement. Let us assume that the process converged. Since the points p_i form an η -net, there exists an index j for which $d(x_0, p_j) \leq \eta$. Since $\eta \leq \eta_j$, the root x_0 is within the ball $B_j = B_{\eta_j}(p_j)$ and thus, due to $|F(x_0)| = 0$ and $|F(x)| \geq 0$ for all x , we have $m_j = \min_{x \in B_j} |F(x)| = 0$. Hence, for the 2^{-k} -approximation $\tilde{m}_{jk}$ to the actual minimum $m_j = 0$, we get $\tilde{m}_{jk} \leq 2^{-k}$. So, according to our algorithm, we then have $\tilde{d}(p_i, p_j) \leq 5\eta$. Since $\tilde{d}(p_i, p_j)$ is an η -approximation to the distance $d(p_i, p_j)$, we conclude that $d(p_i, p_j) \leq \tilde{d}(p_i, p_j) \leq 5\eta + \eta = 6\eta$. From $d(x_0, p_j) \leq \eta$, we can now get $d(x_0, p_i) \leq d(x_0, p_j) + d(p_j, p_i) \leq \eta + 6\eta \leq 7\eta$. Since $\eta = \frac{\delta}{8}$, this implies that $d(x, p_i) < \delta$, i.e., that p_i is indeed the desired δ -approximation to the root x_0 .

To complete the proof, let us show that the algorithm converges. Indeed, since x_0 is the only root, for every ball B_i that does not contain x_0 , the actual minimum m_i is positive. Let m be the smallest of these positive values, and let k be such that $3 \cdot 2^{-k} \leq m$. We will show that for this k , the above algorithm will converge. Indeed, for balls that do not contain x_0 , we have $m_i \geq m \geq 3 \cdot 2^{-k}$. Since the estimate $\tilde{m}_{ik}$ of the actual minimum m_i is 2^{-k} -close to m_i , we get $\tilde{m}_{ik} \geq m_i - 2^{-k} \geq 3 \cdot 2^{-k} - 2^{-k} = 2 \cdot 2^{-k} > 2^{-k}$. Thus, the only points p_i which

will be selected by our algorithm as having $\tilde{m}_{ik} \leq 2^{-k}$ are the points for which the corresponding ball $B_i = B_{\eta_i}(p_i)$ contains x_0 . Thus, for every selected point p_i , we have $d(x_0, p_i) \leq \eta_i$. Since $\eta_i \leq 2\eta$, we get $d(x_0, p_i) \leq 2\eta$.

Let p_i and p_j be two such points. Then, we have $d(x_0, p_i) \leq 2\eta$ and $d(x_0, p_j) \leq 2\eta$ and thus, $d(p_i, p_j) \leq d(p_i, x_0) + d(x_0, p_j) \leq 2\eta + 2\eta = 4\eta$. Hence, the value $\tilde{d}(p_i, p_j)$, which is an η -approximation to the actual distance $d(p_i, p_j)$, satisfies the inequality $\tilde{d}(p_i, p_j) \leq d(p_i, p_j) + \eta \leq 4\eta + \eta = 5\eta$. Thus, the algorithm indeed stops for this value k (if it has not stopped earlier). The proposition is proven.

From uniqueness to approximate uniqueness. In practice, we may not be sure that the desired value is unique, we may only be sure that it is *approximately* unique – in the sense that for some $\varepsilon > 0$, all the roots are ε -close. Our second result extends the above computability from the uniqueness case to this approximate uniqueness case.

Theorem 2. *There exists an algorithm that, given a computable function $F : X \rightarrow \mathbb{R}$, a rational number $\varepsilon > 0$ for which all roots of F are ε -close, and the desired accuracy $\delta > 0$, returns a finite list of points $\ell_1, \dots, \ell_m$ for which $d(\ell_i, \ell_j) \leq \varepsilon + \delta$ and for which every root of F is δ -close to one of these points ℓ_i .*

Proof of Theorem 2. Similarly to the proof of Proposition 5, let us take $\eta = \frac{\delta}{8}$, and build an η -net $\{p_1, \dots, p_k\}$ for the computable compact space X . Let us compute the distances $d(p_i, p_j)$ between the points p_i with accuracy η . As a result, we get the values $\tilde{d}(p_i, p_j)$ for which $|\tilde{d}(p_i, p_j) - d(p_i, p_j)| \leq \eta$.

Similarly for the previous proof, for each $i = 1, \dots, k$, there exists a value $\eta_i \in [\eta, 2\eta]$ for which the ball $B_i \stackrel{\text{def}}{=} B_{\eta_i}(p_i) = \{x : d(x, p_i) \leq \eta_i\}$ is a computable compact. We can therefore compute each minimum $m_i = \min_{x \in B_i} |F(x)|$ with an arbitrary accuracy 2^{-k} . In other words, given an integer k , we can compute a rational value $\tilde{m}_{ik}$ for which $|\tilde{m}_{ik} - m_i| \leq 2^{-k}$.

For each $k = 0, 1, 2, \dots$ we compute these values $\tilde{m}_{ik}$ until for all points p_i and p_j for which $\tilde{m}_{ik} \leq 2^{-k}$ and $\tilde{m}_{jk} \leq 2^{-k}$, we get $\tilde{d}(p_i, p_j) \leq \varepsilon + 5\eta$. Once such a k is reached, we return all the points p_i for which $\tilde{m}_{ik} \leq 2^{-k}$ as the desired list of points $\ell_1, \dots, \ell_m$.

Let us prove that this algorithm always converges, and that once it converges, the produced list has the desired properties. Let us start with the second statement. Let us assume that the process converged. Since for selected points, we have $\tilde{d}(p_i, p_j) \leq \varepsilon + 5\eta$, and the estimate $\tilde{d}(p_i, p_j)$ is an η -approximation to the actual distance $d(p_i, p_j)$, we conclude that

$$d(p_i, p_j) \leq \tilde{d}(p_i, p_j) + \eta \leq (\varepsilon + 5\eta) + \eta = \varepsilon + 6\eta.$$

Since $\eta = \frac{\delta}{8}$, this inequality implies that $d(p_i, p_j) < \varepsilon + \delta$.

Let us now show that each root x_0 is δ -close to one of the selected points p_j . Indeed, since the points p_i form an η -net, for each root x_0 there exists an

index j for which $d(x_0, p_j) \leq \eta$. Since $\eta \leq \eta_j$, the root x_0 is within the ball $B_j = B_{\eta_j}(p_j)$ and thus, due to $|F(x_0)| = 0$ and $|F(x)| \geq 0$ for all x , we have $m_j = \min_{x \in B_j} |F(x)| = 0$. Hence, for the 2^{-k} -approximation $\tilde{m}_{jk}$ to the actual minimum $m_j = 0$, we get $\tilde{m}_{jk} \leq 2^{-k}$. So, the point p_j will indeed be selected. For this point, the inequality $d(x_0, p_j) \leq \eta$ implies that $d(x_0, p_j) \leq 8\eta = \delta$.

To complete the proof, let us show that the algorithm converges. Indeed, for every ball B_i that does not contain any root, the actual minimum m_i is positive. Let m be the smallest of these positive values, and let k be such that $3 \cdot 2^{-k} \leq m$. We will show that for this k , the above algorithm will converge. Indeed, for balls that do not contain any root, we have $m_i \geq m \geq 3 \cdot 2^{-k}$. Since the estimate $\tilde{m}_{ik}$ is 2^{-k} -close to the actual minimum m_i , we get

$$\tilde{m}_{ik} \geq m_i - 2^{-k} \geq 3 \cdot 2^{-k} - 2^{-k} = 2 \cdot 2^{-k} > 2^{-k}.$$

Thus, the only points p_i which will be selected by our algorithm as having $\tilde{m}_{ik} \leq 2^{-k}$ are the points for which the corresponding ball $B_i = B_{\eta_i}(p_i)$ contains a root x_0 . For this root, $d(x_0, p_i) \leq \eta_i$. Since $\eta_i \leq 2\eta$, we get $d(x_0, p_i) \leq 2\eta$.

Let p_i and p_j be two selected points. Then, we have two roots x_0 and x'_0 for which $d(x_0, p_i) \leq 2\eta$ and $d(x'_0, p_j) \leq 2\eta$. Since every two roots are ε -close to each other, we get $d(x_0, x'_0) \leq \varepsilon$ and thus,

$$d(p_i, p_j) \leq d(p_i, x_0) + d(x_0, x'_0) + d(x'_0, p_j) \leq 2\eta + \varepsilon + 2\eta = \varepsilon + 4\eta.$$

Hence, the value $\tilde{d}(p_i, p_j)$, which is an η -approximation to the actual distance $d(p_i, p_j)$, satisfies the inequality

$$\tilde{d}(p_i, p_j) \leq d(p_i, p_j) + \eta \leq (\varepsilon + 4\eta) + \eta = \varepsilon + 5\eta.$$

Since $5\eta < 8\eta = \delta$, for every two selected points p_i , we indeed have $\tilde{d}(p_i, p_j) \leq \varepsilon + \delta$. Thus, the algorithm indeed stops for this value k (if it has not stopped earlier). The theorem is proven.

Acknowledgments. This work was supported in part by the National Science Foundation grants HRD-0734825 and DUE-0926721, by Grant 1 T36 GM078000-01 from the National Institutes of Health.

References

1. Aberth, O.: *Precise Numerical Analysis Using C++*, Academic Press, New York (1998)
2. Beeson, M. J.: *Foundations of Constructive Mathematics*, Springer-Verlag, N.Y. (1985)
3. Bishop, E.: *Foundations of Constructive Analysis*, McGraw-Hill, (1967)
4. Bishop, E., Bridges, D. S.: *Constructive Analysis*, Springer, New York (1985)
5. Bridges, D. S.: *Constructive Functional Analysis*, Pitman, London (1979)
6. Bridges, D. S., Via, S. L.: *Techniques of Constructive Analysis*, Springer-Verlag, New York (2006)

7. Feynman, R. P., Leighton, R. B., Sands, M.: Feynman Lectures on Physics, Addison-Wesley, Boston, Massachusetts (2005)
8. Finkelstein, A. M., Kreinovich, V.: Impossibility of hardly possible events: physical consequences. Abstracts of the 8th International Congress on Logic, Methodology and Philosophy of Science, Moscow 5(2), 25–27 (1987)
9. Kreinovich, V.: Toward formalizing non-monotonic reasoning in physics: the use of Kolmogorov complexity. *Revista Iberoamericana de Inteligencia Artificial* 41, 4–20 (2009)
10. Kreinovich, V., Finkelstein, A. M.: Towards applying computational complexity to foundations of physics. *Notes of Mathematical Seminars of St. Petersburg Department of Steklov Institute of Mathematics* 316, 63–110 (2004); reprinted in *Journal of Mathematical Sciences* 134(5), 2358–2382 (2006)
11. Kreinovich, V., Lakeyev, A., Rohn, J., Kahl, P.: *Computational Complexity and Feasibility of Data Processing and Interval Computations*, Kluwer, Dordrecht (1998)
12. Kushner, B. A.: *Lectures on Constructive Mathematical Analysis*, Amer. Math. Soc., Providence, Rhode Island (1984)
13. Li, M., Vitányi, P.: *An Introduction to Kolmogorov Complexity and Its Applications*, Springer-Verlag, Berlin, Heidelberg, New York (2008)
14. Lifschitz, V. A.: Investigation of constructive functions by the method of filling. *J. Soviet Math.* 1, 41–47 (1973)
15. Misner, C. W., Thorne, K. S., Wheeler, J. A.: *Gravitation*, Freeman, San Francisco, California (1973)
16. Pour-El, M. B., Richards, J. I.: *Computability in Analysis and Physics*, Springer, Berlin (1989)
17. Tikhonov, A. N., Arsenin, V. Y.: *Solutions of Ill-Posed Problems*, V. H. Winston & Sons, Washington, DC (1977)
18. Weihrauch, K.: *Computable Analysis*, Springer-Verlag, Berlin (2000)

A Definability: A Detailed Discussion

To make formal definitions, we must fix a formal theory $\mathcal{L}$ that has sufficient expressive power and deductive strength to conduct all the arguments and calculations necessary for working physics. For simplicity, in the arguments presented in this paper, we consider ZF, one of the most widely used formalizations of set theory.

Using ZF is a little bit of an overkill; a weaker arithmetic system RCA_0 is believed to be quite sufficient to formalize all of nowadays physics. Our definitions and results will not seriously depend on what exactly theory we choose – in the sense that, in general, these definitions and proofs can be modified to fit other appropriate theories $\mathcal{L}$.

A formal definition of definability is given by Definition 1. Crudely speaking, a set is $\mathcal{L}$ -definable if we can explicitly *define* it in $\mathcal{L}$. The set of all real numbers, the set of all solutions of a well-defined equation, every set that we can describe in mathematical terms is $\mathcal{L}$ -definable.

This does not mean, however, that *every* set is $\mathcal{L}$ -definable: indeed, every $\mathcal{L}$ -definable set is uniquely determined by formula $P(x)$, i.e., by a text in the language of set theory. There are only denumerably many words and therefore,

there are only denumerably many $\mathcal{L}$ -definable sets. Since, e.g., in a standard model of set theory ZF, there are more than denumerably many sets of integers, some of them are thus not $\mathcal{L}$ -definable.

A sequence of sets $\{A_n\}$ is, from the mathematical viewpoint, a mapping from the set of natural numbers to set of sets, i.e., a set of all the pairs $\langle n, A_n \rangle$. Thus, Definition 1 leads to the following natural definition of the notion of an $\mathcal{L}$ -definable sequence:

Definition A1. *Let $\mathcal{L}$ be a theory, and let $P(n, x)$ be a formula from the language of the theory $\mathcal{L}$, with two free variables n (for integers) and x . If, in some model of the theory $\mathcal{L}$, the set $\{\langle n, x \rangle \mid P(n, x)\}$ is a sequence (i.e., for every n , there exists one and only one x for which $P(n, x)$), then this sequence will be called $\mathcal{L}$ -definable.*

Our objective is to be able to make mathematical statements about $\mathcal{L}$ -definable sets. Therefore, in addition to the theory $\mathcal{L}$, we must have a stronger theory $\mathcal{M}$ in which the class of all $\mathcal{L}$ -definable sets is a set – and it is a countable set.

Denotation. *For every formula F from the theory $\mathcal{L}$, we denote its Gödel number by $\lfloor F \rfloor$.*

Comment. A Gödel number of a formula is an integer that uniquely determines this formula. For example, we can define a Gödel number by describing what this formula will look like in a computer. Specifically, we write this formula in $\text{\LaTeX}$, interpret every $\text{\LaTeX}$ symbol as its ASCII code (as computers do), add 1 at the beginning of the resulting sequence of 0s and 1s, and interpret the resulting binary sequence as an integer in binary code.

Definition A2. *We say that a theory $\mathcal{M}$ is stronger than $\mathcal{L}$ if it contains all formulas, all axioms, and all deduction rules from $\mathcal{L}$, and also contains a special predicate $\text{def}(n, x)$ such that for every formula $P(x)$ from $\mathcal{L}$ with one free variable, the formula*

$$\forall y (\text{def}(\lfloor P(x) \rfloor, y) \leftrightarrow P(y))$$

is provable in $\mathcal{M}$.

The existence of a stronger theory can be easily proven:

Proposition A1. [10] *For $\mathcal{L}=\text{ZF}$, there exists a stronger theory $\mathcal{M}$.*

Comments. In this paper, we assume that a theory $\mathcal{M}$ that is stronger than $\mathcal{L}$ has been fixed; proofs will mean proofs in this selected theory $\mathcal{M}$.

An important feature of a stronger theory $\mathcal{M}$ is that the notion of an $\mathcal{L}$ -definable set can be expressed within the theory $\mathcal{M}$: a set S is $\mathcal{L}$ -definable if and only if $\exists n \in \mathbb{N} \forall y (\text{def}(n, y) \leftrightarrow y \in S)$.

In the paper, when we talk about definability, we mean this property expressed in the theory $\mathcal{M}$. So, all the statements involving definability (e.g., the Definition 2) become statements from the theory $\mathcal{M}$ itself, *not* statements from metalanguage.