

Information Theory, Evolutionary Computation, and Dembski's "Complex Specified Information"

Wesley Elsberry
Department of Wildlife & Fisheries Sciences
Texas A&M University
2258 TAMU
College Station, Texas 77843-2258
USA
`welsberr@inia.cls.org`

Jeffrey Shallit
School of Computer Science
University of Waterloo
Waterloo, Ontario N2L 3G1
Canada
`shallit@graceland.uwaterloo.ca`

November 7, 2003

Abstract

Intelligent design advocate William Dembski has introduced a measure of information called "complex specified information", or CSI. He claims that CSI is a reliable marker of design by intelligent agents. He puts forth a "Law of Conservation of Information" which states that chance and natural laws are incapable of generating CSI. In particular, CSI cannot be generated by evolutionary computation. Dembski asserts that CSI is present in intelligent causes and in the flagellum of *Escherichia coli*, and concludes that neither have natural explanations. In this paper we examine Dembski's claims, point out significant errors in his reasoning, and conclude that there is no reason to accept his assertions.

1 Introduction

In recent books and articles (e.g., [16, 17, 19]), William Dembski uses a semi-mathematical treatment of information theory to justify his claims about "intelligent design". Roughly

speaking, intelligent design advocates attempt to infer intelligent causes from observed instances of complex phenomena. Proponents argue, for example, that biological complexity indicates that life was designed. This claim is sometimes offered as an alternative to the theory of evolution.

Christian apologist William Lane Craig has called Dembski’s work “groundbreaking” [17, blurb at beginning]. Journalist Fred Heeren describes Dembski as “a leading thinker on applications of probability theory” [38].¹ At a recent conference [53], University of Texas philosophy professor Robert Koons called Dembski the “Isaac Newton of information theory.”² Is such effusive praise warranted?

We believe it is not. As we will show, Dembski’s work is riddled with inconsistencies, equivocation, flawed use of mathematics, poor scholarship, and misrepresentation of others’ results. As a result, we believe few if any of Dembski’s conclusions can be sustained.

Several writers have already taken issue with some of Dembski’s claims (e.g., [23, 71, 72, 92, 78, 22, 94, 32]). In this paper we focus on some aspects of Dembski’s work that have received little attention thus far.

Here is an outline of the paper. First, we summarize what we see as Dembski’s major claims. Next, we criticize Dembski’s concept of “design” and “intelligence”. We then turn to one of Dembski’s major tools, “complex specified information”, arguing that he uses the term inconsistently and misrepresents the concepts of other authors as being equivalent. We criticize Dembski’s concept of “information” and “specification”. We then address his “Law of Conservation of Information”, showing that the claim has significant mathematical flaws. We then discuss Dembski’s attack on evolutionary computation, showing his claims are unfounded. Finally, we issue a series of challenges to those who would continue to pursue intelligent design. Some of our ideas are based on Kolmogorov complexity, so we provide an introduction to this theory as an appendix. The appendix also contains an alternate account of specification and a suggested replacement for CSI.

We note that our criticism is based on *all* of Dembski’s *oeuvre*, not simply his most recent work. We regard this as completely legitimate; all of Dembski’s claims are assumed to be in force unless explicitly retracted, and virtually no retractions have been forthcoming.

2 Dembski’s claims

Dembski makes a variety of different claims, many of which would be revolutionary if true. Here we try to summarize what appears to us to be his most significant claims, together with the section numbers in which we address those claims.

1. The “complexity-specification” criterion/“explanatory filter” is a reliable method for detecting design by intelligent agents, and accurately reflects how humans traditionally infer design. (§5, 6)

¹According to the American Mathematical Society’s online version of *Mathematical Reviews*, a journal which attempts to review every noteworthy mathematical publication, Dembski has not published a single paper in any journal specializing in applied probability theory, and a grand total of one peer-reviewed paper in any mathematics journal at all.

²According to *Mathematical Reviews*, Dembski has not published any papers in any peer-reviewed journal devoted to information theory.

2. There exists a multi-step statistical procedure, the “generic chance elimination argument”, that reliably detects design by intelligent agents. (§4, 6)
3. There is a “souped-up” form of information called “specified complexity” or “complex specified information” (CSI) which is coherently defined and constitutes a valid, useful, and non-trivial measure. (§6)
4. Many human activities exhibit “specified complexity”. (§6)
5. The presence of CSI is a reliable marker of design by intelligent agents. (§5)
6. CSI cannot be generated by deterministic algorithms, chance, or any combination of the two. In particular, CSI cannot be generated either by genetic algorithms implemented on computers, or the process of biological evolution itself. A “Law of Conservation of Information” exists which says that natural processes cannot generate CSI. (§9, 10)
7. Life exhibits specified complexity and hence was designed by an intelligent agent, possibly disembodied. (§11)

3 Design

Dembski’s account of design is inconsistent. On the one hand, he never gives a positive account of design; we do not learn from reading his works what Dembski thinks design *is*. In *The Design Inference* [16] he simply defines design as the complement of regularity and chance, and the possibility that this complement is in fact empty is not seriously addressed. In *No Free Lunch* [19, p. xi] he gives a process-oriented account of design:

(1) A designer conceives a purpose. (2) To accomplish that purpose, the designer forms a plan. (3) To execute the plan, the designer specifies building materials and assembly instructions. (4) Finally, the designer or some surrogate applies the assembly instructions to the building materials.

But this is not a positive account of what constitutes design. Furthermore, the description is problematical. In common parlance, “design” can mean “pattern” or “motif”, and the relationship between “pattern” and “purpose” is unclear. Intelligent design advocates claim that “design implies a designer”, but perhaps this claim owes more to the structure of English than it does to logic. After all, we would not likely say “pattern implies a patterner”.

It is certainly easy to claim a teleological account of biology, but other natural processes produce “design”, in the sense of pattern, without evidently falling under the process-oriented view of design that Dembski provides. Consider, for example, the highly symmetrical 6-sided patterns that appear in snowflakes. If there is any evidence of purpose in the patterns seen in snowflakes, it eludes us. We address this issue in more detail in Section 9.3.

Dembski pleads for more consideration of design as a scientific explanation, but he seems to be of two minds concerning this. On the one hand, he claims “science has largely dispensed with design” and science “repudiates design” [19, p. 3]; on the other hand, just three pages later he cites archaeology [19, p. 6] as an example of a science that is based in part on

inferring design. Contrary to Dembski's assertions, design is *not* arbitrarily ruled out as an element of scientific explanation, even in biology.

Scientists, however, are reluctant to infer "rarefied" design, a design inference based on ignorance of both the nature of the designer and regularities that might explain the observed phenomenon. But this reluctance is well-grounded. Empirically gained knowledge of designers and the artifacts which they create permit us to recognize *regularities of outcomes*, leading us to make an "ordinary" design inference in such cases. With an "ordinary" design inference, a designer becomes just another causal regularity. This is not so with a "rarefied" design inference, which Dembski urges us to make in ignorance of the properties of any putative designer and also of other causal regularities which may be operative. For more details, see [94].

Where Dembski offers examples that have practical application, one finds that the operative mode of inference is to an "ordinary" design inference. The appeal to SETI is such an example, for actual SETI research is based upon knowledge of how intelligent agents (humans) actually use radio wavelengths for communication. Another is Dembski's claim that "the Smithsonian Institution devotes a room to obviously designed artifacts for which no one has a clue what those artifacts do." [19, p. 147]³ Dembski overlooks the fact that artifacts of "ordinary" designers can be recognized not only through something like his concept of CSI, but also by the more prosaic methods long employed in archaeology. These methods include signs of working of an artifact, where a chance explanation is eliminated due to the artifact showing characteristic signs of manipulation that we know by experience are attributable to human artisans. The Smithsonian example turns out to be non-mysterious and unsupportive of Dembski's attempt to justify a "rarefied" design inference.

Dembski claims that, using his methodology, one can infer the existence of an intelligent designer responsible for certain forms of observed design. Sometimes he views this as the first step of a scientific inquiry: "Once specified complexity tells us that something is designed, there is nothing to stop us from inquiring into its production." [19, p. 112]. Later, however, he claims that both the "Intentionality Problem—What was the intention of the designer in producing a given designed object?" and the "Identity Problem—Who or what is the designer?" are not legitimate "questions of science" at all! [19, p. 313] This is especially noteworthy given Dembski's discussion of intentionality in his book, *Intelligent Design* [17, pp. 245–246].

Removing intention and identity from rational inquiry may be legitimate if, as many intelligent design advocates admit when pressed, the designer they have in mind is a disembodied supernatural being [65]. But it is certainly *not* legitimate if the designer is human, or even an extraterrestrial being. "Explaining" crop circles as the product of alien design does not end the inquiry; instead, it enlarges it. Where did the aliens come from? Why did they wish to create the circles? And so forth.

³According to a letter dated March 30, 2002 from Kenneth Burke, Acting Program Coordinator, Public Inquiry Mail Service, Smithsonian Institution, "The Smithsonian has no room such as described in William Dembski's book." Burke goes on to state that in *one showcase* of a 1980–1 exhibition at the then National Museum of History and Technology "a number of unidentified articles were displayed, but there was never a whole room devoted to them."

Furthermore, questions of intention and identity arise all the time in archaeology. To give just two examples: in the 1890's historian Arthur Evans heard of mysterious seal-stones from Crete. The identity of their creators, as well as the script used, was then unknown. Evans went on to identify the stones as the product of a civilization now called Minoan, and eventually one of the scripts, Linear B, was deciphered [24]. Similarly, the intention of the artists of the wall-paintings of the Bronze Age wall paintings from Thera is an active area of scientific controversy, with some arguing that rooms with such paintings were always intended to be shrines, and others disputing this [68, 64, 21]. Despite Dembski wishing to rule identity and intention of designers out of science, archaeologists are quite happily pursuing these questions.

Quoting Jay Richards, Dembski says “If someone explains some buried earthenware as the result of artisans from the second century B.C., no one complains, ‘Yeah, but who made the artisan?’ ” [19, p. 355]. We find this reply altogether too facile. In the case of human artisans from 200 B.C.E, we need no extraordinary explanation to account for their existence — there is abundant evidence for human life and pottery-making culture during that time period. On the other hand, if we found buried earthenware in Devonian strata, and the explanation proffered was “artisans from 300 mya”, scientists certainly *would* want to inquire about the origin of the artisan.

Similarly, Dembski says if we find a scrap of paper with writing on it, we infer a human author and “there is no reason to suppose that this scrap of paper requires a different type of causal story” [19, p. xi]. But surely this depends upon the circumstances of the find and the causal hypothesis which is proposed to account for it. If Neil Armstrong had found a scrap of paper with writing on it on the moon, the remoteness of the location and a hypothesis that the writing was done in situ would conjointly exclude human agency and require a different type of causal story.

These are symptoms of a more general inconsistency in the level of explanation Dembski wishes to pursue. For Dembski, explanation of design always ends at intelligence, and we are not permitted to inquire further about the origin of the intelligence. We continue this line in the next section.

4 Intelligence

Just as Dembski fails to give a positive account of the second half of “intelligent design”, he also fails to define the first half: intelligence. Intelligence, and intelligent agents, are treated as unfathomable mysteries beyond human comprehension, and not explainable by natural causes. He writes

I will argue that intelligent agency, even when conditioned by a physical system that embodies it, cannot be reduced to natural causes without remainder. Moreover, I will argue that specified complexity is precisely the remainder that remains unaccounted for. Indeed, I will argue that the defining feature of intelligent causes is their ability to create novel information and, in particular, specified complexity. [19, p. xiv]

Dembski does not accept that intelligence itself could arise purely through natural processes, via evolution:

Out pop purpose, intelligence, and design from a process that started with no purpose, intelligence, or design. This is magic. [19, p. 369]

But this skepticism is apparently based in part on belief in a sharp distinction between intelligent and non-intelligent causes: agency is always either natural or intelligent, and cannot be both. But what if purpose, intelligence, and design are words we assign to emergent properties of complex systems? What if intelligence is not a binary classification, but a multifactorial gradation, with thermostats and bacteria being only slightly intelligent, and computers and rats more so?

Intelligent agency always receives preferential treatment in Dembski's analysis: in his explanatory filter framework, he never allows hypotheses involving intelligent agents to be eliminated. Consider his analysis of the Nicholas Caputo case. Caputo was an Essex County, New Jersey official charged with deciding assigning the order of political parties on the ballot in local elections. Caputo, a Democrat, chose the Democrats first in 40 of 41 elections. Writing D for Democrat and R for Republican, Dembski proposes considering the string

$c = \text{DDDDDDDDDDDDDDDDDDDDDDDRDDDDDDDDDDDDDDDDDD}$

that represents the sequence of choices to head the ballot. Did Caputo cheat?

When Dembski analyzes this case, he applies his “generic chance elimination argument”, which is supposed to “sweep the field clear” of all relevant chance hypotheses. (Chance hypotheses, in Dembski's idiosyncratic terminology, also include purely deterministic hypotheses in which no chance was actually involved.) What remains is the conclusion that Caputo's selections were due to the mysterious process Dembski calls design.

But in fact the only chance hypothesis that Dembski considers is that Caputo's selections arose by the flipping of a fair coin. He does not consider other possibilities, such as

- (a) Caputo really had no choice in the assignment, since a mobster held a gun to his head on all but one occasion. (On that one occasion the mobster was out of town.)
- (b) Caputo, although he appears capable of making choices, is actually the victim of a severe brain disease that renders him incapable of writing the word “Republican”. On one occasion his disease was in remission.
- (c) Caputo was molested by a Republican at an early age, and the resulting trauma has caused a pathological hatred of Republicans. He therefore tends to favor Democrats, but on one occasion a Republican bought him a beer immediately prior to the ballot assignment.
- (d) Caputo attempted to make his choices randomly, using the flip of a fair coin, but unknown to him, on all but one occasion he accidentally used a two-headed trick coin from his son's magic chest. Furthermore, he was too dull-witted to remember assignments from previous ballots.

- (e) Caputo himself is the product of a 3.8-billion-year-old evolutionary history involving both natural law and chance. The structure of Caputo’s neural network has been shaped by both this history and his environment since conception. Evolution has shaped humans to act in a way to increase their relative reproductive success, and one evolved strategy to increase this success is seeking and maintaining social status. Caputo’s status depended on his respect from other Democrats, and his neural network, with its limited look-ahead capabilities, evaluated a fitness function that resulted in the strategy of placing Democrats first in order to maximize this status.

What are we trying to say in this list of possibilities, some less serious than others? Simply that if Caputo flipping a fair coin is one of the possibilities to be eliminated, it is unclear why Caputo himself cannot figure in *other* chance hypotheses we would like to eliminate. Some of these chance hypotheses, such as (b), involve Caputo, but do not involve design as we understand the word. Others involve design as generally understood. Hypothesis (e), which could well be the correct explanation, is based on a very complex causal chain of billions of steps, most of which we will probably be unable to judge the probability of with any certainty. Currently we cannot rule (e) in or out based *solely* on estimates of probability; we must rely on its consilience with other facets of science, including evolutionary biology, psychology, and neuroscience.

This leads us to what we see as one of the weakest points of Dembski’s argument: if, as he suggests, design is always inferred simply by ruling out known hypotheses of chance and necessity, then *any* observed event with a sufficiently complicated or obscure causal history could mistakenly be assigned to design, either because we cannot reliably estimate the probabilities of each step of that causal history, or because the actual steps themselves are currently unknown. We call this the “Erroneous Design Inference Principle”, or EDIP.

The existence of EDIP receives confirmation from modern research in psychology. For one thing, humans are notoriously poor judges of probability [44]. On the other hand, humans are good detectors of patterns, even when they are not there [7, 96, 39, 80]. Humans also have “agency-detection systems” which are “biased toward overdetection”, a fact some have explained as consonant with an evolutionary history where systems for detecting prey were strongly selected for [4]. Taken together, these factors suggest that it will be common for design to be inferred erroneously, and perhaps explains the large number of cases falling under the EDIP: ghosts, UFO’s, and witchcraft.

Assigning intelligent agency based on ignorance of the precise causal history of an event or the probabilities associated with a hypothesized route seems in opposition to Dembski’s assertion that “frank admissions of ignorance are much to be preferred to overconfident claims to knowledge that cannot in the end be adequately justified” [19, p. 316]. The target of this assertion is “Darwinism”, but it seems to us far more apposite to Dembski’s own conclusions about design.

But back to our analysis of the Caputo case. If the only chance hypothesis that is being considered is that the sequence of ballot assignments resulted from the flips of a fair coin, then Dembski’s analysis has little novelty to it. As Laplace wrote in 1819 [57, pp. 16–17]:

In the game of heads and tails, if heads comes up a hundred times in a row then this appears to us extraordinary, because the almost infinite number of combi-

nations that can arise in a hundred throws are divided in regular sequences, or those in which we observe a rule that is easy to grasp, and in irregular sequences, that are incomparably more numerous.

Laplace’s argument has been updated in modern form to reflect Kolmogorov complexity; see, for example, the wonderful article [51]. The probability that a string x of length n (whose bits are chosen with uniform probability $p = 1/2$) will have $C(x) \leq m$ can be shown to be $\leq 2^{m+1-n}$. The Kolmogorov complexity of $\mathbf{c}$ is very low (we cannot compute it exactly, but let’s say for the sake of argument that $C(\mathbf{c}) \leq 10$). Thus the hypothesis that $\mathbf{c}$ is due flipping a fair coin has probability $\leq 2^{-30}$, or about 1 in a billion, and it seems fair to reject it.

What next? Dembski would have us believe that design is now a purely mathematical implication. But what of the possibilities (a)–(e) given above? Instead of design being a purely eliminative argument, we see that design hypotheses must be considered along with other statistical hypotheses.

To see this, consider Dembski’s discussion of the SETI primes sequence

$$\mathbf{t} := 110111011111011111110 \overbrace{111 \dots 1}^{11} 0 \overbrace{111 \dots 1}^{13} 0 \dots \overbrace{111 \dots 1}^{89} 0 \overbrace{111 \dots 1}^{73} 1,$$

which is a variation on a signal received by fictional researchers in the movie *Contact*. As Dembski describes it, it consists of blocks of consecutive 1’s separated by 0’s, whose lengths encode the prime numbers from 2 to 89, with extra 1’s at the end to make the length 1000. Dembski suggests the specified complexity of this sequence implies a design inference.⁴

Yet is that the case? We know that prime numbers arise naturally in simple predator-prey models [34], so it is at least conceivable that prime number signals could result from some non-intelligent physical process. To infer intelligent design upon receiving $\mathbf{t}$ simply means that we estimate the relative probability of natural prime-number generation as lower than the probability that the signal arises from an intelligence that considers prime numbers an interesting way to communicate. In other words, we explicitly compare two hypotheses, one involving design, one not.

Dembski is fond of fictional examples, so it is instructive to compare Dembski’s treatment of the cinematic SETI sequence with the history of an actual reception of an extraterrestrial signal. Pulsars (rapidly pulsating extraterrestrial radio sources) were discovered by Jocelyn Bell in 1967. She observed a long series of pulses of period 1.337 seconds. In at least one case the signal was tracked for 30 consecutive minutes, which would represent approximately 1340 pulses. Like the SETI sequence, this sequence was viewed as improbable (hence “complex”) and specified (see Section 8), hence presumably it would constitute complex specified

⁴This example is particularly misleading because of the iconic nature of prime numbers in the mind of average person. In common parlance, prime numbers might be thought to be “complex” because their distribution seems mysterious and because there remain so many unsolved conjectures about them (e.g., Goldbach’s conjecture). But in terms of algorithmic information theory, the SETI primes sequence, and lengthier versions of it, are actually quite simple, since they can be generated by a very short program.

In fact, using Dembski’s terminology, receiving a sequence of length 1000 such as $0101 \dots 01$ or even $000 \dots 0$ would be just as “complex” as the SETI primes sequence and would also trigger a design inference. The fact that the sequence he discusses represents prime numbers is actually a bit of a red herring, but one that allows him to trade off on the ambiguous notion of “complex”.

information and trigger a design inference. Yet spinning neutron stars, and not design, are the current explanation for pulsars.

Bell and her research team immediately considered the possibility of an intelligent source. (They originally named the signal LGM-1, where the initials stood for “little green men”.) The original paper on pulsars states “The remarkable nature of these signals at first suggested an origin in terms of man-made transmissions which might arise from deep space probes, planetary radar, or the reflexion of terrestrial signals from the Moon” [41].

However, the hypothesis of intelligent agency was rejected for two reasons. First, parallax considerations ruled out a terrestrial origin. Second, additional signals were discovered originating from other directions. The widely separated origins of multiple signals decreased the probability of a single intelligent source, and multiple intelligent sources were regarded as implausible. In other words, hypotheses involving design were considered at the same time as non-design hypotheses, instead of the eliminative approach Dembski proposes.

4.1 Animal intelligence

As we have stated, Dembski seems to view agents as either intelligent or not. There is no notion of intelligence as a quality which is subject to fine gradation in his writing. In particular, he does not address in any detail whether living creatures other than humans are intelligent. Rats successfully running mazes are accorded intelligence by Dembski, but the larger issue of animal intelligence is not directly addressed [15].

On the other hand, all of Dembski’s examples about design are either about human design or supposed supernatural design. This is not a large evidence base from which to draw. But there is a vast area of pattern construction by agents that Dembski simply ignores: patterns by animals.

The animal world bustles with patterns. From 7-meter high mounds constructed by the termite *Nasutitermes triodiae* to the hexagonal honeycombs of bees, from the bowers of bowerbirds *Ptilonorhynchidae* to the intricate song of the great reed warbler *Acrocephalus arundinaceus* and the woodlark *Lullula arborea*, from dolphin vocalizations to the waggle dance of bees, animals seem to generate patterns that might qualify as showing the properties of Dembski’s “complex specified information”.

If a Shakespearean sonnet represents CSI that identifies an intelligent agent, it seems that the termite mound should identify the termite as an intelligent agent. (But see [90] for an explanation of how termite mounds can be generated through a simple distributed algorithm.) Dembski’s framework for identifying “complex specified information” and his assertion that CSI reliably marks the action of an intelligent agent therefore counts as a proposal of a demarcation criterion for intelligence in general. This would represent a remarkable advance in studies of animal cognition and have general utility in the formulation of public policy regarding management of species whose members emit CSI and thus display this hallmark of intelligence. Remarkably, neither Dembski nor any other intelligent design proponent have promulgated Dembski’s framework for this purpose.

Dembski must tread carefully in this regard. Pattern generation in plants as an example of a natural process at work is disputed by Dembski [19, pp. 13–14]. The patterns seen there, Dembski asserts, could be analogous to a computer generating a pattern according

to a program. The program itself would be due to an external designing intelligence. Deploying this same response in the case of animal intelligence, though, would be disastrous to Dembski's apparent views of human intelligence. If the patterns emitted by animals could be the result of some externally-produced 'program' for which the animals simply act as a substrate for proper execution of its steps, then it would seem that there is nothing in Dembski's framework which would prevent the same explanation being forwarded for the patterns that humans generate. It seems a legitimate question for Dembski to resolve how we might distinguish between the case where a human is considered to "generate" CSI and one where the human merely emits CSI due to a program infused by an external designer. The answer should also be generally applicable to non-human animals.

5 The validity of the design inference

Dembski offers two reasons to believe that his "specified complexity" is a reliable indicator of intelligent design. The first is a novel form of induction:

In every instance where the complexity-specification criterion attributes design and where the underlying causal story is known (i.e., where we are not just dealing with circumstantial evidence, but where, as it were, the video camera is running and any putative designer would be caught red-handed), it turns out design actually is present; therefore, design actually is present whenever the complexity-specification criterion attributes design.

We find this form of induction completely unjustified. Dembski's inductive argument places too onerous a burden on possible counterexamples. In order to overturn this argument, Dembski requires examples with an extraordinary level of detail concerning their causal history. By its arbitrary demand on video-camera-certainty, Dembskian induction rules out considering cases where the underlying causal story cannot be known with this form of certainty because it is either very complicated or occurred long ago. In such cases, in the absence of a time machine, the causal story we develop can be justified *only* through circumstantial evidence. This is often the case in historical sciences. In particular, there is abundant circumstantial evidence that Darwinian processes can account for complexity in nature, but Dembski excludes this evidence because it does not pass his video-camera-certainty test. Indeed, Dembskian induction seems intelligently designed to rule out a naturalistic explanation of biological complexity.

To see the insufficiency of Dembskian induction, let us consider two cases where we try to apply Dembski's reasoning.

1. Let us consider the case of extinction of species, and suppose we wish to argue that all such extinctions are ultimately due to human action. Now in every case where extinction has occurred and the underlying causal story is known with as much certainty as Dembski demands, humans were ultimately responsible: consider, for example, the extinction of the passenger pigeon (*Ectopistes migratorius*), the thylacine (*Thylacinus cynocephalus*), the dodo (*Raphus cucullatus*), and a depressingly long list of avian

species in the Hawaiian islands. Using Dembskian induction, we would conclude that humans must be responsible for all extinctions. But this conclusion can only be retained by throwing out the large body of circumstantial evidence pointing to other factors apparent in the geologic past, such as climate change, tectonic rearrangements, and asteroid or comet impacts. These events cannot pass the video-camera-certainty test.

2. Let us consider the construction of tall pillars made of hard material, such as stone columns. We wish to argue that all such pillars are due to intelligent agents. Now in every case where a pillar appears and the underlying causal story is known with the certainty Dembski demands, these pillars were constructed by intelligent agents (humans). Using Dembskian induction, we would conclude that intelligent agents must be responsible for all such pillars, including the sand pipes at Kodachrome Basin State Park in Utah and the basalt columns at the Giant's Causeway in Ireland. But this conclusion can only be retained by ruling out the circumstantial evidence in favor of accepted geological explanations for these features (ancient geysers and split volcanic flows, respectively; see [43]).

We also note that Dembskian induction has certain consequences Dembski may not have realized. Let us consider the case of receiving communication expressed in a human language, as in a letter, scroll, or book. In every case where such a communication was received and the underlying causal story is known with as much certainty as Dembski demands, humans were ultimately responsible. Therefore, by Dembskian induction, humans are ultimately responsible whenever communication in a human language is received. In other words, the Bible is of human origin. This conclusion is unlikely to receive general assent from Dembski's supporters.

We now turn to the second justification Dembski gives for inferring design from specified complexity: the supposed nature of intelligent agency. Intelligent agents, Dembski tells us, make choices by "actualizing one among several competing possibilities, ruling out the rest, and specifying the one that was actualized" [19, p. 29], and this is what specified complexity detects.

We find this justification weak. First, Dembski gives no reason to suppose that the properties he asserts for intelligent agents are exclusive to those intelligent agents. For example, the actualization-exclusion-specification triad of criteria can be viewed as a description of the process of natural selection.

Second, this reasoning has strange implications. Since the decisions of intelligent agents are supposedly not reducible to chance and natural law, it follows that these decisions are irrational, in the sense of being inexplicable through rational processes. We end up with the paradoxical conclusion that if specified complexity really exists, then it identifies not intelligence, but irrationality. Further, this irrationality extends to the Designer that Dembski argues is responsible for biological complexity, a conclusion hardly likely to receive assent from Dembski's supporters.

An alternate view is that if specified complexity detects anything at all, it detects the output of simple computational processes. This is consonant with Dembski's claim "It is CSI that within the Chaitin-Kolmogorov-Solomonoff theory of algorithmic information identifies

the highly compressible, nonrandom strings of digits” [19, p. 144]. Dembski’s inference of design is then undermined by the recent realization that there are many naturally-occurring tools available to build simple computational processes. To mention just four, consider the recent work on quantum computation [42], DNA computation [47], chemical computing [55, 89, 74], and molecular self-assembly [79]. Furthermore, it is now known that even very simple computational models, such as Conway’s game of Life [3], Langton’s ant [26], and sand piles [33] are universal, and hence compute anything that is computable. Finally, in the cellular automaton model, relatively simple replicators are possible [5].

Under this interpretation, inferring design upon observing specified complexity implicitly ranks “production by unintelligent natural computational process” as less likely than “production by intelligent agent”. Again, this is an explicit comparison of design and non-design hypotheses, which Dembski rejects.

5.1 Indirect design

Dembski defends his concept of specified complexity from the challenge of evolutionary computation by asserting that what results from evolutionary computation (and all other algorithmic processes) is at best *apparent specified complexity*, not actual specified complexity [18]. In all such cases, the specified complexity is asserted to have been present in the inputs to the algorithm or somehow infused by an intelligent agent in the process. This immediately leads to a conclusion that Dembski’s explanatory filter/design inference is incapable of resolving the difference between *apparent specified complexity* and actual specified complexity. In order to accomplish the discrimination of actual and apparent specified complexity, it is absolutely necessary to have information about the actual causation of the event. But Dembski wishes us to utilize his explanatory filter/design inference in precisely those cases where such information is not available. It is obvious that in such cases the explanatory filter/design inference is uninformative as to whether any specified complexity found is actual or only apparent.

An illustration that may be helpful is to consider the “Algorithm Room.” Dembski often refers to examples of cheating by intelligent agents as fruitful places to deploy his explanatory filter/design inference. Imagine that we have a room in which a computational expert (assumed for the purpose of argument to be an intelligent agent) works. The expert also has various computing machinery and an extensive library of works on algorithms in the room with him. The expert has two pay scales, a very high one for consultancy, where he asserts that any work performed is original to himself, and a much more modest pay scale for computer technician work, where the expert simply applies pre-existing algorithms to the task at hand. A problem description and input data is passed into the room through a slot in the door. After a reasonable amount of time for an expert to have solved the problem elapses, a solution is passed back out, along with a bill for time at the consultancy rate. The question is whether Dembski’s explanatory filter/design inference is of any use in determining whether cheating has occurred in this case. The answer, given Dembski’s distinction between actual and apparent specified complexity, is clearly, “No.” In order to realize Dembski’s distinction, we must rely on information about the actual causal process, such as a video camera feed from within the Algorithm Room.

One way in which Dembski could defend his use of a distinction between actual and apparent specified complexity is to assert that algorithms which generate apparent specified complexity serve as proxies for intelligent agents. Under this view, the identification of specified complexity in the output of an algorithmic process is thus a reliable indication of the intelligent agent who instantiated the algorithm.

There are some issues which arise from such a defense, however. Dembski stipulates that algorithms can serve as conduits for prior specified complexity and add a certain amount of information as well [19, p. 160]. Dembski has not produced a demonstration that anything beyond, say, cosmological constant tuning might be needed as the sole role for a distantly removed intelligent designer, perhaps one who did part of what Dembski describes: “The fine-tuning of the universe and irreducibly complex biochemical systems are instances of specified complexity, and signal information inputted into the universe by God at its creation.” [17, p. 233] Even if we accept this nearly Deist construction of events minus the gratuitous inclusion of irreducibly complex systems, natural processes could then have provided the means by which such an initial influx of specified complexity at the beginning of the universe becomes the basis of the whole of biological diversity. Dembski’s framework is, as we have pointed out here, incapable of distinguishing between such a scenario and one requiring more recent interventions by an intelligent agent. In this regard, specified complexity becomes something like the cosmic microwave background radiation: it can be detected almost anywhere one looks.

This also bears upon the discussion in our section on animal intelligence. Dembski’s apparatus does not distinguish the initial appearance of specified complexity from subsequent appearances due to algorithmic processes. We find Dembski’s curious emphasis on the phrase “generating specified complexity” to be out of place. It seems obvious that Dembski intends this to specify those cases where the specified complexity is somehow novel, as when Dembski appeals to the characteristic of intelligent agents as “innovators” to evade a criticism [19, p. 109]. But Dembski’s apparatus has no traction in this regard. The specified complexity found via the explanatory filter/design inference can no more be credited to a supposed proximal intelligent agent than it can an algorithm, if we apply Dembski’s own criteria evenly and fairly. We must always suspect that the specified complexity had its source at some step further removed than the event under analysis. Thus Dembski’s peculiar mode of induction itself is undercut by his insistence that a distinction exists between actual and apparent specified complexity, since he premises that induction upon what he assumes, but cannot prove, are cases of known intelligent agency. This resolves to just a rather complex way of begging the question.

6 Complex specified information

We now turn to addressing Dembski’s mathematical framework for inferring design. He actually seems to have two different frameworks.

The generic chance elimination argument (GCEA) requires the elimination of *all* relevant chance hypotheses. If all such hypotheses are eliminated, Dembski concludes design is the explanation for the event in question. The weakness of this approach seems self-evident, because this method will consistently assign design to events whose exact causal history is

obscure—precisely the events Dembski is most interested in. As this framework has been adequately criticized elsewhere [23, 94], we do not examine it in detail.

Although Dembski spends significant space discussing the GCEA, in practice he rarely uses it. Instead, he employs an alternate approach. This method is a shortcut version of the GCEA, based on eliminating a single chance hypothesis, usually evaluated relative to a uniform distribution. We might call it the “sloppy chance elimination argument”, or SCEA.

It appears that both approaches can detect a certain *property* of events, called “specified complexity” or “complex specified information” (CSI). Dembski insists that “if there is a way to detect design, specified complexity is it.” [19, p. 116] While the GCEA is a statistical procedure that must be followed, CSI seems to be a property that inheres in the *record* of the event in question.

Dembski conflates the *procedure* to eliminate hypotheses with the *property* of CSI [19, p. 73] with no significant explanation. It seems to us a major jump in reasoning to go from eliminating hypotheses about an event E to the positing of a property, CSI, that inheres in E .

Then again, the choice of the term “complex specified information” is itself extremely problematic, since for Dembski “complex” means neither “complicated” as in ordinary speech, nor “high Kolmogorov complexity” as understood by algorithmic information theorists. Instead, Dembski uses “complex” as a synonym for “improbable”.

Not all commentators on Dembski’s work have appreciated that CSI is *not* information in the accepted senses of the word as used by information theorists; in particular, it is neither Shannon’s entropy, surprisal, or Kolmogorov complexity. Although Dembski claims that CSI “is increasingly coming to be regarded as a reliable marker of purpose, intelligence, and design” [19, p. xii], it has not been defined formally in any reputable peer-reviewed mathematical journal, nor (to the best of our knowledge) adopted by any researcher in information theory. A 2002 search of MathSciNet, the on-line version of the review journal *Mathematical Reviews*, turned up 0 papers using any of the terms “CSI”, “complex specified information”, or “specified complexity” in Dembski’s sense. (The term “CSI” does appear, but as an abbreviation for unrelated concepts such as “contrast source inversion”, “conditional symmetric instability”, “conditional statistical independence”, and “channel state inversion”.)⁵

Despite his insistence that his “program has a rigorous information-theoretic underpinning” [19, p. 371], CSI is used inconsistently in Dembski’s own work. Sometimes CSI is a quantity that one can measure in bits: “the CSI of a flagellum far exceeds 500 bits” [17, p. 178]. Other times, CSI is treated as a threshold phenomenon: something either “exhibits” CSI or doesn’t: “The Law of Conservation of Information says that if X exhibits CSI, then so does Y” [19, p. 163]. Sometimes numbers or bit strings “constitute” CSI [17, p. 159]; other times CSI refers to a pair (T, E) where E is an observed event and T is a pattern to which E

⁵A recent paper by creationist Stephen C. Meyer [67] states

Systems that are characterized by both specificity and complexity (what information theorists call “specified complexity”) have “information content”.

The second author was curious about the plural use of “information theorists” and at a recent conference asked Meyer, what information theorists use the term “specified complexity”? He then admitted that he knew no one but Dembski.

conforms [19, p. 141]. Sometimes CSI refers to specified events of probability $< 10^{-150}$; other times it can be contained in “the sixteen-digit number on your VISA card” or “even your phone number” [17, p. 159]. Sometimes CSI is treated as if, like Kolmogorov complexity, it is a property independent of the observer — this is the case in a faulty mathematical “proof” that functions cannot generate CSI [19, p. 153]. Other times it is made clear that computing CSI crucially depends on the background knowledge of the observer. Sometimes CSI inheres in a string regardless of its causal history (this seems always to be the case in natural language utterances); other times the causal history is essential to judging whether or not a string has CSI. CSI is indeed a measure with remarkably fluid properties! Like Blondlot’s N-rays, however, the existence of CSI seems clear only to its discoverer.

Here is a brief catalogue of some of the things Dembski has claimed exhibit CSI or “specified complexity”:

1. 16-digit numbers on VISA cards, [17, p. 159]
2. phone numbers, [17, p. 159]
3. “all the numbers on our bills, credit slips and purchase orders”, [17, p. 160]
4. the “sequence corresponding to a Shakespearean sonnet”, [19, p. xiii]
5. Arthur Rubinstein’s performance of Liszt’s “Hungarian Rhapsody”, [19, p. 95]
6. “Most human artifacts, from Shakespearean sonnets to Durer woodcuts to Cray supercomputers”, [19, p. 207]
7. scrabble pieces spelling words, [19, pp. 172–173]
8. DNA, [19, pp. 151]
9. error-counting function in an evolution simulation, [19, p. 217]
10. a “fitness measure that gauges degree of catalytic function” [19, p. 221]
11. the “fitness function that prescribes optimal antenna performance” [19, p. 221]
12. “coordination of local fitness functions”, [19, p. 222]
13. what “anthropic principles” explain in fine-tuning arguments [19, p. 144]
14. “fine-tuning of cosmological constants” [19, p. xiii]
15. what David Bohm’s “quantum potentials” extract in the way of “active information” [19, p. 144]
16. “the key feature of life that needs to be explained” [19, p. 180]

Dembski also identifies CSI or “specified complexity” with similarly-worded concepts in the literature. But these identifications are little more than equivocation. For example, Dembski quotes Paul Davies’ book, *The Fifth Miracle*, where Davies uses the term “specified complexity”, and strongly implies that Davies’ use of the term is the same as his own [19, p. 180]. This is simply false. For Davies, the term “complexity” means high Kolmogorov complexity, and has nothing to do with improbability. In contrast Dembski himself associates CSI with *low* Kolmogorov complexity:

(Note that in algorithmic information theory, “highly compressible” is synonymous with “low Kolmogorov complexity”; see the Appendix.) Therefore Dembski’s and Davies’ use of “specified complexity” are incompatible, and it is nonsensical to equate them.

- 17. The record

of political parties chosen by election official Nicholas Caputo to head the ballot in Essex County, New Jersey (D = Democrat; R = Republican) [19, pp. 55–58];

representing a variation on a fictional radio signal received from extraterrestrials in the movie *Contact* [19, pp. 6–9; 143–144]; ⁶

- 20. The flagellum of *Escherichia coli*. [19, §5.10].

16

The number of unsupported examples Dembski asserts is much larger than the number of putatively supported examples. Further, we have critiques of the arguments Dembski makes for each of these examples. We examined the Caputo example, #17, and the SETI primes sequence, #18, in Section 4. We continue with the SETI example (#18) in Section 7, and treat the weasel example (#19) in Sections 7 and 10, and the flagellum example (#20) in Section 11. However, we now make one remark about claim #17.

As we have remarked previously, sometimes CSI is treated as if it inheres in the record of events, independent of their causal history. We would like to point out that a record of events isomorphic to $\mathbf{c}$ can be obtained from any number of infrequent natural events. For example, such a record of events might correspond to

- records of whether or not there was an earthquake above 6 on the Richter scale in California on consecutive days (D = no earthquake; R = earthquake);
- records of whether or not overnight temperatures dipped below freezing in Tucson, Arizona on consecutive days (D = above freezing; R = below);
- records of whether or not Venus transited the sun in consecutive years (D = no transit; R = transit).

If Dembski wishes to infer intelligent design from the Caputo *sequence* alone, independent of context, then it seems to us that to be consistent he must also infer intelligent design for the three examples above.

7 Information, complexity, probability

For Dembski, the terms “complexity”, “information” and “improbability” are all essentially synonymous. Drawing his inspiration from Shannon’s entropy, Dembski defines the information contained in an event of probability p to be $-\log_2 p$, and measures it in bits.

It is important to note that Dembski’s somewhat idiosyncratic definition of “complexity” is often at odds with the standard definition as used by algorithmic information theorists. For Dembski the string

1111111111111111111111111111011111111111111111111.

if drawn uniformly at random from the space of all length-41 strings, has probability 2^{-41} and hence is “complex” (at least with respect to a “local probability bound”), whereas for the algorithmic information theorist, such a string is *not* complex because it has a very short description. (See the Appendix for an introduction to algorithmic information theory.)

Even if we accept equating “complexity” with “improbability”, we must ask, probability with respect to what distribution? Events do not typically come with probability spaces already attached, and this is even more the case for the singular events Dembski is interested in studying. Unfortunately, Dembski is quite inconsistent in this regard. Sometimes he computes a probability based on a known or hypothesized causal history of the event; we call this the *causal-history-based interpretation*. Sometimes the causal history is ignored

entirely, and probability is computed with respect to a uniform distribution. We call this the *uniform probability interpretation*.

Dembski's choice of interpretation seems to depend on the nature of the event in question. If the event involves intelligent agency, then he typically chooses the uniform probability interpretation. This can be seen, for example, in his discussion of archery. To compute the probability that an arrow will hit a prespecified target on a wall, he says "probability corresponds to the size of the target in relation to the size of the wall" [19, p. 10], which seems to imply a uniform distribution. Yet arrows fired at a target will almost certainly conform to a normal distribution.

If, on the other hand, the event does not involve intelligent agency, Dembski typically chooses a probability based on the causal history of the event. For example, in his discussion of the generation of the protein URF13, some aspects of causal history *are* taken into account: "First off, there is no reason to think that non-protein-coding gene segments are themselves truly random — as noted above, T-urf 13, which is composed of such segments, is homologous to ribosomal RNA. So it is not as though these segments were produced by sampling an urn filled with loosely mixed nucleic acids. What's more, it is not clear that the recombination is itself truly random." [19, p. 219] Since much of Dembski's argument involves computation and comparison of probabilities (or "information"), this lack of consistency is troubling and unexplained.

This inconsistent use of two approaches can be seen even in Dembski's discussion of a *single* example, his analysis of a version of Dawkins' METHINKS IT IS LIKE A WEASEL program.⁷ In this program, Dawkins shows how a simple computer simulation of mutation and natural selection can, starting with an initially random length-28 sequence of capital letters and spaces, quickly converge on a target sentence taken from *Hamlet*. In one passage of *No Free Lunch*, Dembski writes:

Complexity and probability therefore vary inversely — the greater the complexity, the smaller the probability. It follows that Dawkins's evolutionary algorithm, by vastly increasing the probability of getting the target sequence, vastly decreases the complexity inherent in that sequence. As the sole possibility that Dawkins's evolutionary algorithm can attain, the target sequence in fact has minimal complexity (i.e., the probability is 1 and the complexity, as measured by the usual information measure is 0). Evolutionary algorithms are therefore incapable of generating true complexity. And since they cannot generate true complexity, they cannot generate true specified complexity either. [19, p. 183]

Here Dembski seems to be arguing that we should take into account how the phrase METHINKS IT IS LIKE A WEASEL is generated when computing its "complexity" or the amount of "information" it contains. Since the program that generates the phrase does so with probability 1, the complexity of the phrase is $-\log_2 1$, or 0.

⁷Dembski characterizes Dawkins's "weasel" program as having three steps. The second and third steps which Dembski gives appear nowhere in Dawkins's text and are Dembski's own inventions, upon which he bases a number of criticisms. Dembski proposes a "more realistic" variant later, which is notable for coming much closer to an accurate description of Dawkins's "weasel" program than the one Dembski originally gave. The first author informed Dembski of this problem in October, 2000.

But in other passages of *No Free Lunch*, Dembski seems to abandon this viewpoint. Writing about another variant of Dawkins’ program, he says

...the phase space consists of all sequences 28 characters in length comprising upper case Roman letters and spaces (spaces being represented by bullets). A uniform probability on this space assigns equal probability to each of these sequences—the probability value is approximately 1 in 10^{40} and signals a highly improbable state of affairs. It is this improbability that corresponds to the complexity of the target sequence and which by its explicit identification specifies the sequence and thus renders it an instance of specified complexity (though as pointed out in section 4.1, we are being somewhat loose in this example about the level of complexity required for specified complexity—technically the level of complexity should correspond to the universality probability bound of 1 in 10^{150}). [19, pp. 188–189]

Here the choice of uniform probability is explicit.

Later, he says

It would seem, then, that *E* has generated specified complexity after all. To be sure, not in the sense of generating a target sequence that is inherently improbable for the algorithm (as with Dawkins’s original example, the evolutionary algorithm here converges to the target sequence with probability 1). Nonetheless, with respect to the original uniform probability on the phase space, which assigned to each sequence a probability of around 1 in 10^{40} , *E* appears to have done just that, to wit, generate a highly improbable specified event, or what we are calling specified complexity. [19, p. 194]

In both of these latter quotations, Dembski seems to be arguing that the causal history that produced the phrase METHINKS IT IS LIKE A WEASEL should be ignored; instead we should compute the information contained the result based on a *uniform distribution* on all strings of length 28 over an alphabet of size 27 (note that $27^{28} \doteq 1.197 \times 10^{40}$).

Sometimes the uniform probability interpretation is applied even when a frequentist approach is strongly suggested. For example, when discussing the SETI primes string

$$\mathbf{t} = 1101110111110111111110 \overbrace{111 \dots 1}^{11} 0 \overbrace{111 \dots 1}^{13} 0 \dots \overbrace{111 \dots 1}^{89} 0 \overbrace{111 \dots 1}^{73},$$

Dembski claims its probability is “1 in 2^{1000} ” [19, p. 144], a claim which is viable only under the uniform probability interpretation. But, viewing only the singular instance *t*, there are in fact many possibilities:

- (a) both 0 and 1 are emitted with probability 1/2;
- (b) 1 is emitted with probability .977 and 0 is emitted with probability .023;
- (c) the emitted bits correspond to the unary encodings of 24 numbers between 1 and 100 chosen randomly with replacement;

- (d) the emitted bits correspond to the unary encodings of 24 primes between 1 and 100 chosen randomly with replacement;

(Note: the probabilities in (b) and the choice of the number 24 in (c) and (d) reflect the actual frequencies of the symbols and the number of blocks of 1's in the string as actually printed in Dembski's book; see Section 8.)

We do not see how, in the absence of more information, to distinguish between these possibilities and dozens of others. And the choice is crucial. A purely frequentist approach, as in (b), results in a markedly different probability estimate from (a) — the probability of $\mathbf{t}$ increases from $2^{-1000} \doteq 9.33 \times 10^{-302}$ to

$$.977^{977}.023^{23} \doteq 2.8 \times 10^{-48}.$$

This latter probability, although small, is significantly larger than Dembski's "universal probability bound" of 10^{-150} and would presumably not lead to a design inference. An approach such as (d) gives an even higher probability of $25^{-24} \doteq 2.8 \times 10^{-34}$.

Clearly if Dembski gets to choose whether to apply the causal-history-based interpretation or uniform probability interpretation, as he wishes, little consistency can be expected in his calculations. Furthermore, each of the two approaches has significant difficulties for Dembski's program. The causal-history-based interpretation is the only one that is mathematically tenable; its probability estimates are necessarily based on a thorough understanding of the origin of the event in question. But this very fact makes it essentially *inapplicable* to the kinds of events Dembski wishes to study, which are events where "a detailed causal history is lacking". [19, p. xi] We expand on this in Section 9.

The uniform probability interpretation is, at first glance, easier to apply, and may be viewed as a form of the classical Principle of Indifference. But this principle has long been known to be quite problematical; as Keynes has remarked, "This rule, as it stands, may lead to paradoxical and even contradictory conclusions." [50, p. 42]. We will see in Section 9 that the uniform probability interpretation is incompatible with Dembski's "Law of Conservation of Information".

Further, even the uniform probability interpretation entails subtle choices, such as (when dealing with strings of symbols) the size of the underlying alphabet and appropriate length. If we encounter a string of the form

$$\overbrace{000000000 \dots 0}^{1000}$$

should we regard it as chosen from the alphabet $\Sigma = \{0\}$ or $\Sigma = \{0, 1\}$? Should we regard it as chosen from the space of all strings of length 1000, or all strings of length ≤ 1000 ? Dembski's advice [19, §3.3] is singularly unhelpful here; he says the choice of distribution depends on our "context of inquiry" and suggests "erring on the side of abundance in assigning possibilities to a reference class". But following this advice means we are susceptible to dramatic *overinflation* of our estimate of the amount of information contained in a target. For an example of this, see our discussion of the information content of Dawkins' fitness function in Section 9.

We note in passing that the uniform probability distribution is not the only possible choice in the absence of information. For example, the so-called “universal probability distribution” assigns the probability $2^{-K(x)}$ to an event x represented as a binary string [51]. (Here $K()$ is a prefix-free variant of Kolmogorov complexity; see the first footnote of the Appendix.) Under this choice of distribution, the kinds of events Dembski views as evidence of design, such as the Caputo sequence and the SETI primes sequence, actually occur with much *higher* probability than randomly-chosen strings. The use of this distribution would tend to undermine Dembski’s program.

Because Dembski offers no coherent approach to his choice of probability distributions, we conclude that Dembski’s approach to complexity through probability is very seriously flawed, and no simple repair is possible.

8 Specification

The second ingredient of CSI is specification. By “specification” of an event E , Dembski roughly means a pattern to which E conforms. Furthermore, Dembski demands that the pattern, in some sense, be given *independently* of E .

In its most recent incarnation, specification is formally defined as follows [19, p. 62–63]. An intelligent agent A witnesses an event E and assigns it to some reference class of events Ω . The agent then chooses from its background knowledge K' a set of items of background knowledge K such that K “explicitly and univocally” identifies a rejection function $f : \Omega \rightarrow \mathbb{R}$. Then a target T is defined by either $T = \{\omega \in \Omega : f(\omega) \geq \gamma\}$ or $T = \{\omega \in \Omega : f(\omega) \leq \delta\}$ for some given real numbers γ, δ . If K is “epistemically independent” of E (by this Dembski means that $\mathbf{P}(E|H \& K) = \mathbf{P}(E|H)$), then T is said to be “detachable” from E . (Here H is a hypothesis that E is due to chance.) Finally, if $E \subseteq T$, then T is a specification for E and E is said to be “specified”.

It is important to note that in the generic chance elimination argument, additional considerations such as “specificational resources” and “replicational resources” come into play. However, these considerations seem to be discarded in the definition of CSI.

Dembski’s account of specification has evolved over time. His original definition in *The Design Inference* included a demand that T further be *tractable*, in the sense that A can formulate T within certain constraints on its resources, such as time. This condition is dropped in *No Free Lunch* (though it now appears in Dembski’s definition of his GCEA). Further, the original definition did not restrict T to be of the form $\{\omega \in \Omega : f(\omega) \geq \gamma\}$ or $T = \{\omega \in \Omega : f(\omega) \leq \delta\}$. In this article, however, we will focus on Dembski’s most recent account, as summarized above.

We find Dembski’s account of specification incoherent. Briefly, here are our objections. First, we contend that Dembski has not adequately distinguished between legitimate and illegitimate specifications (which he calls fabrications). Second, Dembski’s notion of specification is too vague. Third, Dembski’s discussion of the generation of the target T and its independence of the event E is problematic.

Now let us look at each of these objections in more detail. When does a specification become illegitimate? To illustrate this objection, consider Dembski’s SETI primes sequence

discussed in section 6. As Dembski describes it, this sequence is of the form

$$\mathbf{t} = 110111011111011111110 \overbrace{111 \dots 1}^{11} 0 \overbrace{111 \dots 1}^{13} 0 \dots \overbrace{111 \dots 1}^{89} 0 \overbrace{111 \dots 1}^{73},$$

which encodes “the prime numbers from 2 to 89 along with some filler at the end” [19, p. 144] to make the length exactly 1000. According to Dembski, this sequence is specified, although he does not actually produce a specification. (What is f , the rejection function? What is R , the rejection region?) And when we try to create a specification, we immediately run into difficulty. What item or items of background knowledge create a legitimate specification (and not a fabrication) for $\mathbf{t}$? Our background knowledge may include prime numbers, the notion of a unary encoding, and the notion of arranging elements of a sequence in increasing order, but it is hard to see that this background knowledge “explicitly and univocally” identifies an appropriate rejection function f . After all, why stop at the prime 89? Why a filler at the end containing 73 1’s? (We suppose the notion of powers of 10 might be background knowledge, but why 1000 as opposed to 100 or 10,000?) We are leading to the following question: how contrived can a specification be and yet remain a specification? Dembski is singularly unhelpful here.

To elaborate on this idea, presumably our background knowledge includes such information as

- Mars has 2 moons;
- water comes in 3 forms, as ice, liquid, or vapor;
- 5 is the smallest odd prime congruent to 1 (mod 4);
- There were 7 wonders of the ancient world;
- In Mark 16:14 Jesus is said to appear to 11 disciples as they are eating;
- 13 is a famous unlucky number;
- 17 is a Fermat prime of the form $2^{2^2} + 1$;

as well as the notion of prime factorization. All of these are clearly epistemically independent of most events (here represented as integers). But, thanks to the unique factorization theorem, we can choose an appropriate subset of our total background knowledge in order to identify 28.7% of all integers between 1 and 1000.⁸ For example, the number 143 is the product of the number of disciples Jesus appeared to, and a famous unlucky number. At what point does specification become nothing more than a form of Kabbalistic numerology? Dembski simply does not address this issue.

⁸As the range of integers we want to identify increases, of course, more primes are needed. We are not proposing this method as a serious way of identifying all large numbers, as a classical theorem on the distribution of prime factors implies that about 6×10^{73} primes would be needed to have about a 30% chance of identifying a randomly-chosen 150-digit number as a product of those primes.

To see this objection in another way, assume we have a specification for a string, perhaps something like “a string of length 41 over the alphabet $\{\mathbf{D}, \mathbf{R}\}$, containing at most one $\mathbf{R}$ ”; this is apparently a valid specification for the Caputo string

`c = DDDDDDDDDDDDDDDDDDDDDDDDR,DDDDDDDDDDDDDDDDDDDDDD.`

Now suppose we witness Mr. Caputo produce yet another choice to head the ballot. If his choice is **D**, it is easy to produce a new specification by changing “41” to “42”. If his choice is **R**, it is easy to produce a new specification by changing “one” to “two”.⁹ But if this is the case, what prevents us from extending the process indefinitely? And if we can extend the process indefinitely, we can produce a specification for *any* string of which **c** is a prefix, a result hardly likely to increase our confidence in specification.

More precisely, suppose we are witnessing a series of events over time. Let $E(t)$ be the record of such a series at time t , viewed as a bit string over the alphabet $\{0,1\}$ and let $T(t)$ be the corresponding target we have chosen. Now suppose we witness the next state of the event, perhaps $E(t+1) = E(t) \times \{a\}$, where $a \in \{0,1\}$. It seems to us churlish to claim that $T(t)$ is a valid specification for $E(t)$, but $T(t+1) = T(t) \times \{a\}$ is not for $E(t+1)$. And if it is valid, what prevents us from continuing this process indefinitely? What we have here, of course, is the classical heap paradox in disguise [82]. Dembski denies that this is a problem for CSI by asserting that CSI is “holistic” [19, pp. 165–166], meaning that incremental additions are not allowed. It certainly follows that adding an event to a time series requires a concomitant adjustment of the specification, but it seems unreasonable to assert that the new form of the time series cannot be found to have the CSI property on that basis alone. We propose an alternate form of specification in the Appendix which resolves this problem.

Along the same lines, we find problematic the fact that all specifications are treated as equal, independent of their length. For example, presumably both

a string containing the unary representations of the first 24 prime numbers, in increasing order, separated by 0's, and followed by enough 1's at the end as to make the string of length 1000

and

a string of 1000 1's

are equally acceptable as specifications (albeit for different strings), despite the dramatic difference in their lengths. But is it permissible for the specification to actually be *longer* than the string it describes? For example, is

A Fuegian word meaning, ‘two people looking at each other, without speaking, each hoping that the other will offer to do something which both parties desire, but neither are willing to initiate’

⁹It's true that this new specification increases the probability that the target is hit, but that is not relevant here; see the next paragraph.

a valid specification for the string “mamihlapinatapai”? As far as we can see, nothing in Dembski’s discussion rules out such a specification. But if we assess a cost based on the length of the specification (a theme we take up in the Appendix, Section A.1), then this problem disappears.

We also believe Dembski’s current notion of specification is too vague to be useful. More precisely, Dembski’s notion is sufficiently vague that with hand-waving he can apply it to the cases he is really interested in with little or no formal verification.

According to its formal definition, a specification is supposed to be a rejection region R of the form $\{\omega \in \Omega : f(\omega) \geq \gamma\}$ or $\{\omega \in \Omega : f(\omega) \leq \delta\}$ for an appropriate choice of a rejection function f and real numbers γ, δ . Now consider Dembski’s discussion of the “specification” of the flagellum of *Escherichia coli*: “...in the case of the bacterial flagellum, humans developed outboard rotary motors well before they figured out that the flagellum was such a machine.” We have no objection to natural language specifications *per se*, provided there is some evident way to translate them to Dembski’s formal framework. But what, precisely, is the space of events Ω here? And what is the *precise* choice of the rejection function f and the rejection region R ? Dembski does not supply them. Instead he says, “At any rate, no biologist I know questions whether the functional systems that arise in biology are specified.” That may be, but the question is not, “Are such systems specified?”, but rather, “Are the systems specified in the precise technical sense that Dembski requires?” Since Dembski himself has not produced such a specification, it is premature to answer affirmatively.

Natural language specification without restriction, as Dembski tacitly permits, seems problematic. For one thing, it results in the Berry paradox [81]. Less trivially, relying on natural language specifications may be an easy way of arriving at a design inference without doing the hard work that could justify that conclusion. While Dembski does assert that a specification should not be tailored to an event being analyzed, his examples do not show that this condition is specifically considered in practice [19, p. 54].

Third, we find Dembski’s account of how the pattern is generated problematic. He says “For detachability to hold, an item of background knowledge must enable us to identify the pattern to which an event conforms, yet without recourse to the actual event.” [19, p. 18]. This is a strangely-worded requirement. For how could anyone verify that the event actually *does* conform to the pattern, without actually examining every bit of the event in question? To illustrate this example, let us return to the sequence $\mathbf{t}$ mentioned above.

Dembski says $\mathbf{t}$ is specified. Let us now restate his specification as S = “a string containing the unary representations of the first 24 prime numbers, in increasing order, separated by 0’s, and followed by enough 1’s at the end as to make the string of length 1000”. Presumably Dembski believes it self-evident that S could enable us to identify $\mathbf{t}$ “without recourse to the actual event”. But we cannot, for in fact, S is *not* a specification of the actual printed sequence! A careful inspection of the string presented on pp. 143–144 of *No Free Lunch* reveals that it is indeed of length 1000, but omits the unary representation of the prime 59. In other words, the string Dembski *actually* presents is

$$\mathbf{t}' = 1101110111110111111110 \overbrace{111 \dots 1}^{11} 0 \overbrace{111 \dots 1}^{13} 0 \dots \overbrace{111 \dots 1}^{53} 0 \overbrace{111 \dots 1}^{61} 0 \dots \overbrace{111 \dots 1}^{89} 0 \overbrace{111 \dots 1}^{73},$$

So in fact our proposed specification S does *not* entail $\mathbf{t}'$, but instead $\mathbf{t}$, and any pretense

that we could have identified S without explicit recourse to $\mathbf{t}'$ vanishes.

We conclude that Dembski’s account of specification is severely flawed. In Section A.1 of the appendix we provide an alternate account that repairs all three of the problems we have identified.

9 The Law of Conservation of Information

Dembski makes many grandiose claims, but perhaps the most grandiose of all concerns his so-called “Law of Conservation of Information” (LCI) which allegedly “has profound implications for science” [19, p. 163]. One version of LCI states that CSI cannot be generated by natural causes; another states that neither functions nor random chance can generate CSI. We will see that there is simply no reason to accept Dembski’s “Law”, and that his justification is fatally flawed in several respects.

Furthermore, Dembski uses equivocation to suggest that his version of LCI is compatible with others in the literature. In the context of a discussion on Shannon information, Dembski notes that if an event B is obtained from an event A via a deterministic algorithm, then $P(A \& B) = P(A)$, where P is probability [19, p. 129]. He then goes on to say “This is an instance of what Peter Medawar calls the Law of Conservation of Information” and cites Medawar’s book, *The Limits of Science*. Dembski repeats this claim when he discusses his own “Law of Conservation of Information” [19, p. 159]. But is Medawar’s law the same as Dembski’s, or even comparable?

No. First of all, Medawar’s remarks do not constitute a formal claim, since they appeared in a popular book without proof or detailed justification. In fact, Medawar acknowledges [66, p. 79], “I attempt no demonstration of the validity of this law other than to challenge anyone to find an exception to it — to find a logical operation that will add to the information content of any utterance whatsoever.”

Second, Medawar is concerned with the amount of information in deductions from axioms in a formal system, as opposed to that in the axioms themselves. He does not formally define exactly what he means by information, but there is no mention of probabilities or the name Shannon. Certainly there is no reason to think that Medawar’s “information” has anything to do with CSI.¹⁰

Let us now turn to Dembski’s claim that functions cannot generate CSI. More precisely, Dembski claims that given CSI $j = (T_1, E_1)$, based on a “reference class of possibilities Ω_1 ”, and a function $f : \Omega_0 \rightarrow \Omega_1$ with $f(i) = j$ for some $i = (T_0, E_0)$, then i is “itself CSI with the degree of complexity in both being identical”. Notice that Dembski makes no restrictions on f at all; it could be known to the agent who observes j , or not known. If the domain of f is strings of symbols, it could map strings of symbols to strings of the same length, or longer or shorter ones. It could be computable or noncomputable.

But Dembski’s “proof” of this claim, given on pages 152–154 of *No Free Lunch*, is flawed in several ways. For the purposes of our discussion, let us restrict ourselves to the case where

¹⁰Medawar’s law, by the way, can be made rigorous, but in the context of *Kolmogorov* information, not Shannon information or Dembski’s CSI; see [10]. As we have already seen above in Section 6 Dembski’s CSI and Kolmogorov complexity, if related at all, are related in an inverse sense.

$\Omega_0 \subseteq \Sigma^*$ and $\Omega_1 \subseteq \Delta^*$, where Σ and Δ are finite alphabets. By this we mean that events are strings of symbols.

First of all, let us consider the uniform probability interpretation of CSI. Dembski justifies his assertion by transforming the probability space Ω_1 by f^{-1} . This is reasonable under the causal-history-based interpretation. But under the uniform probability interpretation, we may not even know that j is formed by applying f to i . In fact, it may not even be mathematically meaningful to perform this transform, since j is being viewed as part of a larger uniform probability space, and f^{-1} may not even be defined there.

This error in reasoning can be illustrated as follows. Given a binary string x we may encode it in “pseudo-unary” as follows: append a 1 on the front of x , treat the result as a number n represented in base 2, and then write down n 1’s followed by a 0. For example, the binary string 01 would be encoded in pseudo-unary as 111110. This encoding is reversible as follows: count the number of 1’s, write the result in binary, and delete the first 1. If we let $f : \Sigma^* \rightarrow \Sigma^*$ be the mapping on binary strings giving a unary encoding, then it is easy to see that f can generate CSI. For example, suppose we consider an 10-bit binary string chosen randomly and uniformly from the space of all such strings, of cardinality 1024. The CSI in such a string is clearly at most 10 bits. Now, however, we transform this space using f . The result is a space of strings of varying length l , with $1025 \leq l \leq 2048$. If we viewed the event $f(i)$ for some i we would, under the uniform probability interpretation of CSI, interpret it as being chosen from the space of all strings of length l . But now we cannot even apply f^{-1} to any of these strings, other than $f(i)$! Furthermore, because of the simple structure of $f(i)$ (all 1’s followed by a 0), it would presumably be easily specified by a target with tiny probability. The result is that $f(i)$ would be CSI, but i would not be.

This example also suggests a paradox inherent in Dembski’s view of specification. Presumably a binary string where all but one of the symbols is 1 is always “specified”, no matter what its length is. (For example, in Dembski’s discussion of the Caputo case, there was no discussion of whether it was important that the length of that sequence (41) be specifiable with background knowledge.) Now let $\Omega_0 = \{0,1\}^n$ be the space of all binary strings of length n , let

$$\Omega_1 = \{\overbrace{111 \cdots 1}^i 0 : 2^n \leq i < 2^{n+1}\},$$

and let $f : \Omega_0 \rightarrow \Omega_1$ be the function which sends a string w to its pseudo-unary encoding

as described in the previous paragraph. Now $f(w)$ is always of the form $\overbrace{111 \cdots 1}^t 0$ for some number t , and hence is presumably always specified. Furthermore, a specification

for $\overbrace{111 \cdots 1}^t 0$ could just be the string itself, as Dembski urges us [19, p. 72] to choose our specifications as narrowly as possible. Now when we apply Dembski’s recipe for computing the rejection region on Ω_0 [19, p. 153] we discover that every element of Ω_0 is specified! We conclude that *every* binary string is specified, and specification becomes a vacuous concept. The only way out of this paradox, we believe, is to insist that not all strings of the form $\overbrace{111 \cdots 1}^t 0$ be specified, and this can be accomplished by taking the description size of the specification into account; see the Appendix.

Another error in Dembski's analysis is as follows. To obtain the detachability of $f^{-1}(T_1)$, Dembski says that " f merely [needs] to be composed with the rejection function on Ω_1 : if g is the rejection function on Ω_1 that induces the rejection region T_1 that is detachable from E_1 , then $g \circ f$, the composition of g and f , is the rejection function on Ω_0 that induces the rejection region T_0 that is detachable from E_0 ". Here Dembski seems to be forgetting that the rejection function is supposed to be "explicitly and univocally" identifiable from background knowledge K . While g is presumed identifiable in this sense relative to K , in what sense is $g \circ f$ identifiable? There is simply no reason for $g \circ f$ to be identifiable, for two reasons. First, in the uniform probability interpretation of CSI, the intelligent agent who identified g may be entirely unaware of f . Recall that Dembski's claim that functions cannot generate CSI was a universal claim about *all* functions f , not just functions specifiable by the intelligent agent's background knowledge K . Second, under both interpretations of CSI, even if the intelligent agent knows f , the composition $g \circ f$ may not be identified "explicitly and univocally" from K , since another function g' identifiable from K , when composed with f , might give a compatible rejection function for T_0 in Ω_0 .

Here is an example illustrating the former error. Suppose j is an English message of 1000 characters (English messages apparently always being specified), $f(i) = j$, and f is a mysterious decryption function which is unknown to the intelligent agent A who identified j as CSI. Perhaps f is computed by a "black box" whose workings are unknown to A , or perhaps A simply stumbles along j which was produced by f at some time in the distant past. The intelligent agent A who can identify j as CSI will be unable, given an occurrence of i , to identify *it* as CSI, since f is unknown to A . Thus, in A 's view, CSI j was actually *produced* by applying f to i . The only way out of this paradox is to *change* A 's background knowledge to include knowledge about f . But then Dembski's claim about conservation of CSI is greatly weakened, since it no longer applies to all functions, but only functions specifiable through A 's background knowledge K .

This error becomes even more important when j arises through a very long causal history, where thousands or millions of functions have been applied to produce j . It is clearly unreasonable to assume that both the initial probability distribution, which may depend on initial conditions billions of years in the past, *and* the complete causal history of transformations, be known to an intelligent agent reasoning about j .¹¹ But in applying the causal-history-based approach, it is absolutely *crucial* that every single step be known; the omission of a *single* transformation by a function f has the potential to skew the estimated probabilities in such a way that LCI no longer holds, as in our example of pseudo-unary encoding.

Finally, there is a third error in Dembski's claim about functions and CSI, which holds in both the causal-history-based interpretation and the uniform probability interpretation. On pages 154–155 of *No Free Lunch* Dembski acknowledges that his proof that functions cannot generate CSI (pages 152–154) is, in fact, not a proof at all. He forgot "the possibility that functions might add information". (Strange, we thought that was what the previous proof was intended to rule out.) To cover this possibility Dembski introduces the universal composition function U , defined by $U(i, f) = f(i) = j$. He then argues that the amount of CSI in the pair (i, f) is at least as much as j . Of course, this claim also suffers from the two

¹¹Dembski seems to admit this when he says that "...most claims are like this (i.e., they fail to induce well-defined probability distributions)..." [19, p. 106].

problems mentioned above, but now there is yet another error: Dembski does not discuss how to determine the CSI contained in f .

This is not a minor or insignificant omission. Recall that under one interpretation of LCI, f is supposed to correspond to some natural law. If f contains much CSI on its own, then by applying f we could accumulate CSI “for free”. Furthermore, since if we consider f to be chosen uniformly from a space of all possible functions with the same choice of domain and range, then the amount of CSI in f could be extraordinarily large.

For example, consider the information contained in a fitness function in Dawkins’ **METHINKS IT IS LIKE A WEASEL** example. A typical such fitness function f might map each string of length 28 into an integer between 0 and 28, measuring the number of matches between a sequence and the target. The cardinality of the space of all such fitness functions is $29^{27^{28}}$, or about $2^{5.816 \times 10^{40}}$. Dembski says “To say that E has generated specified complexity within the original phase space is therefore really to say that E has borrowed specified complexity from a higher-order phase space, namely, the phase space of fitness functions.” [19, p. 195] It is not clear what Dembski thinks the CSI of f is, since he never tells us explicitly. But if the model is uniform distribution over the space of all fitness functions, as his remarks suggest, we are led to conclude that the information in f is given by $-\log_2 p$, where p is the probability of choosing f uniformly from the space of all fitness functions, or 5.816×10^{40} bits. We regard this implication as evidently absurd—the fitness function can be described by a computer program of a few dozen characters—but do not know how else Dembski would evaluate the amount of information in f .

Furthermore, there remains the possibility that large amounts of CSI could be accumulated simply by iterating f a random number of times starting with a short string. If $f : \Sigma^* \rightarrow \Sigma^*$ is a length-preserving map on strings, our objection can be countered simply by considering f^n , the n -fold composition of f with itself. Then f^n would be a map with the same domain and range as f . However, our objection gathers more force if f is a length-increasing map on strings. Then the composition f^n has a larger range than f does, so the amount of CSI added by applying f could itself increase with every iteration of f .

To illustrate this possibility, consider the following procedure: starting from an empty string $x = \epsilon$, we successively choose randomly between applying the transformation $f_0(x) = 0x0$ or $f_1(x) = 1x1$. After n steps we will have produced a string y of length $2n$ that is a *palindrome*, i.e., it reads the same forward and backward. Under the uniform probability interpretation, upon viewing y we would consider it a member of the uniform probability space Σ^{2n} , where $\Sigma = \{0, 1\}$. Assuming our background knowledge contains the notion of palindromes, the specification “palindrome” identifies a target space with 2^n members, and so the probability of a randomly-chosen element of Σ^{2n} hitting the target is 2^{-n} . In other words, y contains n bits of CSI. As n increases in size, we can generate as much CSI as we like.

9.1 Natural law

Dembski states that deterministic natural laws are represented mathematically as functions [19, p. 151]. This seems inaccurate to us. Deterministic natural laws are typically represented as *equations*, not functions. (One speaks of Maxwell’s equations, Schrödinger’s equations,

etc.).

This may seem like a trivial correction, but there is actually something more profound about it. Some relatively simple equations, such as the wave equation, have the property that their solutions may be uncomputable (in a certain technical sense strongly related to Turing’s famous result about uncomputability), even for computable initial data [73]. What is the information content, in Dembski’s sense, of a non-computable function? Dembski never considers this question, but under a reasonable interpretation of his approach, the information content could be infinite. Hence, in principle, an arbitrarily large amount of CSI could be harvested from such equations through natural processes.

9.2 CSI holism

Dembski claims that “CSI is not the aggregate of its constituent items of information” [19, p. 166], but his numerical justification of that claim elevates the trivial to the profound.

Dembski compares the information content of the sentence **METHINKS IT IS LIKE A WEASEL** with the sum of the information contents of the individual words in the set

$$\{\text{METHINKS, IT, IS, LIKE, A, WEASEL}\}.$$

He argues that the sentence **METHINKS IT IS LIKE A WEASEL** is a string of length 28 over an alphabet of size 27 (26 capital letters and a space), and computes an estimate of $-\log_2 27^{-28} \doteq 133$ bits for its information content. He now claims this figure “far exceeds the complexity of ... the sum of the complexities of all items in the set”. A simple calculation shows that

$$-(\log_2 27^{-8} + \log_2 27^{-2} + \log_2 27^{-2} + \log_2 27^{-4} + \log_2 27^{-1} + \log_2 27^{-6}) \doteq 109,$$

and indeed $133 > 109$. But this inequality is not due to any mysterious property of “holism”; rather it is due to the entirely trivial fact that the phrase contains *spaces between the letters* whose information content we measured in the former case, but not the latter! Indeed, using Dembski’s method for computing information, the 5 missing spaces precisely provide the missing $-\log_2 27^{-5} \doteq 24$ bits.

In general, Dembski’s claims for CSI holism seem to us to be merely assertions to close off an avenue of critique that we broached earlier in our discussion of “specification” in Section 8. It seems trivial to use an iterative definition to describe processes which produce outputs that meet the requirements of Dembski’s CSI, and we have produced several such examples. While the specification must change with each increment (even if only to reflect the new length of the string describing the event), it seems to us that each output does indeed reflect CSI. An evident application of CSI holism is in denying that biological processes can incrementally generate CSI. This view critically depends upon a conception that the specification to be applied throughout the process remains static. If this requirement does not hold, the concept of CSI holism as a bar to incremental change vanishes.

9.3 Naturally-occurring CSI

Dembski seems to be of two minds about the possibility of CSI being generated by natural processes. For example, it would seem that the regular patterns formed by ice crystals would

constitute CSI, at least under the uniform probability interpretation. If we consider a piece of glass divided into tiny cells, and each cell either can or cannot be covered by a molecule of water with equal probability, it seems likely even in the absence of a formal calculation that the probability that the resulting figure will have the symmetry observed in ice crystals is vanishingly small. Furthermore, the symmetry seems a legitimate specification, at least as good as specifications such as “outboard rotary motor” that Dembski himself advances. Yet in addressing this claim Dembski falls back on the causal history interpretation, stating that “...such shapes form as a matter of physical necessity simply in virtue of the properties of water (the filter will assign the crystals to necessity and not to design).” [19, p. 12].

Just a paragraph later, Dembski discusses the occurrence of the Fibonacci sequence in phyllotaxis. Once again his discussion is not completely clear, but he seems to be saying (if we understand him correctly) that the occurrence of the Fibonacci sequence is, like the SETI primes sequence, a legitimate instance of CSI. However, he argues that the CSI is not *generated by the plant*, but rather is a consequence of intelligent design of the plant itself. (He compares the generation of the Fibonacci sequence here to the Fibonacci sequence produced by a program, and then asks, “whence the computer that runs the program?”) Here he seems to be invoking not the causal-history-based interpretation, but the uniform probability interpretation.

This seems inconsistent to us. If we apply the uniform probability interpretation consistently, it would seem that many natural processes, including some that are not biological, generate CSI. In a moment we will list some candidates, but first let us note that it seems unlikely Dembski will accept these as invalidating his specified complexity filter. Indeed, in response to one such challenge (the natural nuclear reactors at Oklo) he says

But suppose the Oklo reactors ended up satisfying this criterion after all. Would this vitiate the complexity-specification criterion? Not at all. At worst it would indicate that certain naturally occurring events or objects that we initially expected to involve no design actually do involve design. [19, p. 27]

In other words, Dembski’s claims are unfalsifiable. We find this good evidence that Dembski’s case for intelligent design is not a scientific one.

9.3.1 Dendrites

Dendrites are tree-like or moss-like structures that arise through crystal growth, particularly with iron or manganese oxides. If “tree-like in appearance” is a valid specification, it would seem that such structures could well constitute CSI. Indeed, their tree-like appearance often causes them to be confused with plant fossils. Dendrites were a puzzle until relatively recently [30]. Thus until recently they would have been assigned to design by Dembski’s generic chance elimination argument. Despite the currently accepted physical explanation, they might still constitute CSI under the uniform probability interpretation.

9.3.2 Atmospheric phenomena

Unusual atmospheric phenomena, such as rainbows and 8° and 22° sun halos [35] seem to fulfill the definition of CSI, at least under the uniform probability interpretation. We note

that in the Judeo-Christian tradition, the rainbow is often viewed as a message from a deity.

9.3.3 Triangular ice crystals

Under certain rare conditions snow crystals form triangular plates. Unlike the case of ordinary six-sided snowflakes, there is currently no detailed physical explanation for the formation of triangular plates.

Since there is no detailed causal hypothesis, when trying to infer whether triangular snowflakes are designed, we must fall back on a single hypothesis, the chance hypothesis. Triangular snowflakes would then seem to qualify as CSI, at least under the uniform probability interpretation. They cannot be rejected as “necessity” since no known law accounts for their formation.

Under Dembski’s design inference, we would therefore conclude that triangular plates are the product of design, but ordinary six-sided snowflakes are the product of necessity. This seems like an absurd conclusion to us.

9.3.4 Self-ordering in collections of spheres of different sizes

Under certain conditions, mixtures of small spheres of different sizes will spontaneously self-organize in mysterious ways. This would seem to be an instance of CSI, at least under the uniform probability interpretation. However, this phenomenon has recently been explained as a consequence of entropy [46, 49, 20].

9.3.5 Fairy rings

Fairy rings are circular structures formed by the growth of fungi, particularly the fungus *Marasmius oreades*. They grow outward in a circle, starving the grass above, and sometimes to a diameter of 200 meters. Under the uniform probability interpretation, fairy rings would be considered extremely improbable, and their circular shape would make them specified.

According to legends fairy rings were attributed to intelligent agency, e.g., fairies dancing in the moonlight. While Dembski’s design inference would be less specific in its identification of the identity of the designer, it would be no less erroneous.

9.3.6 Patterned ground

Repeated freeze-thaw cycles in cold environments can generate interesting circular and polygonal patterns. Under a uniform probability interpretation, such patterns would constitute CSI; yet there is now an explanation involving lateral sorting and “stone domain squeezing” [48].

10 Evolutionary computation

As mentioned in Section 2, one of Dembski’s principal claims is that evolutionary computation cannot generate CSI. This is essentially just a corollary of his Law of Conservation of Information, which as we have seen in the previous section, is invalid. More precisely,

he concedes that the “Darwinian mechanism” can generate the “appearance” of specified complexity but not “actual specified complexity” [19, p. 183].

In Chapter 4 of *No Free Lunch*, Dembski examines several examples of genetic algorithms and concludes that none of them generate CSI in his sense. He spends much of his time in this chapter doing detective work, attempting to determine if CSI has been illegitimately inserted (or in Dembski’s terms, “smuggled in”) by genetic algorithm researchers who are presumably considered intelligent agents. Not surprisingly, in each case, he finds that it has.

We remark that it is perfectly legitimate for Dembski to examine existing genetic algorithms in an attempt to see whether they can generate CSI as he understands it. However, since the researchers he discusses do not claim in their articles to have generated anything that falls under Dembski’s idiosyncratic definition of information, the imputation of dishonesty in the choice of the term “smuggling”, not to mention the patronizing analogy of correcting undergraduate mathematics assignments [19, p. 215], seems to us completely unwarranted.

Dembski considers a number of genetic algorithms: variations on Dawkins’s **METHINKS IT IS LIKE A WEASEL** example, an evolution simulation of Thomas Schneider, an algorithm of Altshuler & Linden for the design of antennas, and an evolutionary programming approach to checkers-playing by Chellapilla & Fogel.

In each case he identifies a particular place where he believes CSI has been “smuggled in”. In Dawkins’ weasel example, it is the choice of fitness function. In Schneider’s simulation, it is the error-counting function and “fine tuning” of the simulation itself. In the Altshuler-Linden algorithm, it is the “fitness function that prescribes optimal antenna performance” [19, p. 221]. In the Chellapilla-Fogel example, it is the “*coordination* of local fitness functions” (emphasis his).

It is certainly possible, a priori, that Dembski’s objections might be correct in the context of his particular measure.¹² But to show that his objections have substance, it does not suffice to simply assert that CSI has been “smuggled in”. After all, Dembski’s claim is a quantitative, not a qualitative one: the amount of CSI in the output cannot exceed that in the program and input combined. In order for his objections to be convincing, Dembski needs to perform a calculation, calculating the CSI in output, program and input, and showing that the claimed inequality holds. This he simply fails to do for each of the examples.¹³

10.1 How genetic algorithms can increase Kolmogorov complexity

Before we turn to the question of CSI, we digress to point out that under *accepted* interpretations of the term “information”, such as Kolmogorov complexity, it is easy for algorithms (including genetic algorithms) to generate outputs having more information than that contained in the input. For example, it can be shown that the Kolmogorov complexity of xx exceeds that of x for infinitely many strings x . In other words, a simple program that duplicates strings (i.e., maps x to xx) may increase Kolmogorov information.

¹²However, Schneider [85] argues they are not in the case of Shannon information.

¹³The closest he comes to a quantitative analysis is for the case of Dawkins’ weasel example, where he views the fitness function as an element of the space of all fitness functions. As we have remarked previously, this view implies an absurd estimate for the complexity of the fitness function.

Furthermore, if a single program may be applied an indefinite number of times, arbitrarily high amounts of information can be generated. When a simple duplication program, starting with the string 0, is applied repeatedly, the resulting strings eventually surpass any fixed bound in Kolmogorov complexity.

It might seem that this objection vanishes if both the information in the program and the input are taken into account. In this case, simple randomized algorithms (those with access to a source of genuinely random bits) may generate outputs with arbitrarily more Kolmogorov information than the total information contained in the algorithm and input combined. For example, consider an algorithm **PAL**, defined as follows: on input n , it generates a randomly-chosen palindrome of length $2n$ by starting with the empty string and then randomly applying either $f_0(x) = x0x$ or $f_1(x) = x1x$ n times. The resulting string will nearly always have Kolmogorov complexity close to n , but the Kolmogorov complexity of the input and program combined is bounded by $(\log_2 n) + c$ for some fixed constant c .

Thus information may well increase in genetic algorithms, in one standard sense of the word. But can it increase in the sense of Dembski's CSI? In the next section we examine this.

10.2 CSI and evolutionary computation

In this section we present two algorithms which appear to us to invalidate Dembski's claims about evolutionary computation's inability to generate CSI.

Our first algorithm solves a well-studied problem, and hence is close to the way genetic algorithms are typically used in practice.

Our algorithm is called **TSPGRID**, and takes an integer n as an input. It then solves the traveling salesman problem on a $2n \times 2n$ square grid of cities. Here the distance between any two cities is simply Euclidean distance (the ordinary distance in the plane). Since it is possible to visit all $4n^2$ cities and return to the start in a tour of cost $4n^2$, an optimal traveling salesman tour corresponds to a Hamiltonian cycle in the graph where each vertex is connected to its neighbor by a grid line.

As we have seen above in Section 9, Dembski sometimes objects that problem-solving algorithms cannot generate specified complexity because they are not contingent. In his interpretation of the word this means they produce a unique solution with probability 1. Our algorithm avoids this objection under one interpretation of specified complexity, because it chooses randomly among all the possible optimal solutions, and there are many of them.

In fact, Göbel has proved that the number of different Hamiltonian cycles on the $2n \times 2n$ grid is bounded above by $c \cdot 28^{n^2}$ and below by $c' \cdot 2.538^{n^2}$, where c, c' are constants [31]. We do not specify the details of how the Hamiltonian cycle is actually found, and in fact they are unimportant. A standard genetic algorithm could indeed be used provided that a sufficiently large set of possible solutions is generated, with each solution having roughly equal probability of being output. For the sake of ease of analysis, we assume our algorithm has the property that each solution is equally likely to occur.

Now there are $(4n^2)!$ different ways to list all $4n^2$ cities in order. But, as Göbel proved, there are at most $c \cdot 28^{n^2}$ different ways to produce a Hamiltonian cycle. Let us now compute the specified complexity using the uniform probability interpretation. The probability that

a randomly-chosen list of $4n^2$ cities forms a Hamiltonian cycle is $\leq$

$$\frac{c \cdot 28^{n^2}}{(4n^2)!}$$

and the number of bits of specified information in such a cycle is therefore $\geq$

$$-\log_2 \left(\frac{c \cdot 28^{n^2}}{(4n^2)!} \right).$$

By Stirling's approximation the number of bits of specified information is bounded below by a quantity that is approximately $8n^2 \log_2 n - 2.6n^2$.

Dembski sometimes objects that the CSI produced by algorithms is contained in the program and input. Here the input is n , which has at most $\log_2 n$ bits of information, and the algorithm is of fixed size, and can have at most c bits of information. Since for large n we have $8n^2 \log_2 n - 2.6n^2 \gg (\log_2 n) + c$, we conclude that TSPGRID has indeed generated specified complexity with respect to the uniform probability interpretation.

We now present another program which is intended to illustrate how random chance combined with deterministic laws ("functions") can generate CSI, even if the flaws of the uniform probability interpretation are avoided by using the *true* probabilities of the generated events. Strictly speaking our algorithm may not be a typical genetic algorithm, but it does combine randomness and determinism in a suggestive way.

The basic idea is as follows. We construct a randomized algorithm Q so that

- (a) any particular output of Q occurs with low probability (is "complex");
- (b) every possible output string is specified;
- (c) every output string has a *different* specification, and no two specifications intersect.

Here are the details: first, one can construct a deterministic algorithm P that on input i produces the i 'th string w (in some particular enumeration that we identify below, not necessarily the i 'th string in lexicographic order) such that the Kolmogorov complexity $C(w)$ of w is smaller than any reasonable function of the length $|w|$ of w . For example, $P(i)$ could be the i 'th string w for which $C(w) < |w|/100$, or $C(w) < \sqrt{|w|}$, or $C(w) < (\log |w|)^2$, or anything similar. This can be accomplished by a well-known technique called "dovetailing"; the details are somewhat complicated and are given in the appendix.

Now we construct a randomized algorithm Q , which on input n first generates a randomly-chosen length- n bit string t , using access to a source of genuinely random bits.¹⁴ Next, Q places a "1" in front of the base-2 representation of t , and treats the result as an integer u . (If $t = 5$, or 101 in base 2, then $u = 1101$ in base 2, or 13.) Finally, it outputs $P(u)$.

¹⁴In practice, low-quality random bits can be obtained from environmental sources (e.g., counting keystrokes or time between keystrokes) and high-quality random bits can be obtained from physical sources (e.g., counting radioactive decays). Indeed, there is even a web site, <http://www.fourmilab.ch/hotbits/>, where random bits obtained from a Geiger counter can be downloaded.

We claim Q generates specified complexity in Dembski’s sense with respect to the true probabilities of the generated events. Any string v generated by Q occurs with probability 2^{-n} . Any such string is specifiable, because there is a very short program to generate it, and Dembski says [19, p. 174, n. 23] that high compressibility constitutes a valid specification.¹⁵ If this were all, then we would be generating strings that match the target specification (“compressibility”) with probability 1. But this is not the only possible specification. For any output v , we can instead choose the specification to be the *actual program-input pair* (T, w) such that T outputs v on input w . The probability that any particular output of $Q(n)$ matches *this* specification is 2^{-n} . Hence the output represents n bits of CSI, while the total information in the program and input combined is bounded by $(\log_2 n) + c$, for some constant c .

This discussion has been rather technical, so let us rephrase it. Suppose Q on input n generates w , but we don’t know how Q works; we could, perhaps, call it “Dembski’s Black Box”. As intelligent agents we see an output v of Q and try to fit a pattern to it. If we assume that we will eventually discover a good compression for v (we could, for example, simply do some dovetailing on our own), then v is specified, and the probability that the particular specification we discover matches a random output of Q is 2^{-n} . Thus, v constitutes CSI, and so every output of Q constitutes CSI.

There is a possible objection to this construction, which runs as follows: there is no obvious way to produce the good compression for v in a reasonable length of time, and so perhaps it is contestable whether an intelligent agent could discover it with background knowledge. To counter this criticism, we use a different compression scheme. Instead of having P return the i ’th binary string with respect to low Kolmogorov complexity, we instead return the i ’th binary string with respect to some other compression scheme which is easily computable. One such encoding is run-length encoding, where a binary string is encoded by successively counting the lengths of successive blocks of identical symbols, starting with 0. For example, the run-length encoding of 0001111011111 would be (3, 4, 1, 5). We may then express this encoding in binary using a self-delimiting encoding of each of the terms. Now we could redefine P to return the i ’th bit string w for which the run-length encoding of w is shorter than $|w|/100$, or any easily computable function of $|w|$.

Now it is easy, upon seeing an output v of Q , to compute its run-length encoding and produce that as a specification. (In fact, this is similar to several of Dembski’s examples, such as the Caputo example and the SETI primes example, both of which are notable for their short run-length encodings.) In analogy with Dembski’s remarks about Kolmogorov complexity, we assume these would be valid specifications. So in this case all the specifications would be easily derivable with background knowledge, and they would all be different.

Another objection might be that the “real” specification for any observed output v should be simply “compressible” or “short run-length encoding”, and not the particular specific compression or run-length encoding we produce. But this is hardly fair in light of Dembski’s injunction to make the rejection region as small as possible. Furthermore, this objection would be like seeing both the SETI prime sequence and the Caputo sequence as outputs of some program and saying, “The specification is just that these strings have short run-length

¹⁵Dembski is quite vague about exactly what the definition of “highly compressible” is, but since we allow any computable function to be used, this is not a problem for our analysis.

encodings, so whatever is generating them is just hitting this target with probability 1". We do not believe Dembski would accept this objection for the Caputo sequence and the SETI primes sequence.

We conclude that Dembski's claims about evolutionary algorithms cannot be sustained.

10.3 Dembski and No Free Lunch

Dembski mistakenly invokes the "No Free Lunch" theorems of [95] as justification for his view that evolutionary computation cannot generate CSI. The NFL theorems compare efficiency between algorithms and characterize averaged performance over all cost functions, but Dembski's claim is not about average performance: Dembski makes the universal claim that for all evolutionary computation, no instance of CSI can be attributed to any such computation. Dembski's initial summary of the NFL theorems characterize them as establishing a problem of "regress" to a "higher-order phase space." This does not appear to reflect what is actually discussed in the papers Dembski cites on NFL. (Search for "regress" in the text of "No Free Lunch theorems for search", "No Free Lunch theorems for optimization", and "On the futility of blind search" comes up empty.) The point made by Wolpert and Macready is not that a search of the space of fitness functions must be conducted, but rather that one must examine the fitness function of interest in order to select an algorithm which will perform more efficiently given that fitness function. This is a far cry from Dembski's assertion of some "regress" in finding the source of information seen in the output of an algorithm.

10.4 The displacement problem

Dembski seeks to exclude evolutionary computation as a source of CSI by invoking what he calls "the displacement problem". In this, Dembski claims that whatever CSI exists in the output of an algorithm results from the choice of fitness function, and that this choice is made from an even larger and thus more intractable space. Upon consideration, it is apparent that we have two possible choices, both of which are fatal to Dembski's thesis: either no "displacement problem" exists, or the "displacement problem" applies to *all* possible sources of CSI, including intelligent agents. How might the "displacement problem" not be a problem? When an intelligent agent is credited with generating CSI upon producing a solution to a particular problem, it follows that an algorithm which solves that problem given the same background knowledge and input must also be credited with generating CSI. The only way out of this dilemma for Dembski would be to take the position that intelligent agents also do not generate CSI when given specific problems to solve. In other words, either the "displacement problem" is a trivial digression based upon special pleading, or it is completely generic and applies with equal force to generation of CSI by intelligent agents as well as natural processes. This latter strategy is implicit in Dembski's defense against the critiques of [23] and [14], where he appeals to intelligent agents being innovators as an escape from having to give a rational account of intelligent agency. The appeal to innovation is an appeal to irrationality, as we have previously discussed in Section 5.

11 CSI and Biology

It is no surprise to anyone who has studied the intelligent design movement that the real goal is to cast doubt on the biological theory of evolution. In *Intelligent Design*, Dembski began an attack on evolution which he continues in *No Free Lunch*. However, many of his claims appear suspect.

For example, consider Dembski's claims about DNA. He implies that DNA has CSI [19, p. 151], but this is in contradiction to his implication that CSI can be equated with highly-compressible strings [19, p. 144]. In fact, compression of DNA is a lively research area. Despite the best efforts of researchers, only minimal compression has been achieved [36, 84, 12, 56, 2, 59].

Dembski devotes many pages of *No Free Lunch* to his claim that the flagellum of *Escherichia coli* contains CSI. We have already noted in Section 8 that his treatment of specification in this case leaves much to be desired. But even if one accepts "outboard rotary motor" as a valid specification, is it true that the *E. coli* flagellum matches this specification? There are significant differences. To name a few, a human-engineered outboard rotary motor spins continuously, but the flagellum moves in jerks. An outboard rotary motor drives a propeller, but the flagellum is whip-like. No human-engineered outboard rotary motor is composed entirely of proteins, but the flagellum is.

Specification is just one half of specified complexity; Dembski must also show matching the specification is improbable and thus complex in his framework. To do so, he ignores the causal history and falls back on a uniform probability approach, calculating the probability of the flagellum's origin using a random assembly model. Biologically his calculations verge on the ridiculous, since no reputable biologist believes the flagellum arose in the manner Dembski suggests. Further, even if an *E. coli* flagellum appeared according to the chance causal hypothesis Dembski proposes, it would not establish a heritable trait of flagellar construction in the lineage of *E. coli*, and thus is under no account an evolutionary hypothesis. Dembski justifies his approach by appealing to the flagellum's "irreducible complexity", a term coined by fellow intelligent-design advocate Michael Behe. But Dembski ignores the fact that sequential evolutionary routes for the flagellum have indeed been proposed [77]. True, such routes are speculative and not as detailed as one would like. Nevertheless, they seem far more likely than Dembski's random assembly model.

Even taken as a non-evolutionary account of flagellar construction, the specifics of Dembski's approach reveal a number of problems. Dembski applies the phrase "discrete combinatorial object" to any of the biomolecular systems which have been identified by Michael Behe as having "irreducible complexity." By analogy to the Drake equation from astronomy, Dembski proposes the following equation for estimating the probability of a "discrete combinatorial object" (DCO):

$$p_{\text{dco}} = p_{\text{orig}} * p_{\text{local}} * p_{\text{config}}.$$

This should be read as meaning the probability of the DCO is the product of the probabilities of the origination of its constituent parts, the localization of those parts in one place, and the configuration of those parts into the resulting system. Dembski's calculation of p_{local} is

relatively straightforward:

$$p_{\text{local}} = (prot_{\text{sys}} * subst / prot_{\text{total}})^{prot_{\text{sys}} * copies}$$

where

- $prot_{\text{sys}}$ is the number of proteins in the system being analyzed;
- $subst$ is the number of different proteins which might provide an adequate substitute for each of the proteins in the system;
- $prot_{\text{total}}$ is the total number of different proteins available in context; and
- $copies$ is the number of copies of each protein that will be required to construct the system.

The only number that Dembski provides a citation for in this group is the one for $prot_{\text{total}}$: 4,289. The others are either unreferenced or admittedly made-up. For example, consider $subst$. The number of possible substitutions is not known, and in any case is quite likely highly variable with different proteins being examined. Dembski's equation, though, is highly sensitive to changes in this value. A change from Dembski's recommended value of 10 to a value of 11 produces a change in the probability of about eleven orders of magnitude. If the value were 22 or more, the probability resulting would rise above Dembski's universal probability bound of 10^{-150} .

If we look closely at the calculation Dembski provides for p_{local} , we note that it hides a critical assumption, that the *E. coli* cell should be considered as a grab-bag of proteins, all of them available in equal proportion at any location within the cell. That this assumption is untrue should come as no surprise to the reader.

Moving on to the other factors in Dembski's calculation, we find that variants of what Dembski calls a *perturbation probability* are used for finding both p_{orig} and p_{config} . This concept appears to be original to Dembski. A perturbation probability calculates the ratio of the number of ways that a protein or string of symbols can differ while still preserving functionality to the number of ways which it may differ while still uniquely identifying the function under consideration. This in itself is problematic, for biological proteins commonly serve two or more distinct functions. No time is wasted by Dembski in considering such empirically-verified sloppiness. Dembski's general formula for an approximation of a perturbation probability is

$$\frac{\binom{N}{qN}}{\binom{N}{rN}} (k - 1)^{(N(q-r))}$$

where N is the length of the protein or string, k is alphabet size, q is the perturbation tolerance factor, and r is the perturbation identity factor. Dembski uses the Gettysburg address as an example. If we think of the Gettysburg address as composed of capital letters, the space, and some punctuation marks, there are thirty symbols in the relevant alphabet. 1000 characters of that address could be presented with some proportion of the characters changed around, and it would still convey the meaning to a recipient. The largest proportion of changes to unchanged text which preserves the meaning corresponds to the perturbation

tolerance factor. If some of the characters were missing, a recipient would still be able to identify it as the Gettysburg address. The largest proportion of missing characters to characters present which permits accurate identification corresponds to the perturbation identity factor. Dembski provides arbitrary values of 0.1 and 0.2 for the perturbation tolerance and perturbation identity factors, respectively. These are used both for the case of the English text of the Gettysburg Address and also for the proteins of the *E. coli* flagellum.

There are two things to note about these numbers in Dembski’s calculation. The first is the complete lack of any rigorous justification for the selection of these particular values. In the case of the Gettysburg Address, Dembski completely ignores Claude Shannon’s seminal work on the redundancy of English text, which is highly relevant to the determination of these values [87]. The second is the extreme sensitivity of Dembski’s proffered equation to any change in these values. A change in either value of just one percent of its original amount causes at least two orders of magnitude difference in the calculated probability for the “Gettysburg Address” example. This indicates that for the calculation to have any meaning whatsoever, the values utilized need to be empirically determined to a high degree of precision.

Even if Dembski’s intuition concerning the values he assigned to these factors was proven to be uncannily precise, there remains an interesting observation concerning the application of a perturbation probability to the calculation of p_{orig} for a particular protein. Dembski utilizes an analogy of a supermarket stocked with plenitude of different grocery products. Each of those products, he argues, may have its own p_{orig} value [19, p. 301]. Given Dembski’s values for the perturbation tolerance and identity factors, what one finds without much difficulty is that p_{orig} for any individual protein of length ≥ 1153 is less than Dembski’s “universal small probability.” Further, any collection of proteins with a combined length ≥ 1153 also has p_{orig} less than Dembski’s “universal small probability.” Dembski elsewhere tags biological function as a sufficient stand-in for “specification.” The result is that, using Dembski’s proffered values and equations, any functional protein of length ≥ 1153 has CSI and must be considered to be “due to design.” This is already a low bar for finding CSI in biological systems, but the “universal small probability” is not in any sense a threshold. Dembski merely argues that a probability below the “universal small probability” obviates the need to justify a “local small probability.” By doing so, many shorter proteins may also be found to have CSI and be classed as “due to design.” A Dembskian designer intervening in biology would appear to be exceedingly busy over the course of life’s history.

11.1 Biology and genetic algorithms

Dembski apparently views biological evolution as a natural instantiation of a genetic algorithm. In some respects this is quite unrealistic and a reversal of the actual directionality. For example, when people design genetic algorithms they are typically constructed with the goal of solving some computationally difficult problem (e.g., traveling salesman) which has resisted other approaches. By contrast, there is no known computational problem for which biological evolution is the solution.

Another contrast is that a genetic algorithm typically has a well-defined termination point, but biological evolution is an ongoing process with no pre-set termination point.

Dembski asserts that “evolutionary algorithms” represent the mathematical underpinnings of Darwinian mechanisms of evolution [19, p. 180]. This claim is egregiously backward. A large body of scholarly work is completely ignored by Dembski in order to make this claim, including Ronald Fisher’s 1930 book, *The Genetical Theory of Natural Selection*.¹⁶ It is evolutionary computation which takes its underpinnings from the robust mathematical formulations which were worked out in the literature of evolutionary biology.

While genetic algorithms are useful as problem-solving tools, they are usually a poor computational model of evolution. However, there is a field that attempts to simulate aspects of biological evolution mathematically: the field of artificial life.

11.2 Dembski and artificial life

Artificial life attempts to model evolution not by solving a fixed computational problem, but by studying a “soup” of replicating programs which compete for a resources inside a computer’s memory. Artificial life is closer to biological evolution, since the programs have “phenotypic” effects which change through time.

The field of artificial life evidently poses a significant challenge to Dembski’s claims about the failure of algorithms to generate complexity. Indeed, artificial life researchers regularly find their simulations of evolution producing the sorts of novelties and increased complexity Dembski claims are impossible. Yet Dembski’s coverage of artificial life is limited to a few dismissive remarks. Indeed, the term “artificial life” does not even appear in the index to *No Free Lunch*.

Consider Dembski’s appraisal of of the work of artificial life researcher Tom Ray:

Thomas Ray’s Tierra simulation gave a similar result, showing how selection acting on replicators in a computational environment also tended toward simplicity rather than complexity — unless parameters were set so that selection could favor larger sized organisms (complexity here corresponding to size). [19, p. 211]

We have to wonder how carefully Dembski has read Ray’s work, because this is not the conclusion we drew from reading his papers. One of us wrote an e-mail message to Ray asking if he felt Dembski’s quote was an accurate representation of his work. Ray replied as follows:

No. I would say that in my work, there is no strong prevailing trend towards either greater or lesser complexity. Rather, some lineages increase in complexity, and others decrease. Here, complexity does not correspond to size, but rather, the intricacy of the algorithm.

Dembski also does not refer to papers that demonstrate the possibility of increased complexity over time in artificial life: see, for example [75, 76, 1, 11]. Neither does he cite the pioneering work of Koza, who showed how self-replicating programs can spontaneously arise from a “primordial ooze of primitive computational elements” [54]. Neither does he mention

¹⁶This oversight is all the more curious since Dembski borrows so heavily from Fisher’s other scholarly pursuit, statistics.

the complex adaptive behaviors evolved by Karl Sims' virtual creatures [88], or Lipson and Pollack's work [60] showing how an evolutionary approach can automatically produce electromechanical robots able to locomote on a plane. These omissions cast serious doubt on Dembski's scholarship.

After the publication of *No Free Lunch*, a paper by Lenski, Ofria, Pennock, and Adami [58] offers another reason to reject Dembski's claims. The authors show how complex functions can arise in an artificial life system, through the modification of existing functions.

12 Challenges for intelligent design advocates

Thus far, intelligent design advocates have produced many popular books, but essentially no scientific research. (See, for example, [29, 25].) Future success for the movement depends critically on some genuine achievements. In this section, we provide some challenges for intelligent design advocates.

12.1 Publish a mathematically rigorous definition of CSI

Taking into account the criticisms in this paper and elsewhere, we challenge Dembski to publish a mathematically rigorous definition of CSI and a proof of the Law of Conservation of Information in a peer-reviewed journal devoted to information theory or statistical inference.

12.2 Provide real evidence for CSI claims

We challenge Dembski to either provide a complete, detailed, and rigorous argument in support of Dembski's claim that each of the items #1–16 in Section 6 has CSI, or explicitly retract each unsupported claim. Any supporting argument should describe which of the two methods (causal-history-based or uniform probability) is used to estimate probabilities, and provide a detailed description of the appropriate probability space, the relevant background knowledge, the rejection region, and the rejection function.

12.3 Apply CSI to identify human agency where it is currently not known

Thus far CSI has only been used to assert design in two classes of phenomena: those for which human intervention is known through other means, and those for which a precise step-by-step causal history is lacking. We challenge Dembski or other intelligent design advocates to identify, through CSI, some physical artifact, currently *not* known to be the product of human design, as an artifact constructed by *humans*. After this prediction through CSI, provide confirming evidence for this conclusion, *independent* of Dembskian principles.

Along similar lines, apply CSI to identify a suspicious death, currently thought to be from natural causes, as foul play. Furthermore, also provide confirming evidence for this conclusion, *independent* of Dembskian principles.

We note that Dembski himself has stressed the importance of independent evidence [19, p. 91].

12.4 Distinguish between chance and design in archaeoastronomy

The Anasazi, or ancestral Puebloans, occupied what is now the southwestern United States from about 600 to 1300 C. E.. Several of their buildings, including those at Chaco Canyon, Hovenweep National Monument, and Chimney Rock, have been interpreted as astronomical observatories, with alignments correlated to solstices, equinoxes, lunar standstills and other astronomical events [62]. Using the techniques of *The Design Inference*, provide a rigorous mathematical analysis of the evidence, determining whether these alignments are due to chance or human design.

Similar challenges exist for the claimed astronomical alignments at Stonehenge [37, 69] and Nabta [63], and the enigmatic drawings at Nazca in southern Peru. Which of the proposed alignments were designed, and which are pure coincidence?

12.5 Apply CSI to archaeology

Another interesting question about the Anasazi is the presence of large numbers of pottery shards at certain ruins. Some archaeologists have interpreted the number of these shards as exceeding the amount that could be expected through accidents. Use CSI to determine if the pots were broken through accident, or human intent (possibly in support of some religious ritual).

Archeologists have developed methods for determining whether broken flints cracked due to human intervention or not [13]. Attempt to rederive this classification, or prove it wrong, using the methods of CSI.

Provide a useful means of applying CSI to distinguish early stone tools from rocks with random impact marks.

12.6 Provide a more detailed account of CSI in biology

Produce a workbook of examples using the explanatory filter, applied to a progressive series of biological phenomena, including allelic substitution of a point mutation. There are two issues to be addressed by this exercise. The first is that a series of fully worked-out examples will demonstrate the feasibility of applying CSI to biological problems. The second is to show that small-scale changes are assigned to “chance” and “design” only is indicated for much larger-scale changes or systems already noted as having the attribute of “irreducible complexity.” It is our expectation that application of the “explanatory filter” to a wide range of biological examples will, in fact, demonstrate that “design” will be invoked for all but a small fraction of phenomena, and that most biologists would find that many of these classifications are “false positive” attributions of “design.”

12.7 Use CSI to classify the complexity of animal communication

As mentioned above, many birds exhibit complex songs. We challenge Dembski or other design advocates to produce a rigorous account of the CSI in a variety of bird songs, producing explicit numerical estimates for the number of bits of CSI.

Similar challenges can be issued for dolphin vocalizations, as in providing a definitive test of the “signature whistle” hypothesis [6], and estimation of information of a dolphin biosonar click (to be compared to the information measure suggested by [45]).

12.8 Animal cognition

Apply CSI to resolve issues in animal cognition and language use by non-human animals. Some of these outstanding issues include studies of mirror self-recognition [27, 28] and artificial language understanding in chimpanzees [83], dolphins [40], and parrots [70]. We note the use of examples in Dembski’s work involving a laboratory rat traversing a maze as an indication of the applicability of CSI to animal cognition [16, 17, 19].

13 Conclusions

We have argued that Dembski’s justification for “intelligent design” is flawed in many respects. His concepts of complexity and information are either orthogonal or opposite to the use of these terms in the literature. His concept of specification is ill-defined. Dembski’s use of the term “complex specified information” is inconsistent, and his proof of the “Law of Conservation of Information” is flawed. Finally, his claims about the limitations of evolutionary algorithms are incorrect. We conclude that there is no reason to accept his claims. Finally, we issue several challenges to those who would continue to pursue “intelligent design” as a research paradigm.

14 Acknowledgments

We are grateful to Anna Lubiw, Ian Musgrave, John Wilkins, Erik Tellgren, and Paul Vitányi, who read a preliminary version of this paper and gave us many useful comments. We owe a large debt to Richard Wein, whose original ideas have had significant impact on our thinking. We thank Norman Levitt for suggesting the pulsars example, and Phyllis Forsyth for her examples from ancient Greece.

A Algorithmic Information Theory

In this appendix we give a brief tutorial on algorithmic information theory, provide our replacement for Dembski’s CSI, and provide the details of an algorithm mentioned in the text.

Roughly speaking, algorithmic information theory (AIT) is the study of the complexity of strings of bits (0’s and 1’s). It was invented independently by the Russian mathematician A. N. Kolmogorov [52] and the American mathematician G. Chaitin [8, 9] (while the latter was still in high school). Similar ideas had been proposed earlier by R. Solomonoff. The book by Ming Li & Paul Vitányi [61] is a very good survey of the field, although those without a strong mathematical background will find it quite challenging.

The principal tool of algorithmic information theory is Kolmogorov complexity. Roughly speaking, the Kolmogorov complexity of a string of bits x is the length of the shortest program-input combination (P, i) which will produce x when P is run on i . (By the “length” of (P, i) we mean the number of bits used to write it down.) This complexity is denoted by $C(x)$,¹⁷ and is sometimes called the *information* contained in x . Note that the running time of P does *not* figure at all into our considerations here; P could produce x in one microsecond or one millennium and $C(x)$ would be the same.

A string x has low Kolmogorov complexity if there is a short program P and a short input i such that P prints x when run on input i . For example, the bit string

111111111111111111111111111111

has low Kolmogorov complexity because it can be generated by the program “print ‘1’ n times” together with the input $n = 26$. In fact, the string

$\overbrace{111 \dots 1}^n$

has Kolmogorov complexity $\leq (\log_2 n) + c_1$ for some constant c_1 . Here c_1 is the length of the program “print 1 n times” and $\log_2 n$ is (essentially) the number of bits it takes to write down n in base 2.

Why didn’t we compute c_1 explicitly? For one thing, the number c_1 depends on the particular programming language we use to represent our program. Unfortunately, there is no natural or universally-agreed upon choice. Should we use Java, C, APL, Pascal, FORTRAN, or something else entirely? In fact, mathematicians typically use none of these, preferring a programming model called the Turing machine (named after its inventor, Alan Turing). (In this case the program P is actually an encoding of a Turing machine that can be interpreted by a so-called “universal” Turing machine.) Each programming language might result in a different value of c_1 . This is one reason why computations with Kolmogorov complexity are typically specified only up to an additive constant. Luckily, in our example, the number c_1 isn’t really important, since it is quite small compared to $\log_2 n$ when n gets very large.

Furthermore, an important result called the Invariance Theorem states, roughly speaking, that the Kolmogorov complexity of a string x relative to one programming language P_1 is equal to the complexity relative to another language P_2 , up to a fixed additive constant that depends only on P_1 and P_2 .

Kolmogorov complexity is strongly related to optimal lossless data compression. Lossless data compression may be familiar as the technology which allows you to store a large file in an encoded form that (often) takes up less space on your hard drive, using a command such as **zip**. To recover an exact copy of the original file, one can use **unzip**.

If (P, i) is the shortest program-input pair that produces x , we can think of the (P, i) as the best possible way to compress x . If we want to store x , we could store (P, i) instead, since we could always recover x by running P on i . In the case of strings containing a small

¹⁷There is another notion of Kolmogorov complexity, called prefix complexity, and denoted $K(x)$. These two notions differ only by a small amount and for the purposes of this article we can consider them the same.

amount of information, such as $\overbrace{111 \cdots 1}^n$, it evidently makes sense to store them in some compressed form, rather than writing out all those 1's.

However, not every string can be compressed. It is one of the basic theorems of the Kolmogorov theory that there is at least one string x of each length which is not compressible at all. That is, there is at least one string x such that the “compressed” representation (P, i) has at least as many bits as x itself. Such strings are termed “random”. Note that this is a *definition* of the term “random” and not a theorem. Nevertheless, a string which is random in the Kolmogorov sense possesses many of the properties we associate with being random.

Similarly, it can be shown that there are at least $2^{n-1} + 1$ strings of length n for which the optimal compression has length at least $n - 1$, at least $3 \cdot 2^{n-2} + 1$ strings of length n for which the optimal compression has length at least $n - 2$, and so on. In general there are at least $2^n - 2^{n-k} + 1$ strings of length n for which the optimal compression has length at least $n - k$. It follows from this theorem that “most” strings have relatively high Kolmogorov complexity.

We have seen that $C(x)$ can be very small for highly-patterned strings. Is there a limit on how big it can be? The answer is that we always have $C(x) \leq |x| + c_2$. Here $|x|$ is shorthand for the length of, or number of bits in, the string x . To see this, observe that every string can be “compressed” by outputting the program “**print the input**” together with the input x itself. (This may not be the optimal way to compress x , but we are just trying to find an upper bound on $C(x)$.) Hence we have $C(x) \leq |x| + c_2$, where c_2 represents the length of the program “**print the input**”.

It follows that the quantity $C(x)/|x|$ is a number between 0 and a little more than 1 which measures the complexity of the string x . The following table illustrates how strings can be classified. (For a similar illustration, see, e.g., [97].)

$C(x)/ x $ close to 0	highly compressible ordered produced by simple rules low information nonrandom infrequent a randomly-chosen string has this property with low probability
$C(x)/ x $ close to 1	highly incompressible high information not produced by simple rules random frequent a randomly-chosen string has this property with high probability

Table 1: Classification of strings by Kolmogorov complexity

There is one drawback to measuring the complexity of a specific given string using Kolmogorov theory: it is, in general, uncomputable. More precisely, we can verify the inequality

$C(x) \leq c$ for a string x simply by producing a suitable program-input pair (P, i) . However, lower bounds on $C(x)$ are, in general, very difficult to obtain. The reason is not hard, but somewhat subtle; it depends on Turing’s theory of computability. Thus, there is no computer program that, given x , will unerringly compute the exact value of $C(x)$ for all x .

This is a serious obstacle when we want to apply Kolmogorov theory to a naturally-occurring string, such as the string of bases in a molecule of DNA over the alphabet $\{a, c, t, g\}$. To get around this difficulty, we may use computable approximations to $C(x)$. For example, the length of the result after using the command `compress`, available on many Unix systems, gives an approximation to $C(x)$ using a variant of Ziv-Lempel compression [98, 93]. Other approaches include resource-bounded Kolmogorov complexity [61, Chap. 7] and automatic complexity [86].

A.1 A different account of specification

We now suggest a different account of specification for bit strings, based on Kolmogorov complexity. In our suggested replacement, a specification for a string y is simply a program-input pair (M, x) such that y is output when Turing machine M is run on input x . Now let $n = C(y)$ be the minimum number of bits needed to write down (M, x) , over all pairs (M, x) that generate y ; this is just $C(y)$, the Kolmogorov complexity of y . We now call $\max(0, |y| - n)$ the “specified anti-information” in y , and denote it by $\text{SAI}(y)$. (We call it “anti-information” because it is close to the negation of what algorithmic information theorists usually mean by information.) This definition is consonant with Dembski’s own discussion in *The Design Inference* [16, pp. 171–174].

Since Kolmogorov complexity is uncomputable, we cannot in general know the SAI of y with certainty. If, however, we take our minimum n , not over all pairs (M, x) that output y , but merely over all *currently known* such pairs, then $|y| - n$ is a good lower bound on the SAI of y . We might call this the “known specified anti-information” in y , and denote it by $\text{KSAI}(y)$. Note that if someone later discovers a shorter program-input pair that generates y , this *increases* our estimate $\text{KSAI}(y)$ of $\text{SAI}(y)$; however, discovering a *longer* program never *decreases* it.

As we see it, this account of specification has many advantages over Dembski’s. For one thing, since it is divorced from considerations of probability, we no longer have to engage in the dubious practice of determining the relevant space of events after witnessing an event. Instead, our objects of interest live naturally in the space Σ^* of all finite bit strings over the alphabet Σ . Neither are we forced to assign probabilities to our events, or engage in the pretense that our specification is somehow independent of the event y .

Further, we no longer have to argue about the distinction between “specification” and “fabrication”. Following our suggestion, *every* program that outputs y is a legitimate specification, but some are better (shorter) than others. The shorter the specification, the more SAI we can confidently assert exists in y . Also, there can be no arguments about the validity of our specifications. Anyone who is skeptical can simply run M on x and verify that it produces y .¹⁸

¹⁸If there are concerns about programs that take too long to run, one can substitute a time-bounded version of Kolmogorov complexity instead. [61, Chap. 7]

Finally, our measure shares many of the properties of Dembski's CSI. If y is a Kolmogorov-random string, then by definition $C(y) \geq y$, and so the SAI in y is 0. Hence, like Dembski's CSI, a string of bits formed by flipping a fair coin nearly always has little SAI. On the other hand, strings such as

$$\mathbf{c} = \text{DDDDDDDDDDDDDDDDDDDDDDDRDDDDDDDDDDDDDDDDDDDD}$$

corresponding to the Caputo case and

$$\mathbf{t} = 110111011111011111110 \overbrace{111 \dots 1}^{11} 0 \overbrace{111 \dots 1}^{13} 0 \dots \overbrace{111 \dots 1}^{89} 0 \overbrace{111 \dots 1}^{73},$$

corresponding to the SETI primes sequence, have high SAI compared to their lengths, because they are generated by simple programs.

We now argue that our measure of SAI is strongly related to Dembski's CSI under the uniform probability interpretation. To see this, let us assume that our event space Ω equals Σ^n , the space of all length- n strings over $\Sigma = \{0, 1\}$. Suppose the event E is specified by the target T , and furthermore our description of T is computable, in the sense that there exists a program P which takes elements F of Ω as input, and outputs 1 if $F \in T$ and 0 otherwise. Then, following Dembski, the number of bits I of CSI in E is given by

$$I = -\log_2 \left(\frac{\#T}{\#\Omega} \right) = n - \log_2(\#T), \quad (1)$$

where by $\#S$ we mean the cardinality of the set S . Now we can encode E by providing the program P , together with an index describing the order of occurrence of E in a lexicographically-ordered list of all the elements of T . It follows that

$$C(E) \leq |P| + \log_2(\#T), \quad (2)$$

where $|P|$ denotes the length of a self-delimiting encoding of the program P . Adding together (1) and (2), we get $I + C(E) \leq n + |P|$ and hence

$$I - |P| \leq \text{SAI}(E)$$

It follows that the SAI of E is an upper bound on Dembski's CSI (under the uniform probability interpretation), minus the cost associated with the description P . (Assessing a cost to the length P allows us to avoid the paradoxes discussed in Section 8.)

Note that our framework disqualifies observer-dependent specifications such as 'English sonnet', and vague specifications such as 'like an outboard rotary motor', unless they can be formulated as a legitimate program-input pair. We see this as a significant advantage, although we imagine Dembski may differ.

Also note that Dembski's "Law of Conservation of Information" fails for SAI. Indeed, it is easy to increase SAI through applying functions. We now argue that there exists a constant c such that if $\text{SAI}(y) = n$, then $\text{SAI}(yy) \geq |y| + n - c$. This follows easily from an elementary exercise in Kolmogorov complexity that $C(yy) \leq C(y) + c$ for some constant c

[61, p. 101]. Indeed, by adding the the following inequalities and equalities together, we get the desired result:

$$\begin{aligned} C(yy) &\leq C(y) + c; \\ n &= |y| - C(y); \\ 2|y| - C(yy) &= \text{SAI}(yy). \end{aligned}$$

Provided $|y| > c$, this shows that duplicating a string is guaranteed to increase SAI.¹⁹

A.2 The algorithm P

Let $h(n)$ be any computable function of n . In this section we describe of algorithm P , used in Section 10. This is a deterministic algorithm that, for each integer input n , outputs a string $x \in \{0,1\}^*$ for which for which $C(x) \leq h(|x|)$. The mapping produced by P is injective (in other words, $P(n) \neq P(m)$ if $m \neq n$).

The algorithm P works as follows. Based on some choice of computing model (e.g., Turing machines), P works with an enumeration $P_1, P_2, P_3, \dots$ of all possible programs, and another enumeration of all binary strings as $x_1, x_2, x_3, \dots$. Now P initializes an empty list L of strings. We now do the following for all $N \geq 3$ until the program halts: for every integer $i \geq 1, j \geq 1, k \geq 1$ such that $N = i + j + k$, P runs program P_i on input x_j for k steps. If P_i halts and generates a string y with $|P_i| + |x_j| \leq h(|y|)$, we compare y to see if it is already on L . If not, we append it to L . Now continue, trying the next program (or incrementing N if we are done with all the triples (i, j, k) such that $i + j + k = N$). We continue until the list L is of length n , and at this point we output the last string on the list.

References

- [1] C. Adami, C. Ofria, and T. C. Collier. Evolution of biological complexity. *Proc. Nat. Acad. Sci. U. S. A.* **97** (2000), 4463–4468.
- [2] A. Apostolico and S. Lonardi. Compression of biological sequences by greedy off-line textual substitution. In *Proc. IEEE Data Compression Conference (DCC)*, pp. 143–152, 2000.
- [3] E. R. Berlekamp, J. H. Conway, and R. K. Guy. *Winning Ways, For Your Mathematical Plays*. Academic Press, 1982.
- [4] P. Boyer. *Religion Explained*. Basic Books, 2001.

¹⁹Dembski claims “there is no more information in two copies of Shakespeare’s *Hamlet* than in a single copy. This is of course patently obvious, and any formal account of information had better agree.” [17, p. 158]; [19, p. 129] This is much too glib. We have just shown that yy nearly always contains more SAI than y . Similarly, Kolmogorov complexity itself is a formal account of information, and it can be shown that there exist infinitely many strings y such that $C(yy) > C(y)$. For other formal accounts of information where yy has more information than y , see Vitányi’s quantum information theory [91] and the automatic complexity of Shallit & Wang [86].

- [5] J. Byl. Self-reproduction in small cellular automata. *Physica D* **34** (1989), 295–299.
- [6] M. C. Caldwell, D. K. Caldwell, and T. L. Tyack. Review of the signature-whistle hypothesis for the Atlantic Bottlenose Dolphin. In S. Leatherwood and R. R. Reeves, editors, *The Bottlenose Dolphin*, pp. 199–234. Academic Press, 1990.
- [7] A. C. Catania and D. Cutts. Experimental control of superstitious responding in humans. *J. Experimental Analysis of Behavior* **6** (1963), 203–208.
- [8] G. J. Chaitin. On the length of programs for computing finite binary sequences. *J. Assoc. Comput. Mach.* **13** (1966), 547–569.
- [9] G. J. Chaitin. On the length of programs for computing finite binary sequences: statistical considerations. *J. Assoc. Comput. Mach.* **16** (1969), 145–159.
- [10] G. Chaitin. Information-theoretic limitations of formal systems. *J. Assoc. Comput. Mach.* **21** (1974), 403–424.
- [11] A. Channon. Passing the ALife test: activity statistics classify evolution in Geb as unbounded. In J. Kelemen and P. Sosík, editors, *Proc. 6th European Conference on Advances in Artificial Life (ECAL 2001)*, Vol. 2159 of *Lecture Notes in Artificial Intelligence*, pp. 417–426. Springer-Verlag, 2001.
- [12] X. Chen, S. Kwong, and M. Li. A compression algorithm for DNA sequences and its applications in genome comparison. In *Proc. 10th Workshop on Genome Informatics*, pp. 52–61, 1999.
- [13] J. R. Cole, R. E. Funk, L. R. Godfrey, and W. Starna. On criticisms of “Some Paleolithic tools from northeast North America”: rejoinder. *Current Anthropology* **19** (1978), 665–669.
- [14] R. Collins. An evaluation of William A. Dembski’s *the design inference*: a review essay. *Christian Scholar’s Review* **30** (2001), 329–341.
- [15] W. A. Dembski. Intelligent design as a theory of information. *Perspectives on Science and Christian Faith* **49** (1997), 180–190. <http://www.leaderu.com/offices/dembski/docs/bd-idesign2.html>
- [16] W. A. Dembski. *The Design Inference: Eliminating Chance Through Small Probabilities*. Cambridge University Press, 1998.
- [17] W. A. Dembski. *Intelligent Design: The Bridge Between Science & Theology*. InterVarsity Press, 1999.
- [18] W. A. Dembski. Explaining specified complexity. *Metaviews*, (September 13 1999), # 139 (electronic). http://www.metanexus.net/archives/message_fs.asp?ARCHIVEID=3066

- [19] W. A. Dembski. *No Free Lunch: Why Specified Complexity Cannot Be Purchased Without Intelligence*. Rowman & Littlefield, 2002.
- [20] A. D. Dinsmore, D. T. Wong, P. Nelson, and A. G. Yodh. Hard spheres in vesicles: curvature-induced forces and particle-induced curvature. *Physical Review Letters* **80** (1998), 409–412.
- [21] C. Doumas. *The Wall-Paintings of Thera*. Thera Foundation, 1992.
- [22] T. Edis. Darwin in mind: ‘intelligent design’ meets artificial intelligence. *Skeptical Inquirer* **25**(2) (2001), 35–39.
- [23] B. Fitelson, C. Stephens, and E. Sober. How not to detect design — critical notice: William A. Dembski, *The Design Inference*. *Philosophy of Science* **66** (1999), 472–488.
- [24] J. L. Fitton. *The Discovery of the Greek Bronze Age*. British Museum Press, 1995.
- [25] B. Forrest. The wedge at work: how intelligent design creationism is wedging its way into the cultural and academic mainstream. In R. T. Pennock, editor, *Intelligent Design Creationism and Its Critics*, pp. 5–53. MIT Press, 2001.
- [26] A. Gajardo, A. Moreira, and E. Goles. Complexity of Langton’s ant. *Disc. Appl. Math.* **117** (2002), 41–50.
- [27] G. G. Gallup, Jr. Chimpanzees: self-recognition. *Science* **167** (1970), 86–87.
- [28] G. G. Gallup, Jr. Self-awareness and the emergence of mind in primates. *American J. Primatology* **2** (1982), 237–248.
- [29] G. W. Gilchrist. The elusive scientific basis of intelligent design theory. *Reports of the NCSE* **17**(3) (1997), 14–15.
- [30] E. Glicksman. Free dendritic growth. *Mater. Sci. Eng.* **65** (1984), 45.
- [31] F. Göbel. On the number of Hamiltonian cycles in product graphs. Technical Report 289, Technische Hogeschool Twente, Netherlands, 1979.
- [32] P. Godfrey-Smith. Information and the argument from design. In R. T. Pennock, editor, *Intelligent Design Creationism and Its Critics*, pp. 577–596. The MIT Press, 2001.
- [33] E. Goles and M. Margenstern. Sand pile as a universal computer. *Internat. J. Modern Phys. C* **7**(2) (1996), 113–122.
- [34] E. Goles, O. Schulz, and M. Markus. Prime number selection of cycles in a predator-prey model. *Complexity* **6**(4) (2001), 33–38.
<http://www3.interscience.wiley.com/cgi-bin/fulltext?ID=84502365&PLACEB%0=IE.pdf>
- [35] R. Greenler. *Rainbows, Halos, and Glories*. Cambridge University Press, 1980.

- [36] S. Grumbach and F. Tahi. A new challenge for compression algorithms: genetic sequences. *Information Processing and Management* **30** (1994), 875–886.
- [37] G. S. Hawkins. *Stonehenge Decoded*. Doubleday, 1965.
- [38] F. Heeren. The deed is done. *American Spectator* (2000/1), 28. <http://www.spectator.org/archives/0012TAS/heeren0012.htm>
- [39] R. A. Heltzer and S. A. Vyse. Intermittent consequences and problem solving: the experimental control of “superstitious” beliefs. *Psychological Record* **44** (1994), 155–169.
- [40] L. M. Herman, S. A. Kuczaj, and M. D. Holder. Responses to anomalous gestural sequences by a language-trained dolphin: evidence for processing of semantic relations and syntactic information. *J. Exper. Psych.* **122** (1993), 184–194.
- [41] A. Hewish, S. J. Bell, J. D. H. Pilkington, P. F. Scott, and R. A. Collins. Observation of a rapidly pulsating radio source. *Nature* **217** (February 24, 1968), 709–713.
- [42] M. Hirvensalo. *Quantum Computing*. Springer-Verlag, 2001.
- [43] E. A. Jagla and A. G. Rojo. Sequential fragmentation: the origin of columnar quasi-hexagonal patterns. *Physical Review E* **65** (2002), 026203.
- [44] D. Kahneman, P. Slovic, and A. Tversky. *Judgment Under Uncertainty: Heuristics and Biases*. Cambridge University Press, 1982.
- [45] C. Kamminga, A. B. Cohen Stuart, and M. G. de Bruin. A time-frequency entropy measure of uncertainty applied to dolphin echolocation signals. *Acoustics Letters* **21**(8) (1998), 155–160.
- [46] P. D. Kaplan, J. L. Rouke, A. G. Yodh, and D. J. Pine. Entropically driven surface phase separation in binary colloidal mixtures. *Physical Review Letters* **72** (1994), 582–585.
- [47] L. Kari. DNA computing: Arrival of biological mathematics. *Math. Intelligencer* **19**(2) (1997), 9–22.
- [48] M. A. Kessler and B. T. Werner. Self-organization of sorted patterned ground. *Science* **299** (2003), 380–383.
- [49] D. Kestenbaum. Gentle force of entropy bridges disciplines. *Science* **279** (1998), 1849.
- [50] J. M. Keynes. *A Treatise on Probability*. Macmillan, 1957.
- [51] W. Kirchherr, M. Li, and P. Vitányi. The miraculous universal distribution. *Math. Intelligencer* **19**(4) (1997), 7–15. <http://www.cwi.nl/~paulv/papers/mathint97.ps>

- [52] A. N. Kolmogorov. Three approaches to the quantitative definition of information. *Problemy Peredači Informacii* **1** (1965), 3–11. In Russian. English translation in *Problems Inform. Transmission* **1** (1965), 1-7 and *Internat. J. Computer Math.* **2** (1968), 157–168.
- [53] Robert C. Koons. Remarks while introducing Dembski’s talk at the conference *Design, Self-Organization and the Integrity of Creation*, Calvin College, Grand Rapids, Michigan. May 25 2001.
- [54] J. R. Koza. Artificial life: spontaneous emergence of self-replicating and evolutionary self-improving computer programs. In C. G. Langton, editor, *Artificial Life III*, pp. 225–262. Addison-Wesley, 1994.
- [55] L. Kuhnert, K. I. Agladze, and V. I. Krinsky. Image processing using light-sensitive chemical waves. *Nature* **337** (1989), 244–247.
- [56] J. K. Lanctot, M. Li, and E. h. Yang. Estimating DNA sequence entropy. In *Proc. 11th ACM-SIAM Symp. Discrete Algorithms (SODA)*, pp. 409–418, 2000.
- [57] P. S. Laplace. *A Philosophical Essay on Probabilities*. Dover, 1952.
- [58] R. E. Lenski, C. Ofria, R. T. Pennock, and C. Adami. The evolutionary origin of complex features. *Nature* **423** (2003), 139–145.
- [59] M. Li. Compressing DNA sequences. In T. Jiang, Y. Xu, and M. Q. Zhang, editors, *Current Topics in Computational Molecular Biology*, pp. 157–171. The MIT Press, 2002.
- [60] H. Lipson and J. B. Pollack. Automatic design and manufacture of robotic lifeforms. *Nature* **406** (2000), 974–978.
- [61] M. Li and P. Vitányi. *An Introduction to Kolmogorov Complexity and Its Applications*. Springer-Verlag, New York, 2nd edition, 1997.
- [62] J. McKim Malville and C. Putnam. *Prehistoric Astronomy in the Southwest*. Johnson Books, 1989.
- [63] J. M. Malville, F. Wendorf, A. A. Mazar, and R. Schild. Megaliths and Neolithic astronomy in southern Egypt. *Nature* **392** (1998), 488–491.
- [64] N. Marinatos. *Art and Religion in Thera: Reconstructing a Bronze Age Society*. D. & I. Mathioulakis, 1984.
- [65] Steve Maynard. Life’s intelligent design: Author of ‘Darwin on Trial’ sees a win. *Tacoma News Tribune* (2001), year. <http://www.discovery.org/news/life’sIntelligentDesign.html>
- [66] P. B. Medawar. *The Limits of Science*. Harper & Row, 1984.
- [67] S. C. Meyer. DNA and other designs. *First Things*, No. 102, (April 2000), 30–38.

- [68] L. Morgan. *The Miniature Wall Paintings of Thera: a Study in Aegean Culture and Iconography*. Cambridge University Press, 1988.
- [69] J. D. North. *Stonehenge: Neolithic Man and the Cosmos*. HarperCollins, 1996.
- [70] I. M. Pepperberg. Cognition and communication in an African Grey Parrot (*Psittacus erithacus*): Studies on a nonhuman, nonprimate, nonmammalian subject. In H. L. Roitblat, L. M. Herman, and P. E. Nachtigall, editors, *Language and Communication: Comparative Perspectives*, pp. 221–248. Lawrence Erlbaum, 1993.
- [71] M. Pigliucci. Chance, necessity, and the new holy war against science. A review of W. A. Dembski’s *The Design Inference*. *BioScience* **50** (2000), 79–81.
- [72] M. Pigliucci. Design yes, intelligent no: a critique of intelligent design theory and neocreationism. *Skeptical Inquirer* **25**(5) (2001), 34–39.
- [73] M. B. Pour-El and N. Zhong. The wave equation with computable initial data whose unique solution is nowhere computable. *Math. Logic Quarterly* **43** (1997), 499–509.
- [74] N. G. Rambidi and D. Yakovenchuk. Chemical reaction-diffusion implementation of finding the shortest paths in a labyrinth. *Physical Review E* **63** (2001), 026607.
- [75] T. Ray. Evolution, complexity, entropy, and artificial reality. *Physica D* **75** (1994), 239–263.
- [76] T. Ray. Evolution of complexity: Tissue differentiation in network Tierra. (2001). <http://www.isd.atr.co.jp/~ray/pubs/atrjournal/index.html>
- [77] M. Rizzotti. *Early Evolution: From the Appearance of the First Cell to the First Modern Organisms*. Birkhäuser, 2000.
- [78] D. Roche. A bit confused: creationism and information theory. *Skeptical Inquirer* **25**(2) (2001), 40–42.
- [79] P. W. K. Rothemund and E. Winfree. The program-size complexity of self-assembled squares. In *Proc. Thirty-second Ann. ACM Symp. Theor. Comput.*, pp. 459–468. ACM, 2000.
- [80] J. M. Rudski, M. I. Lischner, and L. M. Albert. Superstitious rule generation is affected by probability and type of outcome. *Psychological Record* **49** (1999), 245–260.
- [81] B. Russell. Mathematical logic as based on the theory of types. *Amer. J. Math.* **30** (1908), 222–262.
- [82] R. M. Sainsbury. *Paradoxes*. Cambridge University Press, 2nd edition, 1995.
- [83] E. S. Savage-Rumbaugh. Language learnability in man, ape, and dolphin. In H. L. Roitblat, L. M. Herman, and P. E. Nachtigall, editors, *Language and Communication: Comparative Perspectives*, pp. 457–473. Lawrence Erlbaum, 1993.

- [84] A. O. Schmitt and H. Herzel. Estimating the entropy of DNA sequences. *J. Theoretical Biology* **188** (1997), 369–377.
- [85] T. D. Schneider. Rebuttal to William A. Dembski’s posting. (June 6 2001). <http://www.lecb.ncifcrf.gov/~toms/paper/ev/dembski/rebuttal.html>
- [86] J. Shallit and M. w. Wang. Automatic complexity of strings. *J. Automata, Languages, and Combinatorics* **6** (2001), 537–554.
- [87] C. Shannon. Prediction and entropy of printed english. *Bell System Tech. J.* **3** (1950), 50–64.
- [88] K. Sims. Evolving 3D morphology and behavior by competition. In R. A. Brooks and P. Maes, editors, *Artificial Life IV: Proceedings of the Fourth International Workshop on the Synthesis and Simulation of Living Systems*, pp. 28–39. MIT Press, 1994.
- [89] O. Steinbock, A. Tóth, and K. Showalter. Navigating complex labyrinths: optimal paths from chemical waves. *Science* **267** (1995), 868–871.
- [90] G. Theraulaz and E. Bonabeau. Coordination in distributed building. *Science* **269** (1995), 686–688.
- [91] P. Vitányi. Quantum Kolmogorov complexity based on classical descriptions. *IEEE Trans. Inform. Theory* **47** (2001), 2464–2479.
- [92] R. Wein. What’s wrong with the design inference. *Metaviews*, (November 6 2000), # 96 (electronic). http://www.metanexus.org/archives/message_fs.asp?ARCHIVEID=2654
- [93] T. A. Welch. A technique for high performance data compression. *IEEE Computer* **17**(6) (1984), 8–19.
- [94] J. Wilkins and W. Elsberry. The advantages of theft over toil: the design inference and arguing from ignorance. *Biology and Philosophy* **16** (2001), 711–724. <ftp://ftp.wehr.edu.au/pub/wilkinsftp/dembski.pdf>
- [95] D. H. Wolpert and W. G. Macready. No free lunch theorems for optimization. *IEEE Transactions on Evolutionary Computation* **1** (1997), 67–82.
- [96] D. R. Yelen. The acquisition and extinction of superstitious behavior. *J. Experimental Research in Personality* **5** (1971), 1–6.
- [97] H. P. Yockey. Behe’s irreducible complexity and evolutionary theory. *NCSE Reports* **21**(3-4) (2001), 18–20.
- [98] J. Ziv and A. Lempel. A universal algorithm for sequential data compression. *IEEE Trans. Inform. Theory* **23** (1977), 337–343.