

Report Information**Author:** [REDACTED] **Folder:** Top Folder**Filename:** Elsberry2011.rtf **Submitted:** Mon, Jun 06 2011, 4:28 PM**Matching:**  94% **Paper ID:** 36888806**Suspected Sources****[1]** Another user's paper: Owner: [REDACTED] Submitted: Mon, Jun 06 2011, 4:28 PM; Filename: Elsberry2003.rtf**[2]** http://diccionario.sensagent.com/Specified_complexity/en-en/**Excluded Sources**http://www.docstoc.com/docs/421325/2005-05-16_Shallit_Ps_expert_rebuttal_readable<http://www.scribd.com/doc/23648185/Elsberry-Shallit-2003-54><http://pandasthumb.org/archives/2010/12/synthese-issue.html><http://www.scribd.com/doc/23648261/Young-Edis-Eds-Why-Intelligent-Design-Fails-A-Scientific-Critique><http://www.antievolution.org/people/wre/papers/eandsdembski.pdf>http://www.iscid.org/boards/ubb-get_topic-f-6-t-000464-p-4.html<http://ncseprojects.org/book/export/html/1859><http://ncse.com/rncse/23/5-6/eight-challenges-intelligent-design-advocates><http://good-without-god.blogspot.com/feeds/posts/default?orderby=updated><http://www.science-meets-religion.org/evolution/probability.php><http://dictionary.sensagent.com/specified+complexity/en-en/><http://diccionario.sensagent.com/SPECIFIED%20COMPLEXITY/en-en/><http://talkorigins.org/indexcc/CI/CI010.html>http://www.iscid.org/boards/ubb-get_topic-f-6-t-000543.html

Another student's paper

Another student's paper

<http://www.ideacenter.org/contentmgr/showdetails.php/id/1488>http://scienceblogs.com/goodmath/2006/06/dembskis_profound_lack_of_comp.phphttp://lofi.forum.physorg.com/%5Bb%5Dthe-Verdict-On-Intelligent-Design%5Bb%5D_19128-200.html<http://ncse.com/book/export/html/2453>

Another student's paper

<http://www.talkorigins.org/design/faqs/nfl/>http://wn.com/specified_complexity<http://austringer.net/wp/index.php/2009/04/>http://www.absoluteastronomy.com/topics/Specified_complexityhttp://scienceblogs.com/goodmath/2009/12/id_garbage_csi_as_non-computab.php

Another student's paper

Paper Text**[1; 100 %]** *Information theory, evolutionary computation, and Dembski's "complex specified information"***[2; 85 %]** *Wesley Elsberry Jeffrey Shallit*

[1; 96 %] *Abstract Intelligent design advocate William Dembski has introduced a measure of information called "complex specified information", or CSI. [1; 100 %] He claims that CSI is a reliable marker of design by intelligent agents. [1; 100 %] He puts forth a "Law of Conservation of Information" which states that chance and natural laws are incapable of generating CSI. [1; 100 %] In particular, CSI cannot be generated by evolutionary computation. [1; 100 %] Dembski asserts that CSI is present in intelligent causes and in the flagellum of*

Escherichia coli, and concludes that neither have natural explanations. [1; 100 %] 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; 79 %] 1 Introduction In recent books and articles (eg. Dembski 1998, 1999, 2002, 2004), theologian and mathematician William Dembski uses a semi-mathematical treatment of information theory to justify his claims about "intelligent design". [1; 100 %] Roughly speaking, intelligent design advocates attempt to infer intelligent causes from observed instances of complex phenomena. [1; 100 %] Proponents argue, for example, that biological complexity indicates that life was designed. [1; 75 %] This claim is usually presented as an alternative to the theory of evolution.

[1; 76 %] Christian apologist William Lane Craig has called Dembski's work "groundbreaking" (Dembski 1999, blurb at beginning). Journalist Fred Heeren describes Dembski as "a leading thinker on applications of probability theory" (Heeren 2000). [1; 90 %] However, according to a 2006 search of MathSciNet, the American Mathematical Society's online version of Mathematical Reviews, a journal that 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 peerreviewed paper in any mathematics journal at all. Dembski's CV (available at <http://www.designinference.com>) lists another paper in Journal of Statistical Computation and Simulation in 1990 that was not reviewed by Mathematical Reviews. These papers have received very few citations, suggesting the lack of mathematical impact. For more details, see Shallit (2004).

University of Texas philosophy professor Robert Koons (2001) called Dembski the "Isaac Newton of information theory." However, according to Mathematical Reviews, Dembski has not published any papers in any peer-reviewed journal devoted to information theory, although recently he has made available some preprints dealing with this topic on his website.

Is the effusive praise of Craig, Heeren, and Koons warranted?

[1; 100 %] We believe it is not. [1; 100 %] As we will show, Dembski's work is riddled with inconsistencies, equivocation, flawed use of mathematics, poor scholarship, and misrepresentation of others' results. [1; 100 %] As a result, we believe few if any of Dembski's conclusions can be sustained.

Many writers have already taken issue with some of Dembski's claims (eg. Fitelson et al. 1999; Pigliucci 2000, 2001; Wein 2000; Roche 2001; Edis 2001; Wilkins and Elsberry 2001; Godfrey-Smith 2001; Shallit 2002; [2; 73 %] Elsberry and Shallit 2003; Perakh 2004; Young and Edis 2004; Forrest and Gross 2004; Olofsson 2007). [1; 88 %] In this paper, we focus on the mathematical aspects of Dembski's work that have received comparatively little attention thus far.

[1; 100 %] Here is an outline of the paper. [1; 100 %] First, we summarize what we see as Dembski's major claims. We examine his generic chance elimination argument (GCEA) and briefly show how it is flawed. [1; 92 %] We then turn to one of Dembski's major concepts, "complex specified information" (CSI), arguing that he uses the term inconsistently and misrepresents the concepts of other authors as being equivalent. [1; 83 %] We criticize Dembski's concepts of "information" and "specification". [1; 100 %] We then address his "Law of Conservation of Information", showing that the claim has significant mathematical flaws. [1; 100 %] We then discuss Dembski's attack on evolutionary computation, showing his claims are unfounded.

Some of the criticisms in this paper have already appeared in an abbreviated form (Shallit 2002) and in a more popular treatment (Young and Edis 2004).

[1; 100 %] 2 Dembski's claims Dembski makes a variety of different claims, many of which would be revolutionary if true. [1; 100 %] 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; 93 %] (1) There exists a multi-step statistical procedure, the "generic chance elimination argument", that reliably detects design by intelligent agents (Sects. 3, 4).

(2) There is a "souped-up" form of information (Dembski 2002, p. [1; 83 %] 142) called "specified complexity" or "complex specified information" (CSI) which is coherently defined and constitutes a valid, useful, and non-trivial measure (Sects. [1; 63 %] 4, 5, 6).

[1; 82 %] (3) Many human activities exhibit "specified complexity" (Sect. 4).

[1; 94 %] (4) CSI cannot be generated by deterministic algorithms, chance, or any combination of the two. [1; 100 %] In particular, CSI cannot be generated either by genetic algorithms implemented on computers, or the process of biological evolution itself. [1; 94 %] A "Law of Conservation of Information" exists which says that natural processes cannot generate CSI (Sects. 7, 8).

[1; 71 %] (5) Life exhibits specified complexity and hence was designed by an intelligent agent (Sect. 9).

3 The generic chance elimination argument

Dembski's generic chance elimination argument (GCEA) exists in at least two different forms (Dembski 1998, 2002). We roughly summarize one version here:

[1; 100 %] An intelligent agent A witnesses an event E, and assigns it to some reference class of events. A lists all possible hypotheses H1, H2, H3, ... involving deterministic and random processes that could account for E. Next, A invents a rejection function f and a rejection region R of a certain special form that includes E. A determines "background knowledge" that "explicitly and univocally identifies" f. (This knowledge must be independent of the hypotheses in a certain technical sense, as discussed in Sect. 6.) A selects a significance level α and computes the

The existence of the Erroneous Design Inference Principle receives confirmation from modern research in

psychology. [1; 78 %] For one thing, humans are notoriously poor judges of probability (Kahneman et al. 1982). [1; 72 %] On the other hand, humans are good detectors of patterns, even when they are not there (Catania and Cutts 1963; Yelen 1971; Heltzer and Vyse 1994; Rudski et al. 1999). [1; 88 %] 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 (Boyer 2001). [1; 87 %] 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 Erroneous Design Inference Principle: [1; 100 %] ghosts, UFO's, and witchcraft.

[1; 100 %] But back to our analysis of the Caputo case. [1; 100 %] 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 remarked,

[1; 97 %] 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 combinations that can arise in a 100 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 1952, pp. 16–17)

[1; 100 %] Laplace's argument has been updated in modern form to reflect Kolmogorov complexity: [1; 65 %] see, for example, the wonderful article (Kirchherr et al. 1997) or our own (Elsberry and Shallit 2004). Let $C(x)$ denote the Kolmogorov complexity of the string of symbols x ; roughly speaking, this is the length of the shortest combination of program P and input i such that P outputs x on input i . [1; 85 %] 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 c is very low: [1; 74 %] we cannot compute it exactly, but let us say for the sake of argument that $C(c) \leq 10$. [1; 99 %] Thus, the hypothesis that c is due to flipping a fair coin has probability $\leq 2^{-30}$, or about 1 in a billion, and it seems fair to reject it.

After ruling out the chance hypothesis that the sequence resulted from flips of a fair coin, what next? [1; 85 %] Dembski would have us believe that design by an intelligent agent is now a purely mathematical implication. [1; 82 %] But what of the possibilities (a)–(d) given above? Dembski does not consider them. We conclude that determining design cannot be a purely eliminative argument as Dembski suggests; instead, hypotheses involving intelligent design must be considered alongside non-design hypotheses.

[1; 67 %] An alternate view is that if specified complexity can be used to detect something, what is detected is the output of simple computational processes. (Of course, it is possible for complicated computational processes to generate simple outputs. The point is that simple outputs do not demand an inference of complicated computational processes; [1; 83 %] simple ones will suffice.) 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" (Dembski 2002, p. 144). [1; 100 %] 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 (Hirvensalo 2001), DNA computation (Kari 1997), chemical computing (Kuhnert et al. 1989; Steinbock et al. 1995; Rambidi and Yakovenchuk 2001), and molecular selfassembly (Rothmund and Winfree 2000). While most of these references deal with how these naturally occurring tools can be adapted to serve human ends, to us they suggest that chemical and physical processes could well perform computation without intelligent intervention.

[1; 62 %] Furthermore, it is now known that even very simple computational models, such as Conway's game of Life (Berlekamp et al. 1982), Langton's ant (Gajardo et al. 2002), and sand piles (Goles and Margenstern 1996) are universal, and hence compute anything that is computable. [1; 83 %] Finally, in the cellular automaton model, relatively simple replicators are possible (Byl 1989).

[1; 81 %] 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.

We now turn to Dembski's second example of the GCEA: [1; 66 %] his discussion of a SETI primes sequence

$t = 110111011111011111110 \dots$

[1; 100 %] which is a variation on a signal received by fictional researchers in the movie Contact. [1; 92 %] As Dembski describes it, t 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 1,000. [1; 94 %] Dembski suggests the specified complexity of this sequence implies a design inference.

[1; 100 %] Yet is that the case? We know that prime numbers arise naturally in simple predator-prey models (Goles et al. [1; 71 %] 2001), so it is at least conceivable that prime number signals could result from some non-intelligent physical process. [1; 93 %] To infer intelligent design upon receiving t simply means that we estimate the relative probability of natural prime-number generation as lower than the probability that the signal arises from some intelligence that considers prime numbers an interesting way to communicate. [1; 93 %] In other words, we compare two hypotheses, one involving design, one not. This decision method is explicitly ruled out by Dembski's method.

[1; 92 %] Dembski is fond of argument based on fictional examples, so it is instructive to compare Dembski's treatment of the cinematic SETI sequence from Contact with the history of an actual reception of an extraterrestrial signal. [1; 100 %] Pulsars (rapidly pulsating extraterrestrial radio sources) were discovered by Jocelyn Bell in 1967. [1; 90 %] She observed a long series of pulses of period 1.337 s. [1; 89 %] In at least one case the signal was tracked for 30 consecutive minutes, which would represent approximately 1,340 pulses. Like the Contact sequence, this sequence was viewed as improbable (hence "complex") and specified (see Sect. [1; 63 %] 6), hence presumably it would constitute complex specified information and trigger a design inference. [1; 100 %]

Yet spinning neutron stars, and not design, are the current explanation for pulsars.

[1; 100 %] Bell and her research team immediately considered the possibility of an intelligent source. [1; 94 %] (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" (Hewish et al. 1968).

[1; 100 %] However, the hypothesis of intelligent agency was rejected for two reasons. [1; 100 %] First, parallax considerations ruled out a terrestrial origin. [1; 100 %] Second, additional signals were discovered originating from other directions. [1; 100 %] The widely separated origins of multiple signals decreased the probability of a single intelligent source, and multiple intelligent sources were regarded as implausible. [1; 100 %] In other words, hypotheses involving design were considered at the same time as non-design hypotheses, instead of the eliminative approach Dembski proposes. In this real-life example, Dembski's approach was not used, which is fortunate, as it would have provided the wrong answer.

4 Complex specified information

[1; 80 %] As we have seen, Dembski's generic chance elimination argument requires the elimination of all relevant chance hypotheses. [1; 100 %] If all such hypotheses are eliminated, Dembski concludes design is the explanation for the event in question.

[1; 100 %] Although Dembski spends significant space discussing the GCEA, in practice he rarely uses it. [1; 100 %] Instead, he employs an alternate approach. [1; 83 %] 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."

[1; 83 %] According to Dembski, both approaches serve to detect a certain property of events, called "specified complexity" or "complex specified information" (CSI). [1; 86 %] Dembski insists that "if there is a way to detect design, specified complexity is it." (Dembski 2002, p. [1; 100 %] 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.

[1; 82 %] Dembski conflates his procedure to eliminate hypotheses with the property of CSI (Dembski 2002, p. [1; 100 %] 73) with no significant explanation. [1; 100 %] 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.

[1; 100 %] 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. [1; 100 %] Instead, Dembski uses "complex" as a synonym for "improbable".

[1; 100 %] 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: [1; 91 %] in particular, it is neither Shannon's entropy, surprisal, nor Kolmogorov complexity. [1; 90 %] Although Dembski claims that CSI "is increasingly coming to be regarded as a reliable marker of purpose, intelligence, and design" (Dembski 2002, p. [1; 100 %] 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. [1; 90 %] A 2006 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. [1; 91 %] (The term "CSI" does appear, but as an abbreviation for unrelated concepts such as "contrast source inversion," "conditional symmetric instability," "conditional statistical independence," "channel state inversion," and "constrained statistical inference.")

[1; 100 %] (A recent paper by creationist Stephen C. Meyer (2000) states

[1; 90 %] Systems that are characterized by both specificity and complexity (what information theorists call "specified complexity") have "information content."

[1; 100 %] 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"? [1; 100 %] He then admitted that he knew no one but Dembski.)

[1; 85 %] Despite his insistence that his "program has a rigorous information-theoretic underpinning" (Dembski 2002, p. [1; 94 %] 371), the term CSI is used inconsistently in Dembski's own work. [1; 100 %] Sometimes CSI is a quantity that one can measure in bits: [1; 81 %] "the CSI of a flagellum far exceeds 500 bits" (Dembski 1999, p. 178). [1; 100 %] Other times, CSI is treated as a threshold phenomenon: [1; 77 %] something either "exhibits" CSI or does not: [1; 87 %] "The Law of Conservation of Information says that if X exhibits CSI, then so does Y" (Dembski 2002, p. 163). [1; 80 %] Sometimes numbers or bit strings "constitute" CSI (Dembski 1999, p. 159); [1; 91 %] other times CSI refers to a pair (T, E) where E is an observed event and T is a pattern to which E conforms (Dembski 2002, p. 141). [1; 74 %] Sometimes CSI refers to specified events of probability < 10⁻¹⁵⁰: [1; 83 %] other times it can be contained in "the 16-digit number on your VISA card" or "even your phone number" (Dembski 1999, p. 159). [1; 67 %] 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 (Dembski 2002, p. 153). [1; 100 %] Other times it is made clear that computing CSI crucially depends on the background knowledge of the observer. [1; 100 %] Sometimes CSI inheres in a string regardless of its causal history (this seems always to be the case in natural language utterances): [1; 100 %] other times the causal history is essential to judging whether or not a string has CSI. [1; 100 %] CSI is indeed a measure with remarkably fluid properties! [1; 100 %] Like Blondlot's N-rays, however, the existence of CSI seems clear only to its discoverer.

a probability of around 1 in 1040, E appears to have done just that, to wit, generate a highly improbable specified event, or what we are calling specified complexity. (Dembski 2002, p. 194)

[1; 100 %] 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; [1; 95 %] instead we should compute the information contained in the result based on a uniform distribution on all strings of length 28 over an alphabet of size 27 (note that $27^{28} \approx 1.197 \times 10^{40}$).

[1; 100 %] Sometimes the uniform probability interpretation is applied even when a frequentist approach is strongly suggested. [1; 82 %] For example, when discussing the Contact primes string

t=11011101111101111110 [...]

[1; 81 %] Dembski claims its probability is "1 in 21000" (Dembski 2002, p. [1; 84 %] 144), a claim which is viable only under the uniform probability interpretation. [1; 100 %] But, viewing only the singular instance t, there are in fact many possibilities:

[1; 100 %] (a) both 0 and 1 are emitted with probability 1/2;

[1; 91 %] (b) 1 is emitted with probability 0.977 and 0 is emitted with probability 0.023;

[1; 100 %] (c) the emitted bits correspond to the unary encodings of 24 numbers between 1 and 100 chosen randomly with replacement;

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

(Note: [1; 100 %] 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 Sect. 6.)

[1; 100 %] We do not see how, in the absence of more information, to distinguish between these possibilities and dozens of others. [1; 100 %] And the choice is crucial. [1; 85 %] A purely frequentist approach, as in (b), results in a markedly different probability estimate from (a)—the probability of t increases from 2×10^{-40} to 9.33×10^{-302} to $0.9779770.02323 \approx 2.8 \times 10^{-48}$.

[1; 100 %] 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. [1; 64 %] (The "universal probability bound," 10^{-150} , is Dembski's estimate for the smallest probability of a specified event that could occur randomly sometime during the history of the universe.) An approach such as (d) gives an even higher probability of $25 \times 10^{-24} \approx 2.8 \times 10^{-34}$.

[1; 100 %] 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. [1; 100 %] Furthermore, each of the two approaches has significant difficulties for Dembski's program. [1; 100 %] The causal-history-based interpretation is the only one that is mathematically tenable; [1; 100 %] its probability estimates are necessarily based on a thorough understanding of the origin of the event in question. [1; 92 %] 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" (Dembski 2002, p. [1; 71 %] xi). We expand on this in Sect. 7.

[1; 94 %] The uniform probability interpretation is, at first glance, easier to apply, and may be viewed as a form of the classical Principle of Indifference. [1; 100 %] But this principle has long been known to be quite problematical; [1; 87 %] as Keynes has remarked, "This rule, as it stands, may lead to paradoxical and even contradictory conclusions." (Keynes 1957, p. 42). We will see in Sect. [1; 78 %] 7 that the uniform probability interpretation is incompatible with Dembski's "Law of Conservation of Information".

[1; 100 %] 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. [1; 88 %] If we encounter a string of the form 1000 [...]

should we regard it as chosen from the alphabet = {0} or = {0,1}? [1; 68 %] Should we regard it as chosen from the space of all strings of length 1,000, or all strings of length $\leq 1,000$? Dembski's advice (Dembski 2002, Sect. [1; 65 %] 3.3) is singularly unhelpful here; [1; 73 %] 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. [1; 86 %] For an example of this, see our discussion of the information content of Dawkins' fitness function in Sect. 7.

[1; 100 %] 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.

6 Specification

[1; 88 %] The second ingredient of CSI is specification. [1; 100 %] By "specification" of an event E, Dembski roughly means a pattern to which E conforms. [1; 100 %] Furthermore, Dembski demands that the pattern, in some sense, be given independently of E. Dembski's initial metaphor for "specification" and "fabrication" is that of an archer loosing an arrow at a wall. If we find that the archer places his arrow into a pre-painted target, Dembski says that this corresponds to his idea of specification. If the archer instead paints his target around the arrow after the fact, the pattern is instead not one from which we may infer design, and this sort of pattern Dembski calls a "fabrication." However, Dembski's metaphor is inapt. Our task in detecting design in the artifacts of biology is not one of observing an agent at work who either uses a target or tries to make it appear that a target existed falsely. We do not have any information bearing upon such an agent. If we were to re-work the metaphor for somewhat better accuracy, Dembski's situation is that upon finding an arrow stuck in a wall, he tries to convince himself that he

To understand specification, at least in one formulation (Dembski 2002, pp. 62–63), we must return to the GCEA and examine it in more detail. **[1: 89 %]** Recall that in the GCEA an intelligent agent A witnesses an event E and assigns it to some reference class of events. **[1: 90 %]** 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 $f: \rightarrow R$. Then a target T is defined by either $T = \{ \bar{u} \in: f(\bar{u}) \geq \bar{a} \}$ or $T = \{ \bar{u} \in: f(\bar{u}) \leq \bar{a} \}$ for some given real numbers $\bar{a}$, $\bar{a}$. **[1: 86 %]** If K is “epistemically independent” of E (by this Dembski means that $P(E|H\&K) = P(E|H)$), then T is said to be “detachable” from E. **[1: 100 %]** (Here H is a hypothesis that E is due to chance.) Finally, if $E \subset T$, then T is a specification for E and E is said to be “specified”.

[1: 82 %] We find Dembski's account of specification incoherent. [1: 100 %] Briefly, here are our objections. [1: 97 %] First, we contend that Dembski has not adequately distinguished between legitimate and illegitimate specifications (which he calls fabrications). [1: 100 %] Second, Dembski's notion of specification is too vague. [1: 100 %] Third, Dembski's discussion of the generation of the target T and its independence of the event E is problematic.

```
t=110111011111011111110 [...]
```

[1; 100 %] 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 {D, R}, containing at most one R”; [1; 100 %] this is apparently a valid specification for the Caputo string

[1: 85 %] Now suppose we witness Mr Caputo produce yet another choice to head the ballot. **[1: 64 %]** 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.” (It is true that this new specification increases the probability that the target is hit, but that is not relevant here.) But if this is the case, what prevents us from extending the process indefinitely? **[1: 90 %]** 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.

[1; 100 %] We also believe Dembski's current notion of specification is too vague to be useful. [1; 100 %] 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. [1; 86 %] According to its formal definition, a specification is supposed to be a rejection region R of the form $\{ \bar{u} \in : f(\bar{u}) \geq \bar{a} \}$ or $\{ \bar{u} \in : f(\bar{u}) \leq \bar{a} \}$ for an appropriate choice of a rejection function f and real numbers $\bar{a}$, $\bar{a}$. Now consider Dembski's discussion of the "specification" of the flagellum of *Escherichia coli*: ". [1; 96 %] 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. [1; 100 %] But what, precisely, is the space of events here? [1; 100 %] And what is the precise choice of the rejection function f and the rejection region R ? [1; 100 %] Dembski does not supply them. [1; 96 %] 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.

[1; 88 %] Third, we find Dembski's account of how the pattern is generated problematic. [1; 93 %] 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." (Dembski 2002, p. 18). [1; 100 %] This is a strangely worded requirement. [1; 100 %] For how could anyone verify that the event actually does conform to the pattern, without actually examining every bit of the event in question? [1; 100 %] To illustrate this example, let us return to the sequence t mentioned above.

[1; 100 %] Dembski says t is specified. [1; 69 %] 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 1,000." Presumably Dembski believes it self-evident that S could enable us to identify t "without recourse to the actual event." But we cannot, for in fact, S is not a specification of the actual printed sequence! [1; 78 %] A careful inspection of the string presented on pages 143–144 of No Free Lunch reveals that it is indeed of length 1,000, but omits the unary representation of the prime 59. [1; 78 %] In other words, the string Dembski actually presents is

[1; 67 %] $t' = 11011101111101111110$ [...]

[1; 91 %] So in fact our proposed specification S does not entail t , but instead t , and any pretense that we could have identified S without explicit recourse to t' vanishes.

[1; 100 %] We conclude that Dembski's account of specification is severely flawed.

[1; 82 %] 7 The Law of Conservation of Information

[1; 93 %] 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" (Dembski 2002, p. 163). [1; 100 %] One version of LCI states that CSI cannot be generated by natural causes; [1; 100 %] another states that neither functions nor random chance can generate CSI. [1; 100 %] We will see that there is simply no reason to accept Dembski's "Law", and that his justification is fatally flawed in several respects.

[1; 100 %] Furthermore, Dembski uses equivocation to suggest that his version of LCI is compatible with others in the literature. [1; 94 %] 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 (Dembski 2002, p. 129). [1; 100 %] 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*. [1; 86 %] Dembski repeats this claim when he discusses his own "Law of Conservation of Information" (Dembski 2002, p. 159). [1; 100 %] But is Medawar's law the same as Dembski's, or even comparable?

No. [1; 100 %] First of all, Medawar's remarks do not constitute a formal claim, since they appeared in a popular book without proof or detailed justification. [1; 69 %] In fact, Medawar acknowledges (Medawar 1984, p. [1; 92 %] 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."

[1; 100 %] 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. [1; 100 %] He does not formally define exactly what he means by information, but there is no mention of probabilities or the name Shannon. [1; 94 %] Certainly there is no reason to think that Medawar's "information" has anything to do with CSI. [1; 100 %] (Medawar's law, by the way, can be made rigorous, but in the context of Kolmogorov information, not Shannon information or Dembski's CSI; see Chaitin (1974). As we have already seen above in Sect. [1; 76 %] 4, Dembski's CSI and Kolmogorov complexity, if related at all, are related in an inverse sense.)

One of Dembski's most important claims is that functions cannot generate CSI. [1; 100 %] More precisely, Dembski claims that given $CSI_j = (T_1, E_1)$, based on a "reference class of possibilities 1", and a function f : [1; 97 %] $0 \rightarrow 1$ 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". [1; 100 %] Notice that Dembski makes no restrictions on f at all: [1; 100 %] it could be known to the agent who observes j , or not known. [1; 100 %] 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. [1; 66 %] It could be computable or non-computable.

[1; 97 %] Dembski's "proof" of this claim, given on pages 152–154 of No Free Lunch, is flawed in several ways. [1; 100 %] For the purposes of our discussion, let us restrict ourselves to the case where...where... [1; 100 %] are finite alphabets. [1; 89 %] By this we mean that events are represented by strings of symbols.

[1; 100 %] First of all, let us consider the uniform probability interpretation of CSI. [1; 83 %] Dembski justifies his assertion by transforming the probability space 1 by f^{-1} . [1; 100 %] This is reasonable under the causal-history-based interpretation. [1; 100 %] But under the uniform probability interpretation, we may not even know that j is formed by applying f to i . [1; 100 %] 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.

[1; 100 %] This error in reasoning can be illustrated as follows. [1; 100 %] Given a binary string x we may encode it in "pseudo-unary" as follows: [1; 100 %] 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. [1; 100 %] For example, the binary string 01 would be encoded in pseudo-unary as 111110. [1; 100 %] This encoding is reversible as follows: [1; 90 %]

count the number of 1's, write the result in binary, and delete the first 1. [1; 80 %] If we let $f: \dots$ [1; 100 %] be the mapping on binary strings giving a unary encoding, then it is easy to see that f can generate CSI. [1; 91 %] For example, suppose we consider an 10-bit binary string chosen randomly and uniformly from the space of all such strings, of cardinality 1,024. [1; 100 %] The CSI in such a string is clearly at most 10 bits. [1; 100 %] Now, however, we transform this space using f . [1; 72 %] The result is a space of strings of varying length l , with $1, 025 \leq l \leq 2, 048$. [1; 100 %] 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 . [1; 100 %] But now we cannot even apply f^{-1} to any of these strings, other than $f(i)$! [1; 97 %] 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 (cf. Sect. 3). [1; 100 %] The result is that $f(i)$ would be CSI, but i would not be.

[1; 100 %] Another error in Dembski's analysis is as follows. [1; 83 %] To obtain the detachability of $f^{-1}(T_1)$, Dembski says that " f merely [needs] to be composed with the rejection function on..." [1; 78 %] if g is the rejection function on 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 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 . [1; 77 %] While g is presumed identifiable in this sense relative to K , in what sense is $g \circ f$ so identifiable? It may not be, for two reasons. [1; 92 %] First, in the uniform probability interpretation of CSI, the intelligent agent who identified g may be entirely unaware of f . [1; 95 %] 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 . [1; 73 %] 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... identifiable from K , when composed with f , might give a compatible rejection function for....

[1; 85 %] Here is an example illustrating this error. [1; 62 %] Suppose j is an English message of 1,000 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. [1; 100 %] 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. [1; 100 %] 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 . [1; 100 %] Thus, in A 's view, CSI j was actually produced by applying f to i . [1; 100 %] The only way out of this paradox is to change A 's background knowledge to include knowledge about f . [1; 95 %] 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 .

[1; 100 %] 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 . [1; 81 %] 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 . (Dembski seems to admit this when he says that "... [1; 75 %] most claims are like this (i.e., they fail to induce well-defined probability distributions).") (Dembski 2002, p. 106.) But in applying the causal-history-based approach, it is absolutely crucial that every single step be known; [1; 100 %] 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.

[1; 100 %] 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. [1; 79 %] On pages 154–155 of No Free Lunch Dembski acknowledges that his proof that functions cannot generate CSI (pp. 152–154) is, in fact, not a proof at all. [1; 100 %] He forgot "the possibility that functions might add information". [1; 100 %] (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 $U(i, f) = f(i) = j$. [1; 72 %] He then argues that the amount of CSI in the pair (i, f) is at least as much as j . [1; 100 %] Of course, this claim also suffers from the two problems mentioned above, but now there is yet another error: [1; 100 %] Dembski does not discuss how to determine the CSI contained in f .

[1; 100 %] This is not a minor or insignificant omission. [1; 100 %] Recall that under one interpretation of LCI, f is supposed to correspond to some natural law. [1; 100 %] If f contains much CSI on its own, then by applying f we could accumulate CSI "for free". [1; 100 %] 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.

[1; 100 %] For example, consider the information contained in a fitness function in Dawkins' METHINKS IT IS LIKE A WEASEL example. [1; 100 %] 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. [1; 89 %] The cardinality of the space of all such fitness functions is 292728 , or about 25.816×1040 . [1; 95 %] 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" (Dembski 2002, p. 195). [1; 95 %] It is not clear what Dembski thinks the CSI of f is, since he never tells us explicitly. [1; 75 %] 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 5.816×1040 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 .

[1; 100 %] 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.... [1; 92 %] 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. [1; 100 %] Then f^n would be a map with the same domain and range as f . [1; 100 %] However, our objection gathers more force if f is a length-increasing map on strings. [1; 100 %] 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*.

[1; 100 %] To illustrate this possibility, consider the following procedure: starting from an empty string $x = \dots$, we successively choose randomly between applying the transformation $f_0(x) = 0x0$ or $f_1(x) = 1x1$. [1; 100 %] 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. [1; 84 %] Under the uniform probability interpretation, upon viewing y we would consider it a member of the uniform probability space $2n$, where $= \{0,1\}$. [1; 84 %] Assuming our background knowledge contains the notion of palindromes, the specification "palindrome" identifies a target space with $2n$ members, and so the probability of a randomly-chosen element of $2n$ hitting the target is 2^{-n} . [1; 100 %] In other words, y contains n bits of CSI. [1; 100 %] As n increases in size, we can generate as much CSI as we like.

7.1 Naturally-occurring CSI

[1; 100 %] Dembski seems to be of two minds about the possibility of CSI being generated by natural processes. [1; 100 %] For example, it would seem that the regular patterns formed by ice crystals would constitute CSI, at least under the uniform probability interpretation. [1; 97 %] 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. [1; 100 %] 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 ". [1; 78 %] 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)" (Dembski 2002, p. 12).

[1; 84 %] Just a paragraph later, Dembski discusses the occurrence of the Fibonacci sequence in phyllotaxis (the arrangement of leaves on plants). [1; 96 %] 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 Contact primes sequence, a legitimate instance of CSI. [1; 100 %] However, he argues that the CSI is not generated by the plant, but rather is a consequence of intelligent design of the plant itself. [1; 100 %] (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.

[1; 100 %] This seems inconsistent to us. [1; 100 %] If we apply the uniform probability interpretation consistently, it would seem that many natural processes, including some that are not biological, generate CSI. [1; 95 %] 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. [1; 100 %] Indeed, in response to one such challenge (the natural nuclear reactors at Oklo) he says

[1; 77 %] But suppose the Oklo reactors ended up satisfying this criterion after all. Would this vitiate the complexity-specification criterion? [1; 100 %] Not at all. [1; 92 %] At worst it would indicate that certain naturally occurring events or objects that we initially expected to involve no design actually do involve design. (Dembski 2002, p. 27)

[1; 70 %] In other words, Dembski's claims are, for him, unfalsifiable. [1; 91 %] We find this good evidence that Dembski's case for intelligent design is not a scientific one.

7.1.1 Dendrites

[1; 100 %] Dendrites are tree-like or moss-like structures that arise through crystal growth, particularly with iron or manganese oxides. [1; 100 %] If "tree-like in appearance" is a valid specification, it would seem that such structures could well constitute CSI. [1; 100 %] Indeed, their tree-like appearance often causes them to be confused with plant fossils. [1; 76 %] Dendrites were a puzzle until relatively recently (Glicksman 1984). [1; 100 %] Thus, until recently, they would have been assigned to design by Dembski's generic chance elimination argument. [1; 100 %] Despite the currently accepted physical explanation, they might still constitute CSI under the uniform probability interpretation.

7.1.2 Triangular ice crystals

[1; 100 %] Under certain rare conditions snow crystals form triangular plates. [1; 100 %] Unlike the case of ordinary six-sided snowflakes, there is currently no detailed physical explanation for the formation of triangular plates.

[1; 100 %] 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. [1; 100 %] Triangular snowflakes would then seem to qualify as CSI, at least under the uniform probability interpretation. [1; 100 %] They cannot be rejected as "necessity" since no known law accounts for their formation.

[1; 100 %] 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. [1; 100 %] This seems like an absurd conclusion to us.

[1; 77 %] 7.1.3 Self-ordering in collections of spheres of different sizes

[1; 100 %] Under certain conditions, mixtures of small spheres of different sizes will spontaneously self-organize in mysterious ways. [1; 100 %] This would seem to be an instance of CSI, at least under the uniform probability interpretation. [1; 73 %] However, this phenomenon has recently been explained as a consequence of entropy (Kaplan et al. 1994; Kestenbaum 1998; Dinsmore et al. 1998).

7.1.4 Fairy rings

[1; 100 %] Fairy rings are circular structures formed by the growth of fungi, particularly the fungus *Marasmius oreades*. [1; 91 %] They grow outward in a circle, starving the grass above, and sometimes to a diameter of 200 m. [1; 92 %] Under the uniform probability interpretation, fairy rings would be considered extremely improbable, and their circular shape would make them specified.

7.1.5 Patterned ground

[1; 87 %] Repeated freeze-thaw cycles in cold environments can generate interesting circular and polygonal patterns. [1; 100 %] Under a uniform probability interpretation, such patterns would constitute CSI: [1; 81 %] yet there is now an explanation involving lateral sorting and "stone domain squeezing" (Kessler and Werner 2003).

[1; 63 %] 8 Evolutionary computation

As mentioned in Sect. [1; 89 %] 2, one of Dembski's principal claims is that evolutionary computation cannot generate CSI. [1; 100 %] This is essentially just a corollary of his Law of Conservation of Information, which as we have seen in the previous section, is invalid. [1; 90 %] More precisely, he concedes that the "Darwinian mechanism" can generate the "appearance" of specified complexity but not "actual specified complexity" (Dembski 2002, p. 183).

In Chap. [1; 96 %] 4 of No Free Lunch, Dembski examines several examples of genetic algorithms and concludes that none of them generate CSI in his sense. [1; 100 %] 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. [1; 83 %] Not surprisingly, in each case, he finds that it has.

[1; 100 %] 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. [1; 92 %] 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 (Dembski 2002, p. [1; 100 %] 215), seems to us completely unwarranted.

[1; 100 %] Dembski considers a number of genetic algorithms: [1; 97 %] variations on Dawkins's METHINKS IT IS LIKE A WEASEL example, an evolution simulation due to Thomas Schneider, an algorithm of Altshuler and Linden for the design of antennas, and an evolutionary programming approach to checkers-playing by Chellapilla and Fogel.

[1; 70 %] 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. [1; 100 %] In Schneider's simulation, it is the error-counting function and "fine tuning" of the simulation itself. [1; 87 %] In the Altshuler-Linden algorithm, it is the "fitness function that prescribes optimal antenna performance" (Dembski 2002, p. 221). [1; 100 %] In the Chellapilla-Fogel example, it is the "coordination of local fitness functions" (emphasis his).

[1; 62 %] It is certainly conceivable, a priori, that Dembski's objections might be correct in the context of his particular measure. [1; 62 %] (However, Schneider (2001) argues they are not in the case of Shannon information.) 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: [1; 100 %] the amount of CSI in the output cannot exceed that in the program and input combined. [1; 100 %] 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. [1; 90 %] This he simply fails to do for each of the examples. [1; 100 %] (The 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. [1; 100 %] As we have remarked previously, this view implies an absurd estimate for the complexity of the fitness function.)

8.1 CSI and computation

Dembski (1999, p. 170) makes a strongly worded claim that no instance of natural causes can produce CSI.

"Since natural causes are precisely those characterized by chance, law or a combination of the two, the broad conclusion of the last section may be restated as follows: Natural causes are incapable of generating CSI. I call this result the law of conservation of information, or LCI for short."

Dembski's claim is broader than he admits. Saying that natural causes cannot produce CSI means that any source processing information by strictly rational means are also barred from producing CSI, including computers, human agents, non-human animals, and disembodied agents. That follows from consideration of deterministic processes, which include symbol manipulation by the rules of logic, as well as many computer algorithms. Given a deterministic process and some input to it, the output is uniquely determined. The amount of information resulting by applying a single deterministic algorithm is bounded by the amount of information in the input, the algorithm, and a small constant. The question then becomes, under what conditions can some agency or process produce substantial new information? A stochastic process is one that uses randomness, leading to the possibility of different outputs each time it is performed. This means that the information concerning a particular output of a stochastic process is not reducible to the information in the process itself and the input to it.

That then leads to another question, can algorithms described in computer science fulfill that role? Yes, they can. As we have seen above in Sect. [1; 63 %] 7, Dembski sometimes claims 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. While we have already noted the facileness with which Dembski adopts

whatever probability distribution best fits his agenda, we can take care to overcome this objection by deploying truly randomized algorithms in response. This means the algorithm uses a source of random numbers, and there is a well-defined probability distribution on the results. We will present two different randomized algorithms that meet and refute a number of objections made by Dembski or others concerning the ability of various forms of computation to produce CSI. We will briefly describe how our algorithms meet the objections, and follow with the technical description of each algorithm.

The first algorithm we call TSPGRID, because it solves the traveling salesman problem (TSP) on a grid layout of cities to be visited by a stochastic process, such as evolutionary computation. The TSP is a well-studied problem in computer science, where the goal is to find the shortest closed path visiting each city in the tour once and only once. The TSP is notable in part because there is no known algorithm that efficiently solves it; the TSP is classified as an NP-hard problem. Neither computers nor humans are privileged when it comes to solving the TSP.

TSPGRID takes one input parameter, n , that determines a grid size to be solved; let us call the total number of cities in a tour k . The output from TSPGRID is a sequence of cities on the grid, which means that there are $k!$ possible tours in the problem space that TSPGRID examines. Given the grid layout, the shortest tour length is shared in common with many different tours. [1; 88 %] TSPGRID avoids the “contingent” objection under one interpretation of specified complexity, because it chooses randomly among all the possible optimal solutions, and there are many of them. There is a bounded range of optimal Hamiltonian cycles of cities on the grid that at once is a large number, but also is a tiny fraction of the total number of possible tours $k!$.

[1; 100 %] Dembski sometimes objects that the CSI produced by algorithms is contained in the program and input. TSPGRID also can demonstrably show that the information contained within the program and its input is much smaller than the CSI of the output. One can select an input, n , such that any optimal tour output has more CSI than the program and input have bits.

The second algorithm we will call Q . It is constructed so that

[1; 72 %] (a) there are many possible outputs, and any particular output of Q occurs with low probability (it is “complex” by Dembski’s standards);

[1; 69 %] (b) every possible output string is specified because it is highly compressible, per Dembski;

[1; 69 %] (c) every output string has a different specification, and no two specifications intersect, that is, every output string is generated by a different program/input pair.

[1; 81 %] Suppose Q on input n generates an output, but we do not know how Q works; [1; 73 %] 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. [1; 91 %] If we assume that we will eventually discover a good compression for v (we could, for example, simply do some dovetailing, a technique well known in computer science), then v is specified, and the probability that the particular specification we discover matches a random output of Q is 2^{-n} . [1; 100 %] Thus, v constitutes CSI, and so every output of Q constitutes CSI.

[1; 100 %] There is a possible objection to this construction, which runs as follows: [1; 79 %] if we assume that we are looking for low Kolmogorov complexity per se in the output string, 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 reasonable background knowledge. To counter this criticism, Q can return output strings that are compressible with respect to some other compression scheme which is easily computable. [1; 100 %] 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. [1; 70 %] 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. [1; 75 %] So the implementation of Q now returns a bit string w for which the run-length encoding of w is shorter than $|w|/100$, or any easily computable function of $|w|$.

[1; 100 %] Now it is easy, upon seeing an output v of Q , to compute its run-length encoding and produce that as a specification. [1; 97 %] (In fact, this is similar to several of Dembski’s examples, such as the Caputo example and the Contact 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. [1; 100 %] So in this case all the specifications would be easily derivable with background knowledge, and they would all be different.

[1; 100 %] 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. [1; 100 %] But this is hardly fair in light of Dembski’s injunction to make the rejection region as small as possible. [1; 63 %] Furthermore, this objection would be like seeing both the Contact prime sequence and the Caputo sequence as outputs of some program and saying, “The specification is just that these strings have short run-length 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 Contact primes sequence.

[1; 73 %] We conclude that Dembski’s claims about natural causes and computation cannot be sustained.

8.1.1 TSPGRID Details

[1; 87 %] The TSPGRID algorithm takes an integer n as an input. [1; 91 %] It then solves the traveling salesman problem on a $2n \times 2n$ square grid of cities. [1; 100 %] Here the distance between any two cities is simply Euclidean distance (the ordinary distance in the plane). [1; 99 %] 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.

[1; 70 %] Göbel has proved that the number of different Hamiltonian cycles on the $2n \times 2n$ grid is bounded above by $c \cdot 28n^2$ and below by $c \cdot 2.538n^2$, where c, c are constants (Göbel 1979). [1; 100 %] We do not specify the

details of how the Hamiltonian cycle is actually found, and in fact they are unimportant. [1; 100 %] 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. [1; 100 %] For the sake of ease of analysis, we assume our algorithm has the property that each solution is equally likely to occur.

[1; 100 %] Now there are $(4n2)!$ [1; 100 %] different ways to list all $4n2$ cities in order. But, as Göbel proved, there are at most $c \cdot 28n2$ different ways to produce a Hamiltonian cycle. [1; 100 %] Let us now compute the specified complexity using the uniform probability interpretation. [1; 87 %] The probability that a randomly chosen list of $4n2$ cities forms a Hamiltonian cycle is $\leq$

[...]

[1; 100 %] and the number of bits of specified information in such a cycle is therefore $\geq$

[...]

[1; 83 %] By Stirling's approximation the number of bits of specified information is bounded below by a quantity that is approximately $8n2 \log_2 n - 2.6n2$.

The CSI produced by TSPGRID can be greater than the information in the program and input. [1; 100 %] 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. [1; 81 %] Since for large n we have $8n2 \log_2 n - 2.6n2 (\log_2 n) + c$, we conclude that TSPGRID has indeed generated specified complexity with respect to the uniform probability interpretation.

8.1.2 Q Details

Here are the details for the implementation of Q: [1; 83 %] 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 (or an easily computable function, as discussed above) $C(w)$ 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. [1; 68 %] This can be accomplished by "dovetailing."

Let $h(n)$ be any computable function of n . The algorithm P works as follows. Based on some choice of computing model (eg, 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.

[1; 75 %] Now we construct our 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. [1; 97 %] (In practice, low-quality random bits can be obtained from environmental sources (eg, counting keystrokes or time between keystrokes) and high-quality random bits can be obtained from physical sources (eg, counting radioactive decays). [1; 72 %] Indeed, there is even a web site, <http://www.fourmilab.ch/hotbits/>, where random bits obtained from a Geiger counter can be downloaded.) Next, Q places a "1" in front of the base-2 representation of t , and treats the result as an integer u . [1; 100 %] (If $t = 5$, or 101 in base 2, then $u = 1101$ in base 2, or 13.) Finally, it outputs $P(u)$. For every different input n , Q outputs a different string, and for large n it becomes highly unlikely that Q will output the same output string more than once if given n as input again.

9 CSI and biology

[1; 93 %] 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. [1; 88 %] In Intelligent Design, Dembski began an attack on evolution which he continues in No Free Lunch. [1; 100 %] However, many of his claims appear suspect.

[1; 100 %] For example, consider Dembski's claims about DNA. [1; 75 %] He implies that DNA has CSI (Dembski 2002, p. [1; 89 %] 151), but this is in contradiction to his implication that CSI can be equated with highly compressible strings (Dembski 2002, p. 144). [1; 100 %] In fact, compression of DNA is a lively research area. [1; 65 %] Despite the best efforts of researchers, only minimal compression has been achieved (Grumbach and Tahi 1994; Schmitt and Herzel 1997; Chen et al. 1999; Lancot et al. 2000; Apostolico and Lonardi 2000; Li 2002).

[1; 100 %] Dembski devotes many pages of No Free Lunch to his claim that the flagellum of *Escherichia coli* contains CSI. We have already noted in Sect. [1; 72 %] 6 that his treatment of specification in this case leaves much to be desired. [1; 100 %] But even if one accepts "outboard rotary motor" as a valid specification, is it true that the *E. coli* flagellum matches this specification? [1; 100 %] There are significant differences. [1; 100 %] To name a few, a human-engineered outboard rotary motor spins continuously, but the flagellum moves in jerks. [1; 100 %] An outboard rotary motor drives a propeller, but the flagellum is whip-like. [1; 100 %] No human engineered outboard rotary motor is composed entirely of proteins, but the flagellum is.

[1; 100 %] Specification is just one half of specified complexity; [1; 100 %] Dembski must also show matching the specification is improbable and thus complex in his framework. [1; 100 %] 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. [1; 100 %] Biologically his calculations verge on the ridiculous, since no reputable biologist believes the flagellum arose in the manner Dembski suggests. [1; 100 %] 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. [1; 93 %] Dembski justifies his approach by appealing to the flagellum's "irreducible complexity", a term coined by fellow intelligent-design advocate Michael Behe. [1; 87 %] But Dembski ignores the fact that sequential evolutionary routes for the flagellum have indeed been proposed (Rizzotti 2000; Pallen and Matzke 2006). [1; 82 %] True, such routes are not as detailed as one might like. [1; 100 %] Nevertheless, they seem far more likely than Dembski's random assembly model. Further, the purported basis of "irreducible complexity" in this case, the uniqueness and interdependence of protein parts in the flagellum, has been shown to be steadily dwindling as research uncovers homologies in other organisms and inessentialness of particular proteins in making functional flagella (Pallen and Matzke 2006).

[1; 100 %] Even taken as a non-evolutionary account of flagellar construction, the specifics of Dembski's approach reveal a number of problems. [1; 97 %] 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):

[1; 78 %] $p_{\text{dco}} = p_{\text{orig}} \cdot p_{\text{local}} \cdot p_{\text{config}}$.

[1; 95 %] 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. [1; 100 %] Dembski's calculation of p_{local} is relatively straightforward:

[1; 75 %] $p_{\text{local}} = (\text{protsys} \cdot \text{subst} / \text{prottotal}) \text{protsys} \cdot \text{copies}$

[1; 94 %] where – protsys is the number of proteins in the system being analyzed;

[1; 77 %] – subst is the number of different proteins which might provide an adequate substitute for each of the proteins in the system;

[1; 100 %] – prottotal is the total number of different proteins available in context; [1; 100 %] and – copies is the number of copies of each protein that will be required to construct the system.

[1; 100 %] The only number that Dembski provides a citation for in this group is the one for prottotal : 4,289. [1; 100 %] The others are either unreferenced or admittedly made-up. [1; 100 %] For example, consider subst . [1; 100 %] The number of possible substitutions is not known, and in any case is quite likely highly variable with different proteins being examined. [1; 87 %] Dembski's equation, though, is exquisitely sensitive to changes in this value. [1; 92 %] 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. [1; 100 %] If the value were 22 or more, the probability resulting would rise above Dembski's universal probability bound of 10^{-150} .

[1; 100 %] 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. [1; 100 %] That this assumption is untrue should come as no surprise to the reader.

[1; 85 %] 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} . [1; 100 %] This concept appears to be original to Dembski. [1; 100 %] 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. [1; 100 %] This in itself is problematic, for biological proteins commonly serve two or more distinct functions. [1; 77 %] No time is wasted by Dembski in considering such empirically verified but mathematically inconvenient sloppiness. [1; 100 %] Dembski's general formula for an approximation of a perturbation probability is

[...]

[1; 100 %] 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. [1; 100 %] Dembski uses the Gettysburg address as an example. [1; 100 %] 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. [1; 92 %] A thousand 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. [1; 100 %] The largest proportion of changes to unchanged text which preserves the meaning corresponds to the perturbation tolerance factor. [1; 100 %] If some of the characters were missing, a recipient would still be able to identify it as the Gettysburg address. [1; 91 %] The largest proportion of missing characters to characters present which permits accurate identification corresponds to the perturbation identity factor. [1; 100 %] Dembski provides arbitrary values of 0.1 and 0.2 for the perturbation tolerance and perturbation identity factors, respectively. [1; 100 %] 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.

[1; 90 %] There are three things to note about these numbers in Dembski's calculation. [1; 90 %] The first is the complete lack of any rigorous justification for the selection of these particular values. [1; 79 %] 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 and suggests that Dembski is far off the mark in his assignment of values (Shannon 1950). [1; 100 %] The second is the extreme sensitivity of Dembski's proffered equation to any change in these values. [1; 100 %] 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. [1; 92 %] 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 for the relevant context. Our third observation is that Dembski's centerpiece calculation based on perturbation probability is wrong. In (Dembski 2002, p. 297) he claims that

[...]

is “on the order of 10–288.” In fact, it is actually about 10–223, an error of about 65 orders of magnitude. (Dembski finally acknowledged this error, more than 3 years after he was informed of it.)

[1; 94 %] Even if Dembski's calculation were right, and his 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 porig for a particular protein. [1; 98 %] Dembski utilizes an analogy of a supermarket stocked with a plenitude of different grocery products. [1; 84 %] Each of those products, he argues, may have its own porig value (Dembski 2002, p. 301). Given Dembski's values for the perturbation tolerance and identity factors, what one finds without much difficulty is that porig for any individual protein of length $\geq 1,153$ is less than Dembski's universal probability bound. Further, any collection of proteins with a combined length $\geq 1,153$ also has porig less than Dembski's universal probability bound. [1; 71 %] 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 $\geq 1,153$ 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 probability bound is not in any sense a threshold. [1; 97 %] Dembski merely argues that a probability below the universal probability bound obviates the need to justify a greater “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.

9.1 Dembski and artificial life

[1; 100 %] 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. [1; 100 %] Artificial life is closer to biological evolution, since the programs have “phenotypic” effects which change through time.

[2; 96 %] The field of artificial life evidently poses a significant challenge to Dembski's claims about the failure of algorithms to generate complexity. [2; 95 %] Indeed, artificial life researchers regularly find their simulations of evolution producing the sorts of novelties and increased complexity Dembski claims are impossible. [1; 100 %] Yet Dembski's coverage of artificial life is limited to a few dismissive remarks. [1; 100 %] Indeed, the term “artificial life” does not even appear in the index to No Free Lunch.

[1; 100 %] Consider Dembski's appraisal of of the work of artificial life researcher Tom Ray:

[1; 100 %] 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). (Dembski 2002, p. 211)

[1; 100 %] We have to wonder how carefully Dembski has read Ray's work, because this is not the conclusion we drew from reading his papers. [1; 100 %] One of us wrote an e-mail message to Ray asking if he felt Dembski's quote was an accurate representation of his work. [1; 100 %] Ray replied as follows:

No. [1; 82 %] I would say that in mywork, there is no strong prevailing trend towards either greater or lesser complexity. [1; 100 %] Rather, some lineages increase in complexity, and others decrease. [1; 100 %] Here, complexity does not correspond to size, but rather, the intricacy of the algorithm.

[1; 100 %] Dembski also does not refer to papers that demonstrate the possibility of increased complexity over time in artificial life: see, for example (Ray 1994, 2001; Adami et al. 2000; Channon 2001). [1; 92 %] 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” (Koza 1994). [1; 83 %] Neither does he mention the complex adaptive behaviors evolved by Karl Sims' virtual creatures (Sims 1994), or the work of Lipson and Pollack (2000), showing how an evolutionary approach can automatically produce electromechanical robots able to locomote on a plane. [1; 100 %] These omissions cast serious doubt on Dembski's scholarship.

After the publication of No Free Lunch, a paper by Lenski et al. (2003) offered another reason to reject Dembski's claims. [1; 100 %] The authors show how complex functions can arise in an artificial life system, through the modification of existing functions.

10 Conclusions

[1; 100 %] We have argued that Dembski's justification for “intelligent design” is flawed in many respects. [1; 100 %] His concepts of complexity and information are either orthogonal or opposite to the use of these terms in the literature. [1; 86 %] His concept of specification is problematic and ill-defined. [1; 100 %] Dembski's use of the term “complex specified information” is inconsistent, and his proof of the “Law of Conservation of Information” is flawed. [1; 68 %] Finally, his claims about the limitations of natural causes and computation are incorrect. [1; 100 %] We conclude that there is no reason to accept his claims.

[1; 86 %] Acknowledgements We are grateful to Anna Lubiw, Ian Musgrave, John Wilkins, Erik Tellgren, Glenn Branch, and Paul Vitányi, who read a preliminary version of this paper and gave us many useful comments. [1; 100 %] We owe a large debt to Richard Wein, whose original ideas have had significant impact on our thinking. [1; 68 %] We thank Norman Levitt for suggesting the pulsars example.