

ARTICLE

Higher-Order Evidence and Normative Contextualism

Darren Bradley

PRHS, University of Leeds, Leeds, UK
Email: d.j.bradley@leeds.ac.uk

Abstract

How should you respond to higher-order evidence which says that you have made a mistake in the reasoning from your first-order evidence? It is highly plausible that you should reduce your confidence in your first-order reasoning. However, attempts to precisely formulate how this works have run into problems. I will argue that we should appeal to an independently motivated normative contextualism. That is, normative words like ‘ought’ and ‘reason’ have a different reference in different contexts. The result is that different answers to our question are true in different contexts.

Keywords: higher-order evidence; normative contextualism; calibrationism; disagreement; subjective and objective reasons

1. Introduction

How should you respond when you are told that you are unreliable in forming beliefs about some topic? Or when a peer disagrees with you? It has proved surprisingly difficult to defend any answer to these questions, as all answers seem to run into difficulties. I will argue that to make sense of these difficulties, we should appeal to an independently motivated normative contextualism. That is, normative words like “ought” and “reason” have a different reference in different contexts. The result is that all the suggested answers can be defended, but their appropriateness depends on the context.

There are three main options regarding what you should believe when faced with peer disagreement. First, the following principle supports not changing your credences:

Right Reason:

If p is the proposition best supported by your first order evidence, then the rational response is to believe p ¹

On the other hand, it is very plausible that you should shift your credence in the light of peer disagreement. Call this Calibrationism.² I will focus on a simple version of Calibrationism which says that your credence should match your expected degree of reliability.³ Following this line of

¹Titelbaum (2015) gives an extensive defense. For related views, according to which what ought to be done/believed depends on the facts, see Williamson (2002), Lasonen-Aarnio (2014), and Weatherson (2019). See Steel (2019) for discussion.

²Calibrationism comes from Schoenfield (2014) and is closely connected to *conciliationism*. For defenses, see Christensen (2007, 2010); Elga (2007); and Sliwa and Horowitz (2015).

³Isaacs (2019) points out that this ignores the base rate. Still, he agrees (p. 255) that there are cases where the base rate can be ignored (e.g., for propositions with prior probability .5 and uniform expected reliability). The peer disagreement case on which we focus is plausibly such a case. At any rate, the core issue is whether one should shift credence in the light of higher-order evidence (calibrationism) or ignore higher-order evidence (right reason). See Pittard (2019) for a related discussion.

thought, if you first judge that *p*, then learn that such judgments are correct 50% of the time, your credence that *p* should end up at 50%.

Consider two versions of Calibrationism:

Judgment-calibrationism:

If *p* is your judgment, and *r* is your expected degree of reliability, then the rational credence for you to assign to *p* is *r*.

Evidence-calibrationism:

If *p* is the proposition best supported by your first order evidence, and *r* is your expected degree of reliability, then the rational credence for you to assign to *p* is *r*.⁴

The difference between them is that *Judgment-calibrationism* starts with your *judgment* about what the evidence supports, while *Evidence-calibrationism* starts with what the evidence *really* supports. These agree when one judges correctly, but disagree when one makes a mistake. To see the difference, compare the following rules:

Make your best attempt at multiplying 5 and 6
Correctly multiply 5 and 6

Agents who make no mistakes do exactly the same thing in following these rules, but an agent who makes a mistake might follow the first without following the second. The second builds in a correctness condition, making it harder to successfully follow. Similarly, *Evidence-calibrationism* builds in a correctness condition, making it harder to successfully follow than *Judgment-calibrationism*. Schoenfield (2014), building on Kelly (2010), argues that neither of the calibrationism principles is correct. I will argue that versions of all three principles are correct, but are relevant in different contexts.

Section 2 explains the background and motivates the three principles, Section 3 explains contextualism, Section 4 offers a defense of *Right Reason*, Section 5 offers a defense of *Judgment-calibrationism*, Section 6 offers a defense of *Evidence-calibrationism*, Section 7 discusses the problem of disagreement, and Section 8 concludes.

2. Disagreement and Three Responses

In this section, I will explain the problem and discuss some arguments for and against the three responses mentioned above. Let us start with the following well-known example from Christensen (2007):

Restaurant

Suppose that five of us go out to dinner. It's time to pay the check, so the question we're interested in is how much we each owe. We can all see the bill total clearly, we all agree to give a 20 percent tip, and we further agree to split the whole cost evenly, not worrying over who asked for imported water, or skipped desert, or drank more of the wine. I do the math in my head and become highly confident that our shares are \$43 each. (As it happens, I'm correct.)

⁴My formulations follow Schoenfield (2014). Similar distinctions are made by Sliwa and Horowitz (2015) and Christensen (2016) (in his simple vs. idealized thermometer model). I leave it implicit that your expected degree of reliability is determined independently of your own first-order reasoning about the proposition under consideration. This assumption is called *Independence* (see Christensen, 2019).

Meanwhile, my friend does the math in her head and becomes highly confident that our shares are \$45 each. How should I react, upon learning of her belief?
I think that if we set the case up right, the answer is obvious.⁵

Setting up the case right involves adding that my friend and I know that we are epistemic peers, that is, we have the same evidence and are equally reliable. The obvious answer Christensen arrives at is:

I should lower my confidence that my share is \$43 and raise my confidence that it's \$45. In fact, I think (though this is perhaps less obvious) that I should now accord these two hypotheses roughly equal credence.⁶

Judgment-calibrationism delivers Christensen's answer. Recall:

Judgment-calibrationism:

If p is your judgment, and r is your expected degree of reliability, then the rational credence for you to assign to p is r .

p is "my share is \$43." We can assume that the expected degree of reliability is 50% given that both agents are epistemic peers and setting aside the possibility that you are both wrong. The result is that the rational credence to assign to "my share is \$43" is 50%.

Let us pause to note that the term "judgment" in *Judgment-calibrationism* is used in a technical way, as the proposition the agent regards as most likely to be correct on the basis of *first-order evidence* alone.⁷ This use of "judgment" presupposes that we can neatly divide up first-order and higher-order evidence. This is contestable,⁸ but for the purposes of our discussion, it is a harmless simplification. We can focus on cases like *Restaurant*, which can be divided into two stages—a first stage where the agent works through a specific set of evidence, for example, the contents of a bill, and a second stage in which she learns the judgment of a peer.

I take it that *Judgment-calibrationism* provides an intuitive answer. In light of the disagreement, it would be unacceptably dogmatic for either agent to maintain a strong belief in their original answer.

However, *Judgment-calibrationism* has two odd implications which Kelly (2010) drew attention to and which continue to be a source of controversy.⁹ The first oddity is that *Judgment-calibrationism* says that the agent who initially judged correctly and the agent who initially judged incorrectly should make equally extensive revisions to their beliefs (Kelly, 2010, p. 123). However, surely, says Kelly, there is an asymmetry. Should not the agent who initially judged correctly be required to make less extensive revisions than the agent who initially judged incorrectly?

The second oddity is that *Judgment-calibrationism* says that the agent who initially judged incorrectly *becomes* rational by following *Judgment-calibrationism* and assigning r to p . To see this, recall that the consequent of *Judgment-calibrationism* is "the rational credence for you to assign to p is r ." *Judgment-calibrationism* can be understood as a sufficiency condition for rationality, that is, the agent that assigns r to p becomes rational. However, surely it should not be that easy to become rational. After all, the agent who made the incorrect judgment was irrational. How could this irrationality drop out of the picture? Indeed, how could the *evidence*, which supports the belief that each owes \$43, drop out of the picture?

⁵Christensen (2007, p. 193). I take this example to be a placeholder for various cases in which an agent receives higher-order evidence.

⁶Ibid.

⁷A referee points out that, technically, the disjunction of all possible answers would be most likely to be correct. We can assume that the relevant propositions form an exclusive and exhaustive set of answers.

⁸See Hedden and Dorst (2022).

⁹The two reasons are not explicitly distinguished by Kelly.

These challenges might be answered by replacing *Judgment-calibrationism* with:

Evidence-calibrationism:

If p is the proposition best supported by your first order evidence, and r is your expected degree of reliability, then the rational credence for you to assign to p is r .

Evidence-calibrationism cannot be followed by the person who initially judges incorrectly, as they will be wrong about which proposition is best supported by their first-order evidence. So, *Evidence-calibrationism* will not deliver the odd verdicts that both agents should make equally extensive revision, nor that an agent that judges incorrectly can become rational by following *Evidence-calibrationism*.

However, the problem now is that following *Evidence-calibrationism* seems to require having access to which proposition is best supported by your first-order evidence, and if the agent has access to this, then they should believe that proposition (the one best supported by their first-order evidence) and ignore their expected degree of reliability.

Given these problems for Calibrationism, we might resort to:

Right Reason:

If p is the proposition best supported by your first order evidence, then the rational response is to believe p .

This says that you should believe your share is \$43 if \$43 is the result of correctly doing the arithmetic, and you should believe your share is \$45 if \$45 is the result of correctly doing the arithmetic.

The problem here is that *Right Reason* seems to imply that you should ignore relevant evidence. It seems unacceptably dogmatic to receive evidence from a peer that you have made a mistake, and respond by refusing to shift your opinion one iota.

The right reason theorist will say that the person who got the correct answer is right not to shift opinion. They might emphasize that it is only the person who has done the arithmetic correctly who should stick with their opinion—the person who made a mistake should reconsider their opinion. However, this is puzzling, as neither agent knows whether they made a mistake. Their epistemic positions are symmetrical in this key sense. In my view, we can see that there is something right about all three responses once we embrace contextualism.

3. Contextualism

In this section, I will explain the core ideas of normative contextualism. It is a familiar thought that whether someone is correctly described as tall depends on the details of the conversation. For example, Michael Jordan, at 1.98 m, is tall for an ordinary person, but not tall for a basketball player. So the truth of “Michael Jordan is tall” depends on the conversational context. It is true in the context of discussing ordinary people, but false in the context of discussing basketball players.

A popular theory in linguistics is that normative terms like “ought” and “reason” are context-sensitive in a similar way. The traditional Kratzerian semantics posits two parameters which are determined by the conversational context (Kratzer, 1981).¹⁰

The *first* parameter is a modal base, which corresponds to a proposition, a set of information, or a set of live (partial¹¹) possible worlds. “It must be that P ” means, roughly, that in all the live worlds, P . The *second* parameter determines a ranking of the live worlds. “It ought to be that P ” means roughly that in all the *best* live worlds, P . So “ S ought to believe P ” means, roughly, that S believes P

¹⁰See Finlay (2014) for a helpful study.

¹¹The possible worlds need to be partial because a full possible world determines the truth, the belief, the first-order evidence, and so forth. We need partial possible worlds (propositions) that determine only one of these.

in all the best live worlds. The modal base can be made explicit using “given,” for example, S ought to believe P given proposition B.

Candidate values for the second parameter include a moral ranking, which ranks worlds in order of how morally good they are, and an epistemic ranking, which ranks worlds in order of how epistemically good the beliefs/credences are. In this paper, we will focus exclusively on the epistemic ranking. The first parameter, the modal base, is the only moving part we need.

What are the candidates to be the modal base? The mechanics allow any proposition, but we will focus on four: all and only the true propositions; all and only the propositions believed; all and only the propositions supported by the first-order evidence; all and only the propositions supported by the total evidence. These correspond to the following “oughts”:¹²

Truth-relative ought

What you ought to believe given the truth.

Belief-relative ought

What you ought to believe given your beliefs (e.g., what you should infer).

First-order-evidence-relative ought

What you ought to believe given your first-order evidence.

Total-evidence-relative ought

What you ought to believe given your total evidence.

To make these explicit, instead of using “ought” (e.g., S ought to believe P), we can use one of these four (e.g., “S truth-relative-ought to believe P”). I will talk about these being used in conversation, but of course, they are all pronounced the same in English.

The familiar distinction between the objective and subjective “ought” can be plugged into this framework.¹³ Suppose there will be dancing at a party, and Dancing Dave loves dancing while Talking Tim hates dancing. Dancing Dave correctly believes that there will be dancing at the party, while Talking Tim falsely believes that there will be no dancing, only talking. There are senses in which both ought to go to the party, but the “oughts” seem to be of different types. We can say that Dave objectively ought to go to the party because there will be dancing, but Tim only subjectively ought to go to the party because he only believes there will be dancing. As Schroeder (2004, p. 348) writes, “On a natural view, subjective reasons are simply things that you believe such that, if they are true, they are reasons for you to do something.” The objective ought will correspond to our truth-relative ought; the subjective ought identified by Schroeder will correspond to our belief-relative ought.

What determines the modal base? I suggest that part of what determines the modal base is the aim of the conversation. Several writers have distinguished two aims (among others) that people might have when using “ought” or other normative terms:¹⁴

- i) Expressing standards.
- ii) Deliberating.

¹²For related distinctions, see Feldman (1988, pp. 407–408), Worsnip (2021, p. 31).

¹³The example is based on Schroeder (2007, p. 1). See Jackson (1991) for the *locus classicus* and <https://philpapers.org/browse/subjective-and-objective-reasons> for recent papers.

¹⁴See Steinberger (2019: 7) for a helpful discussion. Relatedly, Bales (1971) distinguishes decision procedures from right-making characteristics; Arpaly (2003: 34) distinguishes a rational agent’s manual from an account of rationality; Schroeder (2011, pp. 1–2) distinguishes deliberative from evaluative oughts; McHugh (2012: 9–10) distinguishes prescriptive norms from evaluative norms; and Schoenfeld (2018: 690) distinguishes plans to make from procedures to conform to.

These most clearly come apart when the agent is missing important information. To take a well-known example, suppose there is petrol in Bernard's glass but he believes, with good reason, that it contains gin.¹⁵ What ought he to do? It is very plausible that there is no univocal answer—there is a sense in which he ought to drink and a sense in which he ought not to. These fit with the aims of expressing standards and deliberating which we can make explicit with “standards-ought” and “deliberative-ought,” that is, he standards-ought not to drink and deliberatively-ought to drink.¹⁶ We get these results if the modal base in contexts where standards are expressed includes information the agent does not have (e.g., that the glass contains petrol), while the modal base in deliberative contexts includes only information the agent has.

We can make more fine-grained distinctions to recover the four senses of “ought” above. Starting with (i), we can distinguish two standards. One standard is to believe all and only truths.¹⁷ One might say “one ought to believe all and only truths” and we can make sense of this “ought” as the truth-relative ought.

However, as we often do not have full information, in most contexts it is not appropriate to simply say that you ought to believe all and only truths. This leads to a second standard, which is to form the appropriate beliefs given the first-order evidence. We can think of this as the standard of not making any mistakes in reasoning and ignoring any higher-order evidence. (Higher-order evidence can only be misleading given that one has formed the appropriate beliefs given the first-order evidence.) Such contexts invoke the first-order-evidence-relative ought.

Moving to (ii), we can distinguish two types of deliberative context, which I will call a prospective-deliberative context and a retrospective-deliberative context.

In some deliberative contexts, the question is how the agent should proceed, holding fixed their current beliefs. The agent might have made mistakes in the past, but the question is how to proceed given what they currently believe. Call this a prospective-deliberative context. Such contexts invoke the belief-relative ought.

In other deliberative contexts, current beliefs are not held fixed. Agents have the time and cognitive capacity to reconsider their beliefs in the light of their total evidence. Call this a retrospective-deliberative context. Such contexts invoke the total-evidence-relative ought.

I suggest that these oughts allow us to identify what is true in *Right Reason*, *Judgment-calibrationism* and *Evidence-calibrationism* as follows:

- a) *Right reason* is true for first-order-evidence-relative ought
- b) *Judgment-calibrationism* is true for belief-relative ought
- c) *Evidence-calibrationism* is true for total-evidence-relative ought

I will argue in the next three sections, respectively, that these principles explain the conflicting intuitions.

4. Right Reason is True for First-Order-Relative-Ought

Recall:

Right Reason:

If *p* is the proposition best supported by your first order evidence, then the rational response is to believe *p*

¹⁵See Williams (1981). Williams' focus is on reasons and mine is on 'ought', but, as I say above, the contextualist framework can be naturally extended to other normative terms such as 'reasons'.

¹⁶The conflict could be avoided by denying that a belief can be rational and false (e.g., Sutton, 2007).

¹⁷See McHugh (2012).

This says that the person who did the arithmetic correctly should ignore the evidence that their peer came to a different answer. Recall that the main objection was that this seems to ignore relevant evidence.

However, if the context invokes the standard of forming the appropriate beliefs given the first-order evidence, then we get (a):

Right Reason First-Order Evidence:

If *p* is the proposition best supported by your first order evidence, then you first-order-evidence-relative-ought to believe *p*

There can be no debate over the truth of *Right Reason First-Order Evidence*. It is true by definition. I will argue that the arguments and intuitions of those who defend *Right Reason* are best understood as invoking the first-order-evidence-relative ought and *Right Reasons First-Order Evidence*. We will look at Field (2000), Titelbaum (2015), and Weatherson (2019).¹⁸

Hartry Field (2000) gave an early defense of what was later called *Right Reason*. His specific claim was that a priori justification is indefeasible by empirical evidence. In response to the argument that one could get evidence against propositions that have a priori justification, Field wrote:

[W]hile the non-ideal credibility of, say, a complex logical truth can certainly be lowered by empirical evidence that well-respected logicians didn't accept it, *ideal* credibility can't be lowered in this way. (p. 118 *Italics added*)

It is not entirely clear what "ideal credibility" means, but one possibility is that it expresses a standard. In a context in which we are discussing standards, we might say that the agent ought to make no mistakes in their reasoning and ignore higher-order evidence, so first-order-evidence-relative-ought is relevant. *Right Reason* will therefore be true because it will be understood as *Right Reason First-Order Evidence*. So Field's defense of *Right Reason* can be understood as a defense of *Right Reason First-Order Evidence*.

Relatedly, Titelbaum's (2015) defense of *Right Reason* relies on:

Assets

All agents, in all possible situations, possess a priori propositional justification for the rational requirements that is indefeasible

In response, Claire Field (2019, p. 179) writes, "*Assets* requires some explanation. Titelbaum provides no such explanation when he introduces the claim, and as it stands, it is somewhat surprising. More recent advocates of a priori justification have typically thought of it as defeasible at best (see BonJour, 1998)." I agree with Field that *Assets* is not very plausible. The point I want to make here is that, to the extent that I can get myself to see how someone might accept *Assets*, it seems most plausible as expressing a standard. Thus, *Right Reason* is most plausible if understood as *Right Reason First-Order Evidence*.

Finally, Weatherson (2019) offers a book length defense of a *Right Reason* style theory. He distinguishes making normative claims about actions from giving normative advice. Put in our terms, 'making normative claims' seems to be relevant to the aim of expressing standards and 'giving normative advice' seems to be relevant to the aim of deliberating (or helping someone else deliberate). In considering the view he opposes, according to which normative claims depend on the agent's higher-order evidence, he writes:

¹⁸See also Tal (2020), who explicitly appeals to ideal agents, in the sense that ideal agents always draw the correct inferences from their first-order evidence.

[My theory, Right Reason] is most strongly opposed to the view about actions, and least strongly opposed to the view about advice. P.2

It is not entirely clear what he means in saying that his theory is “least strongly opposed” to the view about advice. However, one plausible reading is that Weatherson is least *confident* that *Right Reason* is correct when it comes to advice. In addition, this would fit with my view that *Right Reason* is true in some contexts where agents express standards, but not true in contexts where agents deliberate.

So far I have argued that *Right Reason* can be motivated by thinking about the first-order-evidence-relative-ought. I add that it can also be motivated by thinking about the truth-relative-ought. For example, in *Restaurant*, the first-order evidence has guided the agent to the truth. So, in a context in which the relevant “ought” is what one truth-relative-ought to believe, *Right Reason* delivers the correct verdict. The agent truth-relative-ought to believe they owe \$43; the agent also first-order-evidence-relative-ought to believe they owe \$43. These “oughts” do not diverge in this example. This bolsters my argument that defenders of *Right Reason* can be understood to have in mind contexts where standards are being expressed—either the standard of forming the appropriate beliefs given the first-order evidence, or the standard of believing all and only truths.¹⁹

5. Judgment-Calibrationism is True for Belief-Relative-Ought

In this section, I will argue that *Judgment-calibrationism* is true for contexts which invoke prospective deliberation, and that this explains calibrationist intuitions. Recall:

Judgment-calibrationism:

If *p* is your judgment, and *r* is your expected degree of reliability, then the rational credence for you to assign to *p* is *r*.

(b) says this is true for the belief-relative-ought:

Judgment-calibrationism Belief-relative:

If *p* is your judgment, and *r* is your expected degree of reliability, then you belief-relative-ought to assign credence *r* to *p*.

Which set of beliefs are relevant? Contextualism offers a natural answer—the beliefs referred to in the antecedent.²⁰ That is, we start with the belief/judgment that *p* and the expected reliability of *r*, then ask what the agent ought to believe given that starting point. In the restaurant case, the result is that the agent belief-relative-ought to assign 50% credence to their share being \$43. *Judgment-calibrationism Belief-relative* sets aside whether the initial beliefs referred to in the antecedent are rational and simply tells us what to do given that we have them.

This fits with contexts that invoke prospective deliberation (which holds fixed current beliefs). We can create a context that invokes prospective deliberation by asking what you, the reader, would do in *Restaurant*. Imagine being in the situation described in *Restaurant*. What procedure would you follow? Notice you have not been told how you arrived at \$43, so it is impossible to reconsider the way in which you arrived at that answer. I suggest that the most plausible procedure to follow is *Judgment-calibrationism Belief-relative*. Thus, we can explain intuitions in favor of *Judgment-calibrationism Belief-relative* as invoking contexts of prospective deliberation.²¹

¹⁹Gonzalez de Prado (2020) makes a suggestion that fits my analysis. He thinks *Right Reason* fails to take into account that what matters is not what reasons there *are*, but what reasons are *possessed*. I understand this shift as moving from truth-relative-ought or first-order-evidence-ought to belief-relative-ought or total-evidence-relative-ought.

²⁰This is the “restrictor” analysis of indicative conditionals. See Charlow (2015) for discussion.

²¹Schoenfield (2018) suggests that what is right about Calibrationism is that it is the best principle to *plan* to follow. This is compatible with my position, but I think the appeal to plans obscures something important. Schoenfield emphasizes the

Of course, the proponent of *Right Reason* will reply that the agent does have access to which proposition is best supported by the evidence. They will claim that the first-order evidence is undefeated. My point is that this move is less plausible when we invoke a deliberative-ought than it is when we invoke a standards-ought. We can accommodate the intuitions of the *Right Reason* theorist without being committed to the view that *Right Reason* should be used in deliberation.²²

Relatedly, the proponent of *Right Reason* might complain that the evidence that supports \$43 should not have dropped out of the picture. Even in a deliberative context, why shouldn't the agent use the total evidence, and why shouldn't that total evidence support for the answer of \$43? In response, the Calibrationist need not concede that the evidence that supports \$43 has dropped out of the picture. Quite the contrary—the evidence for \$43 provided the support for the agent's initial judgment that they each owed \$43. The evidence is still there, supporting this belief.

To be clear, it is not my aim to show that the Calibrationist is right about this. My aim is to show that the intuitions motivating Calibrationists (and *Right Reason* theorists) can be explained by a contextualist framework.

How does *Judgment-calibrationism Belief-relative* fare against Kelly's objections? Kelly's first objection was that *Judgment-calibrationism* misses the fact that an agent who irrationally judges that *p* should be required to make more extensive revisions to their opinions than an agent who rationally judges that *p*.

In response, *Judgment-calibrationism Belief-relative* does say that both agents are required to make equally extensive revisions to their opinions, but that is because we are starting with their judgments and asking how to proceed from there. Any mistake an irrational agent made in working out the bill has fed into their judgment. If we hold fixed those judgments and ask how to proceed from them, each *should* make equally extensive revisions.

Kelly's second objection was that an agent who irrationally judges does not become rational merely by calibrating. Yet Calibrationism seems to imply that they do.

In response, *Judgment-calibrationism Belief-relative* does not make any claim about the overall rationality of the agent, so it does not provide a way of becoming rational. It says what the agent belief-relative-ought to believe given their prior beliefs, but does not offer a judgment on the rationality of those prior beliefs, so it does not offer a judgment on their overall rationality. The belief-relative-ought only offers a judgment on belief-relative rationality.²³

Let us consider an objection to *Judgment-calibrationism* from David Christensen (2016).²⁴ If solid, it would be just as serious an objection to *Judgment-calibrationism Belief-Relative*. Christensen argues that *Judgment-calibrationism* is incomplete. The problem is that *Judgment-calibrationism* refers to an agent's judgment, but there are cases where the agent has not formed a judgment. Christensen writes, "there are cases where agents have not formed any initial credence on the basis of a batch of first-order evidence before they get the higher-order evidence" (p. 406). And he points out that there are cases where the agent is told of a peer's opinion before thinking through the first-order evidence. In such cases, the agent will never form a credence on the basis of

diachronic nature of plans, but even if we focus on a single time, if the agent has learnt of the disagreement, then it is Calibrationism they should be guided by rather than *Right Reason*.

²²Compare contextualism about knowledge (Lewis, 1996). A skeptic can always claim that even in ordinary contexts, you do not know you have hands. The contextualist response is that skeptical intuitions are less compelling in an ordinary language context than in a philosophical/skeptical context. However, there are not going to be any knockdown arguments here, and you cannot account for everyone's intuitions. Similarly, my aim is to show that *Right Reason* intuitions are more plausible in standards contexts than deliberative contexts.

²³This fits with Schoenfield's (2015 section 5) comments that *Judgment-calibrationism* might be thought of as a principle for belief transitions.

²⁴Christensen discusses three problems with *Judgment-calibrationism*. He argues that two of the problems can be solved. I agree with Christensen's solution to the first two problems, so I will not discuss them here.

just the first-order evidence. Christensen uses this to motivate the transition from *Judgment-calibrationism* to *Evidence-calibrationism*.

However, I don't think this is a fatal objection, for an agent can work through what the first-order evidence *seems to support*, even if they believe they are unreliable. They can then take into account their believed unreliability to arrive at a considered view. Compare a case where you are first told that your color perception has been inverted. Then you appear to see a blue table. You can judge that your first-order evidence seems to support that the table is blue, then take into account your believed unreliability and move to the considered verdict that the table is orange. One can go through a similar process in cases like *Restaurant*. Recall, the term "judgment" in Calibrationism is used in a technical way, as the proposition the agent regards as most likely to be correct on the basis of *first-order evidence* alone. You might first get your peer's opinion, then work out what the first-order evidence *seems to show*, then come to a considered judgment based on these two pieces of evidence in the manner suggested by *Judgment-calibrationism*.

6. Evidence-Calibrationism is True for Total-Evidence-Relative-Ought

Let us now turn to:

Evidence-calibrationism:

If p is the proposition best supported by your first order evidence, and r is your expected degree of reliability, then the rational credence for you to assign to p is r .

Evidence-calibrationism seems to have been motivated largely by the problems with *Judgment-calibrationism* (Christensen, 2016). In light of my defense of a version of *Judgment-calibrationism*, this motivation is less compelling. Still, what should we say about it?

Contextualists can allow that there is a reading that makes *Evidence-calibrationism* true. It is most plausible with the *Total-evidence-relative-ought*:

Evidence-calibrationism Total-evidence:

If p is the proposition best supported by your first order evidence, and r is your expected degree of reliability, then you total-evidence-relative-ought to assign credence r to p .

Is *Evidence-calibrationism Total-evidence* a principle of deliberation or a standard? Perhaps both.

Starting with deliberation, I suggested at the end of Section 3 that when we create a context that invokes retrospective deliberation (which does not hold current beliefs fixed), *Evidence-calibrationism Total-evidence* is intuitive. That is, when we do not hold the agent's beliefs fixed and allow them to consider afresh what they should believe, *Evidence-calibrationism Total-evidence* seems like the correct rule for deliberation. Unlike *Judgment-calibrationism Total-evidence*, *Evidence-calibrationism Total-evidence* builds in the correctness condition that p is the proposition best supported by the first-order evidence, making it more demanding and harder to follow. Nevertheless, perhaps it could still be used in deliberation. If not, *Evidence-calibrationism Total-evidence* can be thought of as expressing a standard.

Evidence-calibrationism Total-evidence seems intuitive when the relevant standard is: believing what the total evidence supports, while allowing that you might have made a mistake. It tells the agent to come to the right judgment about what the first-order evidence supports, but still take into account that you might have made a mistake about what the first-order evidence supports.

Someone might object that *Evidence-calibrationism Total-evidence* seems imperfect both as a principle of deliberation and as a standard. I am happy to concede the point. Perhaps *Evidence-calibrationism Total-evidence* occupies an awkward halfway house between a principle for expressing standards and a principle for deliberation. Still, I see no reason to exclude this kind of "ought" from the language.

7. Contextualism and the Disappearance of Disagreement

A familiar objection is that contextualism wrongly predicts that there is no genuine disagreement between proponents of *Right Reason*, *Judgment-calibrationism* and *Evidence-calibrationism*.²⁵ Once the parameters are made explicit, each participant in the debate should accept what the other is saying and disagreement disappears—and that looks like the wrong verdict.²⁶

A first response is that it would not be so surprising to find no genuine disagreement. It is common for people to appear to disagree, only to find out that they have been using language slightly differently. Chalmers (2011) argues that this is a typical feature of philosophical debates. In addition, when it comes to normative terms like “ought,” we even have a widely accepted linguistic theory that tells us where the hidden parameters are and how they might lead to the appearance of disagreement.

Nevertheless, it would be unsatisfying to simply say that all participants in the debate are just confused by language and talking past each other. The challenge for the contextualist is to make sense of the intuition that there is genuine disagreement. Let me offer four compatible ways to retain disagreement.

First, even if statements about what one *ought* to believe are context-sensitive in English, there might be a metaphysically privileged parameter that picks out a normatively privileged property. Worsnip writes:

we should be careful to separate the question of whether (e.g.) the law ...has genuine normative authority from whether there is a robustly normative usage of the legal “ought.” The former requires the law to actually possess normative authority, whereas the latter only requires there to be speakers who take the law to possess normative authority. (Worsnip, 2019, p. 3104)

Worsnip is working with a primitive concept of “normative authority.” He is allowing that there might be lots of “oughts,” just as contextualism predicts, but that not all of them have normative authority. Perhaps only one “ought” has normative authority. Indeed, Broome (2013, p. 24) talks about the “central ought” and Kiesewetter (2017, p. 9) talks about the “deliberative ought,” and these might be taken to be attempts to refer to one privileged “ought” with normative authority.²⁷ Disagreement thus remains when speakers disagree about which “ought” has normative authority.

Someone might object that contextualism would be much less interesting if there were a single “ought” with normative authority. And indeed it would be. However, contextualism reduces the motivation for the thesis that there is a normatively authoritative “ought.” If there were a normatively authoritative “ought,” then we would face tricky metaphysical (what makes an “ought” privileged?) and epistemic questions (how do we know which is privileged?). Dispensing with a normatively authoritative “ought” allows us to avoid such questions and to explain apparent disagreements by appealing to different speakers using different parameters. What’s not to like?

Furthermore, all that’s needed for genuine disagreement is differing *beliefs* about which “ought” is privileged (even if none of them actually are privileged). Compare: If you think Santa Claus is married and I think he is a bachelor, we disagree, even though we are both wrong. So, genuine disagreement is compatible with the absence of a privileged “ought.”

The second way to retain disagreement is to hold that sometimes disagreement involving contextualist terms should be understood as disagreement about *which parameters to use*. To motivate this view, note that contextualism says that the parameters are determined by the conversational context, and which parameters are operative might be in flux in a conversation.

²⁵Gibbons’ (2013 ch. 3) objection to contextualism about the norms of belief is that there must be genuine disagreement.

²⁶To be clear, the disagreement at issue is between proponents of *Right Reason*, *Judgment-calibrationism*, and *Evidence-calibrationism* rather than the disagreement between the people adding up the bill.

²⁷A similar view has been discussed in detail in Eklund (2017).

For example, consider the following conversation when all parties know that Michael Jordan is 1.98 m tall and Sun Mingming is 2.45 m tall:

- A. Michael Jordan is tall
- B. No he isn't, Sun Mingming is tall

It is plausible that this conversation is best understood as an implicit negotiation about the standards for the word "tall."

Lewis (1979; 1996) suggests that conversations should conform to the principle of accommodation, which says that when a speaker makes an utterance involving a context-sensitive term, the parameters shift to make the sentence true. In Lewis's terms, the utterance changes the "conversational score." Thus, B is attempting to shift the conversational score to one in which the relevant sense of "tall" is "tall-for-a-basketball player." Applied to our example, defenders of *Right Reason* are arguing that the best norm of belief in some context is *Right Reason* and defenders of *Judgment-calibrationism* are arguing that the best norm of belief in some context is *Judgment-calibrationism*.

However, what if the hearer does not want to change the conversational score? Then, we get what Plunkett and Sundell (2013) call *meta-linguistic negotiation*, that is, an exchange in which speakers tacitly negotiate the proper deployment of a linguistic expression in a given context. There is no guarantee that there will be a successful conclusion to the metalinguistic negotiation, resulting in an unresolved disagreement.

This differs from the first way to retain disagreement because the speakers need not believe that any parameter is metaphysically privileged. They might still disagree about the aim of the conversation. That is, one person wants a conversation about the standards, and another wants a conversation about deliberation. This framework allows for genuine disagreement about which context to create.

This leads to the third way to retain disagreement—there might be disagreement about *what to do*. Björnsson and Finlay (2010) argue that the hidden goal of many conversations involving contextualist terms is to establish what to do. Thus, there need be no disagreement over which propositions are true, but still disagreement over what to do.²⁸ (If determining the parameters is something we do, then the second way to retain disagreement could be considered a version of this third way.)

The fourth way to retain disagreement is to separate the metaphysics from the epistemology.²⁹ One might hold that the full conversational context *determines* the values of the parameters (metaphysical), but the *inference* (epistemic) from the full conversational context to the values of the parameters is nonobvious. For example, the full conversational context might be determined by the intentions of the speakers, but speakers' intentions are not always transparent, so hearers might infer false conclusions about the context. This leaves room for genuine disagreement about which beliefs one "ought" to have.

Therefore, the contextualist can make sense of the intuition that there is genuine disagreement between proponents of *Right Reason*, *Judgment-calibrationism*, and *Evidence-calibrationism*. However, there may well be less disagreement than originally appeared, and the remaining disagreement might not be what it initially seemed to be.

8. Conclusion

It is plausible that many debates in philosophy come down to verbal disputes. Sometimes we have an independently motivated framework that predicts these verbal disputes, but the dispute

²⁸For criticism, see McKenna (2014) and Bolinger (2022).

²⁹To the best of my knowledge, this has not been suggested before.

continues regardless. I have argued that this is the situation in the debate about disagreement and higher-order evidence. The various positions are all compatible once we take into account the context-sensitivity of normative terms. Thus, *Right Reason*, *Judgment-calibrationism*, and *Evidence-calibrationism* are all true in different contexts.

Funding statement. The author is grateful to the Templeton Foundation for partial funding of this research (Grant No. 62603).

Competing interests. The author declares no competing interests exist.

Darren Bradley is a Professor of Metaphysics and Epistemology at the University of Leeds. He works primarily in epistemology, and especially in places where it overlaps with metaphysics and metaethics.

References

- Bales, R. E. (1971). Act-utilitarianism: Account of right-making characteristics or decision-making procedure? *American Philosophical Quarterly*, 8(3), 257–265.
- Broome, J. (2013). *Rationality through reasoning*. Oxford: Blackwell.
- Chalmers, D. J. (2011). Verbal disputes. *Philosophical Review*, 120(4), 515–566.
- Charlow, N. (2015). Triviality for restrictor conditionals. *Noûs*, 50(3), 533–564.
- Christensen, D. (2007). Does murphy's law apply in epistemology? Self-doubt and rational standards. *Oxford Studies in Epistemology*, 2, 3–31.
- Christensen, D. (2016). Disagreement, drugs, etc.: From accuracy to Akrasia. *Episteme*, 13, 397–422.
- Christensen, D. (2019). Formulating independence. In M. Skipper, & A. Steglich-Peterse (Eds.), *Higher-order evidence: New essays* (pp. 13–34). Oxford/New York: Oxford University Press. <https://doi.org/10.1093/oso/9780198829775.003.0001>.
- Elga, A. (2007). Reflection and disagreement. *Nous*, 41(3), 478–502.
- Feldman, R. (1988). Subjective and objective justification in ethics and epistemology. *The Monist*, 71(3), 405–419.
- Field, C. (2019). It's ok to make mistakes: Against the fixed point thesis. *Episteme*, 16(2), 175–185. <https://doi.org/10.1017/epi.2017.33>.
- Finlay, S. (2014). *Confusion of tongues: A theory of normative language*. New York: Oxford University Press.
- Gibbons, J. (2013). *The norm of belief*. New York: Oxford University Press.
- González de Prado, J. (2020). Dispossession defeat. *Philosophy and Phenomenological Research*, 101, 323–340. <https://doi.org/10.1111/phpr.12593>.
- Hedden, B., & Dorst, K. (2022). (Almost) all evidence is higher-order evidence. *Analysis*, 82(3), 417–425.
- Jackson, F. (1991). Decision-theoretic consequentialism and the nearest and dearest objection. *Ethics*, 100(3), 461–482.
- Kelly, T. (2010). Peer disagreement and higher-order evidence. In R. Feldman & T. Warfield (Eds.), *Disagreement*. Oxford: Oxford University Press.
- Kiesewetter, B. (2017). *The normativity of rationality*. Oxford: Oxford University Press.
- Kratzer, A. (1981). The notional category of modality. In *Words, worlds, and contexts*. Berlin, New York: de Gruyter, 38, 74.
- Lasonen-Aarnio, M. (2014). Higher-order evidence and the limits of defeat. *Philosophy and Phenomenological Research*, 88, 314–345.
- McHugh, C. (2012). The truth norm of belief. *Pacific Philosophical Quarterly*, 93, 8–30.
- Pittard, J. (2019). Fundamental disagreements and the limits of instrumentalism. *Synthese*, 196(12), 5009–5038.
- Schoenfield, M. (2014). A dilemma for Calibrationism. *Philosophy and Phenomenological Research*, 91(2), 425–455.
- Schoenfield, M. (2018). An accuracy based approach to higher order evidence. *Philosophy and Phenomenological Research*, 96(3), 690–715.
- Schroeder, M. (2004). The scope of instrumental reason. *Philosophical Perspectives* 18(1), 337–364.
- Sliwa, P., & Horowitz, S. (2015). Respecting all the evidence. *Philosophical Studies*, 172(11), 2835–2858.
- Steel, R. (2019). Against right reason. *Philosophy and Phenomenological Research*, 99, 431–460. <https://doi.org/10.1111/phpr.12491>.
- Steinberger, F. (2019). Three ways in which logic might be normative. *The Journal of Philosophy*, 116(1), 5–31.
- Titelbaum, M. (2015). Rationality's fixed point (or. In defense of right reason). *Oxford Studies in Epistemology*, 5, 253–294.
- Weatherson, B. (2019). *Normative externalism*. Oxford: Oxford University Press.
- Williamson, T. (2002). *Knowledge and its limits*. Oxford: Oxford University Press.
- Worsnip, A. (2019). 'Ought'-contextualism beyond the parochial. *Philosophical Studies*, 176(11), 3099–3119.
- Worsnip, A. (2021). *Fitting things together: Coherence and the demands of structural rationality*. Oxford: Oxford University Press.