

Fine-Tuning Should Make Us More Confident that Other Universes Exist

Bradford Saad (Global Priorities Institute, University of Oxford)

Forthcoming in *American Philosophical Quarterly*

Abstract: This paper defends the view that discovering that our universe is fine-tuned should make us more confident that other universes exist. My defense exploits a distinction between ideal and non-ideal evidential support. I use that distinction in concert with a simple model to disarm the most influential objection—the this-universe objection—to the view that fine-tuning supports the existence of other universes. However, the simple model fails to capture some important features of our epistemic situation with respect to fine-tuning. To capture these features, I introduce a more sophisticated model. I then use the more sophisticated model to show that, even once those complicating factors are taken into account, fine-tuning should boost our confidence in the existence of other universes.

Keywords: fine-tuning; multiverse hypothesis; inverse gambler’s fallacy; arguments for a multiverse; ideal rationality

Word count: 6,083 (main text)

Acknowledgements: For helpful feedback, I am grateful to participants in the 2022 Rutgers philosophy of religion group—in particular, I recall helpful comments from Frederick Choo, Brian Leftow, Avi Sommer, and Dean Zimmerman; special thanks to Daniel Berntson for extended discussion. I am also grateful to anonymous reviewers for extremely helpful comments, some of which led to significant improvement in the Uncertainty* model.

1. Introduction

Does fine-tuning support the existence of other universes? That is, does the discovery that our universe's fundamental cosmological parameters seem to fall within the exceedingly narrow, life-permitting portion of parameter space confirm the *multiverse hypothesis* (M)?¹ Many philosophers and physicists think so. After all, if there is such a suitably vast and varied multiverse, a fine-tuned universe is to be expected. In contrast, such a universe would seem remarkable and extraordinarily unlikely if there is only one universe—at least absent (something like) a life-favoring designer.² However, according to an influential objection—the *this-universe objection*—this reasoning is fallacious: while the fact that some universe is fine-tuned would be more likely on—and hence confirm—M, our relevant evidence is instead that *this* universe is fine-tuned. But that more specific fact is not more likely on—and hence does not confirm—M.³

The task of this paper is to defend the view that fine-tuning confirms M. My defense will focus on the this-universe objection, since it is by far the most discussed objection to taking fine-tuning to support M and I regard it as the most plausible objection with that target. However, my defense also offers a general recipe for

¹ For an overview of the literature on fine-tuning, see Friederich (2021).

² If this argument goes through conditional on the non-existence of a designer, it also goes through unconditionally, provided that the probability of design hypotheses on which fine-tuning disconfirms M are sufficiently low, i.e. low enough so as not to cancel the confirmation that fine-tuning confers to M if no such design hypothesis is true. In what follows, I assume that the probability of such design hypotheses is indeed sufficiently low. Anyone who would contest this assumption can take my discussion as conditional on it.

³ Variations of this objection are advocated by Hacking (1987), Olding (1990), Dowe (1999), White (2000; 2003), Sober (2004), Draper et al. (2007), Collins (2009), Plantinga (2011), Landsman (2016), Draper (2020), and Goff (2021).

withstanding other attempts to show that fine-tuning does not support M. (I'll use "support" and "confirm" interchangeably.)

My defense starts with a distinction that generally goes undrawn in discussions of fine-tuning and the this-universe objection. The distinction concerns two ways of understanding "does fine-tuning support M?":

Ideal Support Question: Would our fine-tuning evidence justify an *ideally rational agent*—i.e., an agent that is not susceptible to errors in reasoning and whose memory and reasoning capacities are unlimited—in boosting her confidence in M?

Non-Ideal Support Question: Does our fine-tuning evidence justify *us*—i.e., non-ideal agents who are trying to figure out what to think in light of fine-tuning—in boosting our confidence in M?

For short, I'll use "ideal confirmation" to talk about confirmation relative to an ideally rational agent and "non-ideal confirmation" to talk about confirmation relative to non-ideal agents like ourselves. In what follows, I'll argue that even if fine-tuning fails to ideally confirm M, it still non-ideally confirms M—at least for all the this-universe objection says.⁴

Here's the plan. §2 motivates the distinction between ideal and non-ideal support. §3 motivates a project of answering the Non-Ideal Support Question. §4 argues that upon adopting an ideal support reading of the this-universe objection, we should be uncertain whether it is sound. §5 presents a simple model of our epistemic

⁴After this paper was accepted, I learned that Dorst and Dorst (forthcoming) argue that fine-tuning should boost our confidence in a design hypothesis, given that we should be uncertain about which of two indifference principles would apply to ideally rational agents. While the two arguments use uncertainty about ideal rationality and concern what fine-tuning supports, they differ in various ways. For example, they have different, though compatible, conclusions about what fine-tuning supports; indifference principles play no role in my argument; and the this-universe objection is set aside in theirs.

situation—which includes rational uncertainty about the this-universe objection—and uses the model to argue that our fine-tuning evidence non-ideally confirms M. Constructing the simple model sets the stage for addressing complicating factors that are not captured by the model (§6) and introducing a better model that takes these factors into account (§7). §8 argues that, even after these factors are taken into account, fine-tuning should still lead us to be more confident in the existence of other universes.

2. Ideal Versus Non-Ideal Support

My argument will rely on the noted distinction between ideal and non-ideal support. However, ideal and non-ideal support are usually not distinguished in discussions of fine-tuning.⁵ So one might doubt that this is a serviceable distinction.

My response is that reflection on mundane cases yields a grip on the distinction. To take a stock example, suppose you are reasonable in having a middling credence in whether a certain mathematical conjecture (Goldbach's conjecture, say) can be proven. You then read from your most trustworthy news source that the Fields Medal is being awarded this year for the proof of that conjecture. Clearly, this news makes it rational for you to increase your credence in the provability of that conjecture. This is so even on the assumption that an ideally rational agent would always be rationally certain about whether the conjecture is provable and hence never receive evidence that merits a boost in their confidence on this score. Of course, there is theoretical work to be done in elucidating the distinction. But the grip on the distinction afforded by such cases will suffice for the purposes of this discussion.

⁵ They are sometimes distinguished in epistemology. For example, see Field (2000: 117), Schoenfield (2012: *passim*), and Smithies (2015).

For anyone who remains skeptical of non-ideal support, I suggest the following as a way to recast my discussion: read my claims about non-ideal support as claims about ideal support relative to bodies of evidence that warrant belief in one's being a non-ideal agent and my claims about ideal support as claims about ideal support relative to bodies of evidence that warrant belief in one's being an ideal agent.⁶ I believe that my argument goes through just as well on this framing and that little hangs on which of these two framings is adopted.

3. The Two Questions and Two Projects

We've seen that there is a distinction between ideal and non-ideal support. Thus, the Ideal and Non-Ideal Support Questions are distinct. Even so, there is a notable connection between them: settling what an ideal agent should make of our fine-tuning evidence is a way for us to settle what we should make of it. While the Ideal Support Question is also of intrinsic interest, I primarily find it interesting because of its connection with the Non-Ideal Support Question. Hence there is a project of trying to answer the Ideal Support Question as a means to answering the Non-Ideal Support Question. By and large, participants in the debate about the this-universe objection can be charitably interpreted as engaging in this project.

However, I submit that this is not the only project in the vicinity that is worthy of engagement. There's also the project of evaluating what we should make of our fine-tuning evidence given an informational background that includes uncertainty about whether the this-universe objection succeeds. At any rate, this project is worth engaging in provided that some of the people who are curious about what to make of

⁶ Cf. Christensen (2008; 2010).

fine-tuning either are or should be uncertain about whether fine-tuning ideally supports M.⁷

4. We Should Not Be Certain that the This-Universe Objection Is Sound

I've suggested that the debate between proponents of the this-universe objection and their opponents can be understood as concerning ideal support in the first instance. Given this understanding, we should be uncertain whether the this-universe objection is sound.⁸ This is so for two reasons.

First, despite extensive investigation, there remains a deep and recalcitrant divide among relevant experts who work on the objection about whether it succeeds. Everyone who advocates the this-universe objection thus has higher-order evidence that the reasoning in the objection is mistaken.⁹ In light of such evidence, we should be uncertain whether the this-universe objection is sound—or at least this is so given a rather weak form of conciliationism about epistemic rationality.¹⁰

⁷ N.B. There is no need to choose between these projects. While I am here only investigating what to make of fine-tuning given uncertainty about the Ideal Support Question, I would welcome further work aimed at resolving that uncertainty.

⁸ More precisely: we should be all-things-considered uncertain about whether the this-universe objection succeeds. This is compatible with our having rational insulated convictions about the matter that result from considering our first-order evidence while bracketing higher-order evidence from, for example, disagreement. Arguably, such insulated attitudes play an important role in collective truth-seeking enterprises such as philosophy—see Barnett (2019). I hereafter leave this qualification implicit.

⁹ A reviewer notes: while the this-universe objection is primarily debated among philosophers, relevant experts include physicists and cosmologists; however, the latter rarely address the objection when discussing the multiverse hypothesis and tend not to agree with it when they are presented with it (Manson, (2020)). I take this fact to be an additional piece of evidence that tells against certainty in the soundness of the this-universe objection.

¹⁰ For defenses of conciliationism, see, for example, Christensen (2007) and Elga (2007). The form of conciliationism required here is weak in that it need not apply universally or require according equal weight to epistemic peers in cases where it does apply.

Admittedly, this is not the only possible moral of the disagreement among relevant experts. For instance, disagreement can be taken as evidence that the debate is ill-posed or merely verbal. Notice, however, that there is no tension between drawing that moral and taking the disagreement to require uncertainty about the soundness of the this-universe objection. Here, as elsewhere, disagreement can at once support the hypothesis that a debate is defective and mandate uncertainty in views at issue in the debate.¹¹

Second, to be certain that the this-universe objection is sound, one would need to be certain that every reply to it that has been defended in the literature fails. However, the literature contains a range of seemingly independent replies that both enjoy at least a modicum of plausibility and turn on difficult philosophical matters about which we should be uncertain.¹² Hence, we should not be certain that every reply fails, and so should not be certain that the this-universe objection succeeds.

My argument does not turn on the specific content of replies to the this-universe objection. But to see that it would be unreasonable to be certain that every reply fails, it may help to consider some of them. Some replies contend that the this-universe objection is at best inconclusive, since advocates of the objection have not adequately defended the premise that the multiverse hypothesis fails to raise the probability of our

¹¹ Taking disagreement as evidence that the debate is defective does cast doubt on the enterprise pursued by this paper of assessing whether the this-universe objection succeeds. This doubt is encouraged by two themes of this paper: the important but neglected role of the distinction between ideal and non-ideal support in the debate and the need for careful attention to choice of informational background relative to which support is evaluated. But it should be borne in mind that the relevant experts seem to be in wide agreement that the debate is well-posed and substantive (though see Friederich (2019a)). As a result, even if the disagreement provides some evidence that the debate is ill-posed or merely verbal, I think it will remain more plausible that the debate is well-posed and substantive. In any case, I will assume this in what follows.

¹² For a review of the literature, see Manson (2022).

(total relevant) evidence.¹³ Developments of this response identify versions of the multiverse hypothesis or informational backgrounds on which the existence of multiple universes would raise the probability of our fine-tuning evidence.¹⁴ Other replies argue on independent grounds that we should reject the total evidence requirement in favor of the predesignation requirement (to evaluate evidence in a manner that can be specified before acquiring that evidence) and show how adopting the latter blocks the objection.¹⁵ Another reply is that even if M does not raise the probability of this universe being fine-tuned, it does raise the probability of *our* observing it to be fine-tuned, which is enough for our cosmological fine-tuning evidence to support M.¹⁶ Still other replies give reasons for thinking the objection goes wrong somewhere. For instance, it's hard to see how the multiverse hypothesis could (as the objection grants) raise the probability of some universe being fine-tuned without raising the probability of a particular universe being fine-tuned; yet, evidently, the this-universe objection could be run with respect to any particular universe.¹⁷ Another reply in this genre is that the objection overgeneralizes.¹⁸ Others have motivated approaches to self-locating belief on independent grounds and argued that they predict that our being in this cosmologically fine-tuned universe supports M.¹⁹ Yet another response is that our fine-tuning evidence

¹³ See Friederich (2019a), Isaacs et al. (forthcoming), Juhl (2005), and Saad (forthcoming).

¹⁴ See Juhl (2005); cf. Bradley (2009).

¹⁵ See Epstein (2017) and Manson and Thrush (2003); cf. Dawkins (1986: Chapter 1). For responses, see Barrett and Sober (2020) and Draper (2020).

¹⁶ See Bradley (2009) and Huemer (2018: Ch. 11).

¹⁷ See Bostrom (2002: 22).

¹⁸ See Manson and Thrush (2003), Juhl (2005), and Epstein (2017); cf. Friederich (2019a).

¹⁹ See Bradley (2012) and Isaacs et al. (forthcoming).

supports M by enhancing its theoretical virtues and amplifying the theoretical vices of the single universe hypothesis.²⁰

While I'm not certain of any of these responses, I find most of them at least somewhat plausible. Your plausibility judgements may differ. Even if so, we should still agree that it would be unreasonable to be certain that every such response fails. So we should not be certain that the this-universe objection is sound.²¹

5. A Simple Model

Proponents of the this-universe objection hold that fine-tuning provides no support for M. Their opponents claim that fine-tuning provides some support for M. Interpreting this debate as concerning ideal support, I've claimed that we should be uncertain about which side is correct. My next task is to offer a simple model of our epistemic situation with respect to fine-tuning and M. After presenting the model, I'll use it to argue that fine-tuning non-ideally supports M. In later sections, I'll note some complicating factors that the simple model does not capture, modify the model to capture them, and argue that fine-tuning also supports M by the lights of the modified model.

The simple model uses the below case in which uncertainty about the outcome of one coin flip stands in for uncertainty about M, uncertainty about the outcome of another coin flip stands in for uncertainty about the ideal support targeted by the this-universe objection, your awakening stands in for finding ourselves in a fine-tuned universe, and dice rolls stand in for universes.

Uncertainty: You know:

²⁰ See Friederich (2019b); cf. Leslie (1989) and Parfit (2011: Appendix D).

²¹ This is a cumulative case argument. While I do not know of any other cumulative case arguments that have been used to support the multiverse hypothesis (but see Lewis (1986)), such arguments have been used in efforts to support its main rival, the design hypothesis—see, for example, Swinburne (2004), Draper (2010), and Audi (2017).

- A fair coin was flipped to determine the size of a team of dice rollers, each member of which will roll a pair of dice.
 - If it lands on the side labeled “ONE ROLLER”, the team has just one member.
 - If it lands on the side labeled “TWO ROLLERS”, the team has exactly two members.
- Another fair coin was flipped.
 - If that coin lands on the side labeled “NEUTRAL”, then you are assigned a roller; you will wake up just in case she rolls a double six.
 - If that coin lands on the side labeled “CONFIRMATION”, you will wake up just in case at least one dice roller rolls a double six.

You awaken.²²

Notice that conditional on NEUTRAL, your waking up is probabilistically independent of how many people are in the group. Hence, conditional on NEUTRAL, your evidence neither confirms nor disconfirms TWO ROLLERS. In contrast, given CONFIRMATION, it’s clear that the more rollers there are, the more likely it is that you’ll wake up—hence, conditional on CONFIRMATION, your waking up supports TWO ROLLERS. In general, if you are rationally uncertain which of exactly two outcomes obtains and your evidence confirms a hypothesis conditional on one of the outcomes and neither confirms nor disconfirms the hypothesis conditional on the other outcome, then your evidence confirms that hypothesis.²³ Therefore, in this case, your evidence confirms the hypothesis that there were multiple rollers.

²² This case combines two cases from White (2000; 2003), one of which is due to McGrath (1988: 265). I’ll discuss these below. For ease of exposition, I’ll focus on White’s rendition of McGrath’s case rather than the original.

²³ Some readers have objected that this principle overgeneralizes by favoring one side in a wide range of philosophical debates. In reply, I think a relatively small portion of philosophical debates chiefly concern or turn on whether a particular piece of evidence confirms a hypothesis. So I think the principle’s interesting ramifications for other philosophical debates are limited; nor am I aware of any applications of the principle that clearly generate an incorrect result.

This general analysis holds even if we dispense with certain features of the case: the fairness of the coins, the fairness of the dice, the number of rollers being one or two. These dispensable features are included for concreteness and ease of illustration. Readers who are so inclined may wish to modify the case to reflect their degree of uncertainty about whether fine-tuning ideally confirms M and their views concerning the potential number of universes. That said, let's illustrate the general analysis by running the numbers.

To start, notice that there's a 25% probability of each of the following: NEUTRAL and ONE ROLLER, NEUTRAL and TWO ROLLERS, CONFIRMATION and ONE ROLLER, and CONFIRMATION and TWO ROLLERS. Conditional on NEUTRAL and ONE ROLLER, there is a $1/36$ probability that you will be woken up. The same is true conditional on NEUTRAL and TWO ROLLERS and on CONFIRMATION and ONE ROLLER. However, conditional on CONFIRMATION and TWO ROLLERS, the probability that you will be woken up is $1 - (35/36 \times 35/36)$. So, the probability that you will be woken up in ONE ROLLER is $(.25 \times 1/36) + (.25 \times 1/36) = \sim 1.4\%$. And the probability that you will be woken up in TWO ROLLERS is $(.25 \times 1/36) + (.25 \times (1 - (35/36 \times 35/36))) = \sim 2.1\%$. Thus, your being woken up is more probable on the hypothesis that there were multiple rollers; so, it should lead you to boost your confidence in that hypothesis.

Our situation is analogous to Uncertainty. Granting that we should reserve some credence for the hypothesis that the this-universe objection is sound, we should reserve some credence for our situation being like that of the person who is unknowingly in the NEUTRAL condition. Likewise, we should reserve some credence for the this-universe

objection failing, and hence for our situation being like the person who is unknowingly in the CONFIRMATION condition. Therefore, since our situation is analogous to yours in Uncertainty and your evidence supports the multiple rollers hypothesis, our fine-tuning evidence supports M.

While I think this argument is essentially correct, I'll offer a more careful version of it to address some complications in §8. But two points are worth noting at this juncture. One is that the argument generalizes: if there are cases to be made in addition to the this-universe objection for thinking that fine-tuning is evidentially irrelevant to M, then Uncertainty will also model our situation with respect to them—provided, of course, that we should be uncertain about them. Indeed, Uncertainty will model our situation with respect to such arguments taken collectively, provided that we should be uncertain whether at least one of them succeeds.

The second point is that Uncertainty importantly differs from other cases that have been offered as models of our epistemic situation. To illustrate, consider an instructive pair of cases from White and McGrath. They disagree about which of the following captures our situation with respect to our fine-tuning evidence.

Case B: Jane knows that an unspecified number of players will simultaneously roll a pair of dice just once, and that she will be woken if, and only if, a double six is rolled. Upon being woken she infers that there were several players rolling dice.

Case B*: Jane knows that she is one of an unspecified number of sleepers each of which has a unique partner who will roll a pair of dice. Each sleeper will be woken if and only if *her* partner rolls a double six. Upon being woken, Jane infers that there are several sleepers and dice rollers. [emphasis White's (2003: 236-7)]

In their essentials, these cases are tantamount to the CONFIRMATION and NEUTRAL conditions in Uncertainty. Neither case models our rational uncertainty about whether the this-universe objection succeeds. If White and McGrath are concerned with the Non-Ideal Support question, each of these cases thus misses the mark. If White and McGrath are instead concerned with the Ideal Support Question, then the cases' failure to model our rational uncertainty may well be appropriate. But then a different model is needed to capture our uncertainty about the this-universe objection. On this score, Uncertainty is at least a step in the right direction.

6. Some Complications

My argument used a simple case to model our uncertainty about the this-universe objection. That model ignored several complications that I'll address in this section:

- Whether the fact that we're observers is included in the informational background relative to which the import of fine-tuning evidence is evaluated,
- whether fine-tuning plays a crucial role in enabling our evidence to support M, and
- the model's failure to capture uncertainty about whether fine-tuning *disconfirms* M.

We'll see that taking these factors into account will require adjustments to the argument. While undertaking these adjustments will provide a deeper understanding of how fine-tuning non-ideally supports M, the main moral of my argument—that fine-tuning non-ideally supports M—will remain intact.

These three factors may seem separate. But they bear on fine-tuning's import in an intertwined fashion. We can begin to disentangle them by considering a passage from White (2003) in which all three are operative. The passage is part of White's defense of his view that fine-tuning does not support M. It comes in response to the following objection:

The more universes there are, the more living creatures there are. So the more opportunities I had to be picked out of the pool of “possible beings,” and hence the greater the likelihood that I should be observing anything. (*ibid*: 244)

White responds as follows.

...we can see that something must be wrong with this line of reasoning... If the current objector’s argument is cogent, then... regardless of... fine-tuning... we could still argue along these lines that the more universes there are the more opportunities I had for existing and observing, and hence that my observations provide evidence for multiple universes.

Indeed, if the objector’s argument is sound, then the discovery that a universe must meet very tight constraints in order to support life should *diminish* the strength of the case for multiple universes. For if every universe is bound to produce life, then by increasing the number of universes we rapidly increase the number of conscious beings, whereas if each universe has a slim chance of producing life, then increasing the number of universes increases the number of conscious beings less rapidly, and hence (by the objector’s argument) increases the likelihood of my existence less. I would be surprised if anyone wants to endorse an argument with these consequences, but, at any rate, it is not the standard one that takes the *fine-tuning* data to be crucial in the case for multiple universes. (*ibid*: 244-5; emphasis his)

Let’s consider the three factors in turn. After doing so, I’ll modify the simple model to take them all into account.

Informational background. Notice the informational background relative to which White’s imagined objector is contending that observations support M: they are in effect considering the import of updating on observations relative to priors that leave open whether one makes any observations. The informational background is often left unspecified in discussions of fine-tuning. But this exchange illustrates how choice of informational background can be important for evaluating the import of fine-tuning.²⁴ After all, the objector’s argument is unavailable relative to priors on which it is settled

²⁴ Others have emphasized this point—see, for example, Juhl (2005) and Friederich (2019b).

that one makes observations—for a multiverse cannot raise the probability that one does so relative to priors on which it is certain that one makes observations.²⁵ Further, in practice, people who are drawn to taking fine-tuning to support M learn about fine-tuning against an informational background that includes knowledge that they make observations.²⁶ So we’d expect a fine-tuning argument that captures what draws people to thinking fine-tuning supports M to be compatible with that knowledge. The objector’s argument frustrates this expectation.

Herein lies a limitation of the simple model, as well as White’s and McGrath’s cases. In each of these cases, the subject’s antecedent informational background leaves open whether she will make an observation while building in the fact that if she makes an observation, it will require the analog of fine-tuning: a double six. Thus, none of these cases serve to model learning about fine-tuning against an informational background that encodes the fact that we’re observers.

Admittedly, given that the import of a body of evidence is insensitive to the order in which pieces of evidence are acquired, we should arrive at the same doxastic state by updating on fine-tuning given an informational background that includes the fact that we’re observers as we would by updating on both fine-tuning and the fact that we’re

²⁵ This is not to say that we cannot consult “ur” priors on which it is uncertain whether we are observers in order to assess the evidential import of our being observers. Indeed, some such consultation may be needed to solve the problem of old evidence—see, for example, Juhl (2007) and Isaacs et al. (forthcoming). Here my point is the modest one that since the objector’s argument aims to support M by showing that M raises the probability of our being observers, the argument does not work when evaluated against a background on which M cannot raise that probability because it is already one.

²⁶ Friederich (2019b) defends another fine-tuning argument that takes life-friendliness as part of the background against which the import of fine-tuning is evaluated. Despite this similarity, our arguments are different beasts. Rather than relying on the distinction between ideal vs. non-ideal support, his argument contends that fine-tuning confirms M by eroding a predictive advantage of the single universe hypothesis. For objections to Friederich’s argument, see Metcalf (2021).

observers against the same informational background modulo the fact that we're observers. Here, however, we're investigating whether acquiring fine-tuning evidence should lead to increased confidence in M, not which level of confidence we should have in M given our fine-tuning evidence. So, the commutativity of evidence is beside the point: these cases are inapt to model the impact of fine-tuning evidence against an informational background that includes the fact that we're observers.

According fine-tuning a crucial role. White is right that standard fine-tuning arguments take the fine-tuning part of our evidence to play a crucial role and that fine-tuning plays no such role in the objector's argument.²⁷ Further, a slight extension of the objection shows that to mount an argument for M in which fine-tuning plays a crucial role, it is not enough to show that a body of evidence includes fine-tuning and confirms M. For the objector's argument can be run by evaluating *any* observation relative to priors that leave open whether one makes observations. Hence, while the objection can be run using the observation that our universe is fine-tuned, it equally well could have been run using the observation that our universe is not fine-tuned if we had observed that. A challenge for fine-tuning enthusiasts is therefore to show not only that a body of evidence that includes fine-tuning supports M but also that the fine-tuning component of the evidence is crucial in securing the support.

In the simple model and White's and McGrath's cases, the subject simultaneously learns both a fine-tuning fact *and* the fact that she is making some observation rather than none. The objector's argument points to a way in which the fact that one is an observer is potentially relevant to M. That fact is thus a potentially confounding factor,

²⁷ For further discussion of this issue, see Isaacs et al. (forthcoming).

one that we need to tease apart from fine-tuning in order to assess whether fine-tuning plays a crucial role in supporting M.

Uncertainty about Disconfirmation. The simple model conjures confirmation from uncertainty. What enables the conjuring feat is an asymmetry in the model: uncertainty between confirmation and non-confirmation but not disconfirmation. However, it might be objected that our epistemic situation is instead a symmetric one in which we should be uncertain both about whether fine-tuning confirms M and about whether fine-tuning disconfirms M (as well as about whether fine-tuning has no evidential bearing on M). After all, as a rule our epistemic limitations demand uncertainty about non-trivial matters. And whether fine-tuning disconfirms M is not a trivial matter. Therefore, a better model of our situation would be a symmetric one in which we are uncertain between three outcomes, one of which results in confirmation of a multiple rollers hypothesis, another of which results in disconfirmation of the multiple rollers hypothesis, and one of which results in neither confirmation nor disconfirmation of the multiple rollers hypothesis. But nothing in this model licenses a rational boost in your confidence in the multiple rollers hypothesis.

In response, I grant that we cannot rule out the disconfirmation hypothesis with certainty. Nonetheless, I maintain that there is an asymmetry in our situation and so I deny that the proposed symmetric model is a better model of our epistemic situation than the simple (asymmetric) model. The simple model partially captures an asymmetry in the debate about whether fine-tuning supports M: the debate concerns whether fine-tuning confirms M or whether, instead, fine-tuning lacks evidential bearing on M. That fine-tuning disconfirms M is a possibility that is almost always passed over in

silence. This is no accident. The hypothesis that fine-tuning confirms M enjoys prima facie plausibility and seems defensible on reflection (witness the defenses of it cited in §4). In contrast, the hypothesis that fine-tuning disconfirms M is prima facie implausible and, on reflection, there seems to be very little to be said in its favor.

As far as I know, the literature contains no explicit argument for the disconfirmation hypothesis. The sole implicit argument for the disconfirmation hypothesis that I am aware of can be gleaned from the White passage quoted above.²⁸

For ease of reference, I reprint the relevant portion here:

if the objector's argument is sound, then the discovery that a universe must meet very tight constraints in order to support life should *diminish* the strength of the case for multiple universes. For if every universe is bound to produce life, then by increasing the number of universes we rapidly increase the number of conscious beings, whereas if each universe has a slim chance of producing life, then increasing the number of universes increases the number of conscious beings less rapidly, and hence (by the objector's argument) increases the likelihood of my existence less. (White, 2003: 244-5)

White frames the argument in terms of diminished strength rather than disconfirmation. But the argument can be recast to concern disconfirmation through a suitable choice of informational background. White is responding to an objector who assumes an informational background that leaves open whether we make observations. The argument grants that our observations confirm M relative to that background and then contend that the fine-tuning aspect of our evidence diminishes the strength of that confirmation. We've seen above that the import of fine-tuning should be evaluated against a background that includes the fact that we're observers. If updating on fine-tuning against the former background diminishes the extent to which our being

²⁸ But see Juhl (2005: 344) for a fanciful case that he uses to show that fine-tuning could tell against M relative to some informational backgrounds.

observers confirms M, then against the latter it should decrease our confidence in M. In other words, in shifting from the one background to the other, a decrease in confirmation becomes disconfirmation. An analogy: suppose I cast aspersions on the character of someone testifying against a defendant. The testimony provides evidence of guilt. Relative to pre-testimonial evidence, my aspersions decrease that confirmation. Relative to the credibility of the guilty hypothesis just after the testimony, my aspersions disconfirm that hypothesis.

So, White's argument for diminished strength is equivalent to an argument for disconfirmation. White stops short of endorsing this argument, as it relies on assumptions that White is spotting the objector. However, the argument fails for a reason White does not consider: the argument illicitly focuses on how fine-tuning affects the likelihood of his evidence on M while ignoring how fine-tuning affects the likelihood of his evidence on the single-universe hypothesis. This is illicit because fine-tuning affects both likelihoods and the ratio between them is what determines what the evidence (dis)confirms. As Isaacs et al. (forthcoming) concisely explain:

White's reasoning only gives half the story: it is correct that the smaller the fine-tuning parameter, the smaller the expected number of conscious beings given a multiverse. But also, in the same way, the smaller the fine-tuning parameter, the smaller the expected number of conscious beings given a single universe.... These factors cancel...

In light of the arguments for fine-tuning confirming M, failure of White's argument for disconfirmation, and the absence of other arguments for disconfirmation, I conclude that we are provisionally warranted in taking it to be more plausible that fine-tuning ideally confirms M than that it ideally disconfirms M, and, indeed, that the expected

degree of ideal confirmation exceeds the expected degree of ideal disconfirmation.²⁹

This is a genuine epistemic asymmetry. It shows that the objection from symmetry is mistaken. But the asymmetry is also far from fully captured by the simple model. In the next section, I'll offer another case that better models this asymmetry.

7. A Better Model

In §6, we identified three desiderata for a defense of fine-tuning supporting M: showing that fine-tuning evidence yields support for M against an informational background that includes the fact that we are observers, showing that fine-tuning plays a crucial role in enabling that support, and taking into account uncertainty about disconfirmation. I'll now show how these desiderata can be met by modifying Uncertainty as follows:

Uncertainty*: You awaken. You know:

- A fair coin has been flipped to determine the size of a team of dice rollers.
 - If it lands on the side labeled "ONE ROLLER", the team has just one member. (50%)
 - If it lands on the side labeled "TWO ROLLERS", the team has exactly two members. (50%)
 - You will observe a dice roll that contributed to your awakening.
 - A three-sided coin was flipped. The possible outcomes were as follows.
 - N: you are assigned a roller; you will be woken up just in case her roll meets the awakening condition. (40%)
 - C: you will be woken up just in case at least one dice roller's roll meets the awakening condition. (40%)
 - D: you will be woken up just in case every dice roller's roll meets the awakening condition. (20%)
- You then observe a double six.
- You discover that another fair coin was flipped to determine whether the conditions for awakening would be coarse-grained or fine-grained.
- If the coin landed on the side labeled "COARSE", then your awakening required a non-(double six). (50%)

²⁹ The *expected* degree of (dis)confirmation of E for H is the average of the confirmation and disconfirmation of H by E posited by hypotheses we countenance weighted by our confidence in those hypotheses.

- If the coin landed on the side labeled “FINE”, then your awakening required a double six. (50%)
-You conclude from your observation and discovery that the coin landed FINE and hence that your awakening required a double six.³⁰

This case differs from Uncertainty in two respects that enable it to satisfy the noted desiderata.

First, in Uncertainty*, you begin awake and only later acquire evidence that stands in for fine-tuning, namely learning that the coin landed FINE. Uncertainty* mirrors the actual order of discovery: we discover that we’re observers before finding out about fine-tuned physical parameter values and before learning that they’re fine-tuned.³¹ Uncertainty* thereby models the kind of informational background against which we want to evaluate fine-tuning: we want to know how fine-tuning bears on M given an informational background that includes the fact that we’re observers. This feature of Uncertainty* also enables it to model evidentiary work that fine-tuning is doing in our fine-tuning evidence: since the part of the evidence corresponding to fine-tuning comes after the part corresponding to the fact that we’re observers (your awakening), any boost to the multiple rollers hypothesis in the later stage will be due to the part of your evidence concerning fine-tuning. Thus, if Uncertainty* otherwise adequately models our epistemic situation with respect to fine-tuning, we can use

³⁰ To capture residual uncertainty about whether the physical parameters are fine-tuned rather than coarse-tuned, one could set the fine and coarse awakening conditions so that they overlap (e.g. by making the coarse condition anything except a double one). For simplicity, I bracket such uncertainty by using non-overlapping conditions in the model.

³¹ Whereas some agents learn the values of the fine-tuned parameters and that they are fine-tuned, others merely learn that there are parameters with certain values (without learning what they are) and that they are fine-tuned. Without substantively affecting the analysis, Uncertainty* could be adapted to the latter case by supposing that you observe a dice outcome without being able to tell what it is and learn that just that outcome (whatever it is) is the FINE awakening condition.

whether the multiple rollers hypothesis accrued support from the coin landing FINE as a guide to whether M is non-ideally confirmed by our fine-tuning evidence in virtue of its fine-tuning component.

Second, Uncertainty* models three possibilities concerning ideal confirmation rather than two. More specifically, it models M's being neither ideally confirmed nor ideally disconfirmed (with N), M's ideal confirmation (with C), and M's ideal disconfirmation (with D). Uncertainty* thus improves on Uncertainty by modeling uncertainty about the ideal disconfirmation of M. However, recall that such disconfirmation needs to be modeled with care: it would be a mistake to model it using a case that treats confirmation and disconfirmation symmetrically. Instead, as argued in §6, we need a case that captures the plausibility edge enjoyed by the ideal confirmation hypothesis over the ideal disconfirmation hypothesis. Uncertainty* meets this need by modeling ideal confirmation with an outcome that is twice as likely as the outcome that models ideal disconfirmation.³² By my lights, this is a conservative estimate of the plausibility edge—but the argument should go through just as well with higher or lower estimates, provided that one grants that there is an edge. More generally, readers should feel free to adjust the assigned chances in the scenario to reflect their confidence levels in the corresponding hypotheses, subject to the constraints for which I've argued.

8. Fine-Tuning Still Supports Boosting Our Confidence in Other Universes

We've seen how Uncertainty* improves on Uncertainty by taking into account several complicating factors in our epistemic situation. Our final task is to determine what you

³² And by not using a disconfirmation outcome whose expected degree of disconfirmation swamps the expected degree of confirmation associated with the confirmation outcome.

should think in Uncertainty* and bring that to bear on what we should make of our fine-tuning evidence.

To start, we partition the outcome space and calculate outcome probabilities.

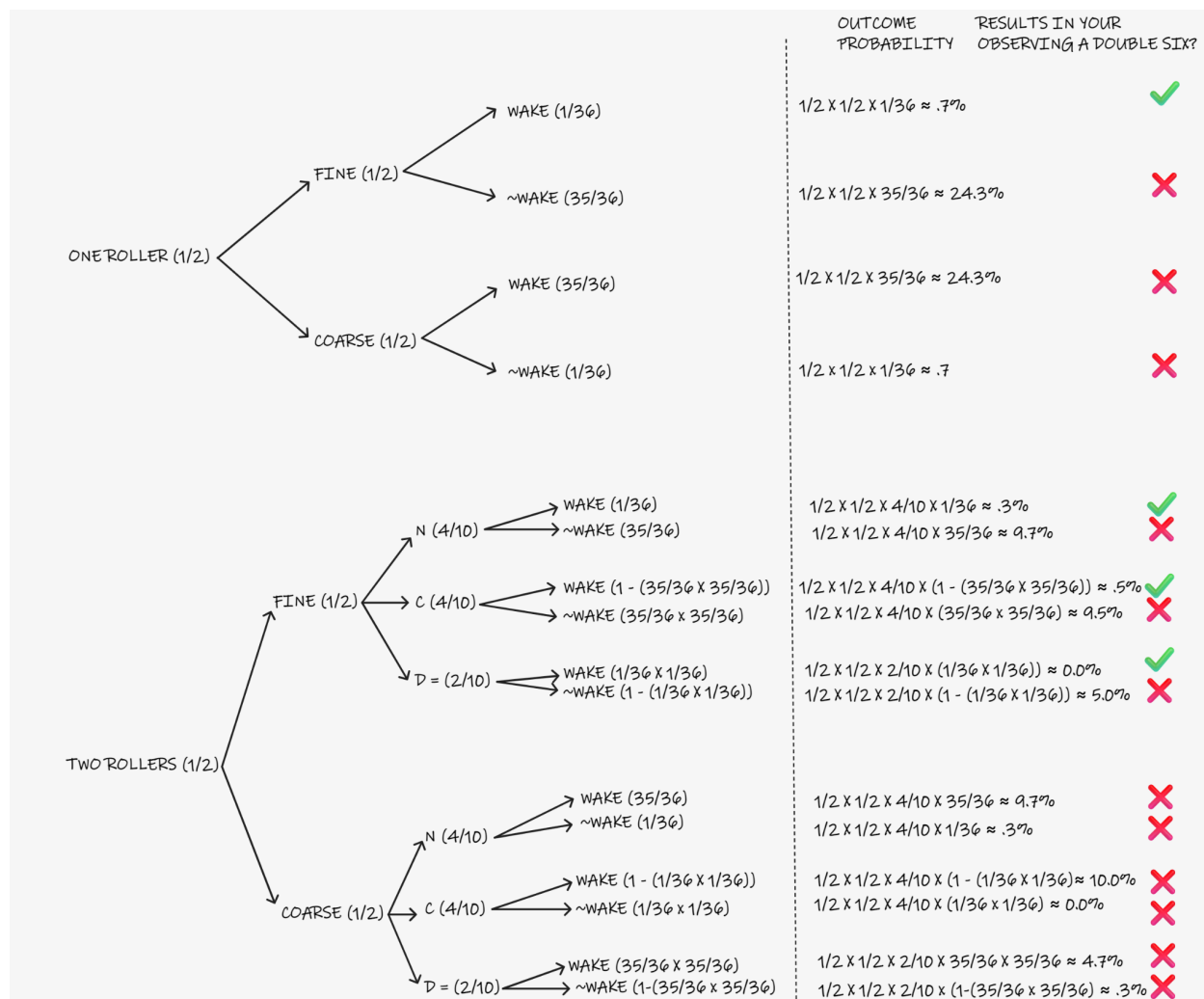


Figure 1.

Next, to test whether discovering that the coin landed FINE supports the multiple rollers hypothesis, you might think we just need to check whether a FINE awakening is more likely on TWO ROLLERS than it is on ONE ROLLER. In fact, it is: .8% (= .3% + .5%) vs. .7%. However, for several reasons, this suggestion is untenable. First, the test assesses evidential import against an informational background that

leaves open whether you are awake. But, as I have been at pains to emphasize, we want to determine whether your evidence supports the many rollers hypothesis against an informational background that includes the fact that you are awake. Second, we also want to determine whether the fine-tuning component of the evidence plays a crucial role in securing support for TWO ROLLERS. The suggested test is silent on this. Moreover, relative to an informational background that leaves open whether you are awake, such crucial support is in fact lacking: a COARSE awakening is *also* more likely on TWO ROLLERS than on ONE ROLLER: 24.4% ($= 9.7\% + 10\% + 4.7\%$) vs. 24.3%.³³

So we need a different test: we need to ask, *after updating on the fact that you are awake*, does learning FINE support the multiple rollers hypothesis? Well, you wake up in (roughly) 50.2% of the outcome space—25% of the space in ONE ROLLER and 25.2% in TWO ROLLERS. So your post-awakening priors in ONE ROLLER and TWO ROLLERS are respectively 49.8% ($= 25/50.2$) and 50.2% ($= 25.2/50.2$). Now, suppose you observe a double six and hence that the coin landed FINE. The portions of the outcome space in which this happens in ONE ROLLER and TWO ROLLERS are respectively .7% and .8%. So, upon learning FINE you should boost your confidence from 50.2% in the multiple rollers hypothesis to 53.3% ($= .8/(.7 + .8)$). In contrast, if you had observed a non-(double six) and hence been in a COARSE awakening, then you should have decreased your confidence in TWO ROLLERS from 50.2% to 50.1% ($= 24.4/(24.3 + 24.4)$). Thus, after updating on the fact that you are awake, learning FINE rather than COARSE should boost your confidence in the multiple rollers hypothesis.

³³ Admittedly, FINE confirms the TWO ROLLERS *more* than would COARSE. But this only shows that, relative to an informational background that leaves open whether you are awake, the fine-tuning component of your evidence is crucial to the *degree* of confirmation for TWO ROLLERS, not that it's crucial to *whether* your evidence confirms TWO ROLLERS.

One perhaps surprising prediction of this model is that learning FINE only modestly confirms the multiple rollers hypothesis. But the modesty of this confirmation turns out to be an artifact of the model, namely its using a multiple rollers hypothesis with a small number of rollers. To illustrate, recall that the portion of the outcome space allotted to your awakening on FINE and TWO ROLLERS is .8%. In contrast, holding other parameters fixed, the portion of the outcome space allotted to awakening on FINE and the multiple rollers hypothesis approaches 10.3% as the number of rollers on the multiple rollers hypothesis approaches infinity. But the portion of the outcome space allotted to your awakening on FINE and ONE ROLLER would remain .7%. So increasing the number of rollers on the multiple rollers hypothesis would increase the degree to which learning FINE confirms that hypothesis, conditional on your being awake. Thus, the proposed model captures the intuition that fine-tuning lends more support to the existence of a plenitudinous multiverse than it does to a sparse one.³⁴

Suppose I am right that Uncertainty* reflects the essential features of our epistemic situation with respect to uncertainty about the ideal support targeted by the

³⁴ A second aspect of the model that artificially limits how much learning FINE supports the multiple rollers hypothesis is that the probability of any given roll meeting the fine-tuning condition—that is, of being a double six—is only modestly low. A more realistic model would lower this probability to reflect the astronomical improbability of a universe being fine-tuned. Intuitively, if fine-tuning confirms M, that confirmation should increase with the improbability of fine-tuning. The model also captures this intuition, as FINE’s confirmation of the multiple rollers hypothesis increases with the improbability of individual rolls meeting the fine-tuning condition (given that the probability remains positive). However, these intuitions leave open exactly how confirmation of M scales with the number of universes and the improbability of fine-tuning. In contrast, when the number of rollers on that hypothesis is m and the probability of any given roll meeting the fine-tuning condition is $1/n$, the factor by which FINE confirms the multiple rollers hypothesis is: $((1/2 \times 1/2 \times 4/10 \times (1/n)) + (1/2 \times 1/2 \times 4/10 \times (1 - ((n-1)/n)^m)) + (1/2 \times 1/2 \times 2/10 \times (1/n)^m)) / (1/2 \times 1/2 \times 1/n)$. One way to arrive at a still more realistic model would be to determine how (dis)confirmation of the multiverse hypothesis should scale and adjust the model’s C and D conditions accordingly.

this-universe objection. Then discovering that we inhabit a fine-tuned universe should likewise boost our confidence in the existence of other universes. Further, that boost should both occur against an informational background that includes the fact that we are observers and depend crucially on the fine-tuning component of our fine-tuning evidence.

References

- Audi, R. (2017). Cumulative Case Arguments in Religious Epistemology. *Oxford Studies in Philosophy of Religion*. 8:1-15.
- Barnett, Z. (2019). Philosophy without belief. *Mind*, 128(509)109-138.
- Barrett, M., and Sober, E. (2020). The Requirement of Total Evidence: A Reply to Epstein's Critique. *Philosophy of Science*, 87(1),191-203.
- Bostrom, N. (2002). *Anthropic Bias: Observation Selection Effects in Science and Philosophy*. Routledge.
- Bradley, D. (2009). Multiple universes and observation selection effects. *American Philosophical Quarterly*, 46, 61–72.
- (2012). Four problems about self-locating belief. *Philosophical Review*, 121(2),149-177.
- Christensen, D. (2007). Epistemology of disagreement: The good news. *Philosophical Review* 116(2):187-217.
- (2008). Does Murphy's law apply in epistemology? *Oxford Studies in Epistemology* 2:3-31.
- (2010). Higher-order evidence. *Philosophy and Phenomenological Research*, 81(1), 185-215.
- Collins, R. (2009) "The teleological argument: an exploration of the fine-tuning of the cosmos", in W.L. Craig and J.P. Moreland (eds.), *The Blackwell Companion to Natural Theology*, Oxford: Blackwell, pp. 202–281.
- Dawkins, R. (1986) *The Blind Watchmaker*. London: Penguin Books.
- Draper, P. (2020). In Defense of the Requirement of Total Evidence. *Philosophy of Science*, 87(1),179-190.
- (2010). Cumulative cases. In C. Taliaferro, P. Draper and P.L. Quinn (eds.), *A Companion to Philosophy of Religion*. Malden, MA, USA: pp. 414-424.
- Dorst, C. and Dorst, K. (forthcoming). Splitting the (In)Difference: Why Fine-Tuning Supports Design. *Thought: A Journal of Philosophy*.
- Draper, K., Draper, P., and Pust, J. (2007). Probabilistic arguments for multiple universes. *Pacific Philosophical Quarterly*, (88)288–307.
- Dowe, P. (1999). Response to Holder: Multiple Universes Are Not Explanations. *Science & Christian Belief*, (11)67-8.
- Elga, A. (2007). Reflection and disagreement. *Noûs*, 41(3),478-502.
- Epstein, P. F. (2017). The Fine-Tuning Argument and the Requirement of Total Evidence. *Philosophy of Science*, 84(4):639-658.
- Field, H. (2000). A priority as an evaluative notion. In Paul A. Boghossian and C. Peacocke (eds.), *New Essays on the A Priori*. OUP.
- Friederich, S. (2019a). Reconsidering the Inverse Gambler's Fallacy Charge Against the Fine-Tuning Argument for the Multiverse. *Journal for General Philosophy of Science*, 50(1), 29-41.
- (2019b). A new fine-tuning argument for the multiverse. *Foundations of Physics*, 49(9),1011-1021.
- (2021) "Fine-Tuning", *The Stanford Encyclopedia of Philosophy* (Winter 2021 Edition), Edward N. Zalta (ed.), URL = <<https://plato.stanford.edu/archives/win2021/entries/fine-tuning/>>.
- Goff, P. (2021) "Is fine-tuning evidence for a multiverse?" manuscript.
- Hacking, I. (1987). The inverse gambler's fallacy: The argument from design. The anthropic principle applied to Wheeler universes. *Mind* 96(383):331-340.
- Huemer, M. (2018). *Paradox lost: Logical solutions to ten puzzles of philosophy*. Springer.

- Isaacs, Y.; Hawthorne, J. and Russell, J.S. (forthcoming). Multiple Universes and Self-Locating Evidence. *Philosophical Review*.
- Juhl, C. (2005). Fine-tuning, many worlds, and the 'inverse gambler's fallacy'. *Noûs* 39 (2):337–347.
- (2007). Fine-tuning and old evidence. *Noûs* 41(3):550–558.
- Landsman, K., (2016), “The fine-tuning argument: exploring the improbability of our own existence”, in K. Landsman and E. van Wolde (eds.), *The Challenge of Chance*, Springer, pp. 111–129.
- Leslie, J. (1989). *Universes*. Routledge.
- Lewis, D. (1986). *On the Plurality of Worlds*. Wiley-Blackwell.
- Manson, N.A. (2020). The Multiverse: What Philosophers and Theologians Get Wrong. *Theology and Science*, 18(1)31-45.
- Manson, N.A. (2022). Cosmic fine-tuning, the multiverse hypothesis, and the inverse gambler's fallacy. *Philosophy Compass* 17 (9).
- Manson, N.A., and Thrush, M. (2003). Fine-tuning, multiple universes, and the ‘this universe’ objection. *Pacific Philosophical Quarterly*, 84, 67–83.
- McGrath, P.J. (1988). The inverse gambler's fallacy and cosmology--a reply to Hacking. *Mind*, 97(386), 265-268.
- Metcalf, T. (2021). On Friederich’s New Fine-Tuning Argument. *Foundations of Physics*, 51(2)1-15.
- Parfit, D. (2011). *On What Matters: Vol. 2*. OUP.
- Plantinga, A. (2011). *Where the Conflict Really Lies: Science, Religion, & Naturalism*. OUP.
- Olding, A. (1990). *Modern Biology & Natural Theology*. Routledge.
- Saad, B. (forthcoming) Harmony in a Panpsychist world. *Synthese*.
- Schoenfield, M. (2012). Chilling out on epistemic rationality. *Philosophical Studies*, 158(2), 197-219.
- Smithies, D. (2015). Ideal rationality and logical omniscience. *Synthese*, 192(9)2769-2793.
- Sober, E. (2004). The design argument. In W. E. Mann (Ed.), *The Blackwell guide to the philosophy of religion* (pp. 117–147). Blackwell.
- Swinburne, R. (2004). *The Existence of God*. OUP.
- White, R. (2000). Fine-tuning and multiple universes. *Noûs* 34(2):260–276.
- White, R. (2003) “Fine-Tuning and Multiple Universes”. Chap. 13 in *God and Design: The Teleological Argument and Modern Science*, ed. by Neil A. Manson, 229–250. London; New York: Routledge.