

Timothy Williamson

RECOGNITIONAL CAPACITIES AND THEIR USES

Abstract: A recognitional capacity for a kind is a capacity to come to know whether an appropriately presented object is a member of that kind; similarly, there are recognitional capacities for properties and for particular individuals. Recognitional capacities are typically non-reflective and fallible. Although often associated with words of a natural language that refer to the kind, property, or individual, they are not part of the semantics of those words and are not required for linguistic competence with them. Recognitional capacities for properties expressed by moral terms are argued to play a major role in moral epistemology, and even in moral perception. Arguably, David Chalmers' recently proposed methodology for identifying verbal disputes and identifying bedrock concepts depends on neglecting the epistemic role of recognitional capacities associated with ordinary words.

Keywords: recognitional capacities, moral epistemology, moral perception, verbal disputes, David Chalmers.

In this paper, I will briefly explain what recognitional capacities are and how we humans and other animals use them, and derivatively how we philosophers can use the category of recognitional capacities to better understand everyday phenomena that they have managed to be puzzled by.

To illustrate: a recognitional capacity for horses is a capacity, when you see something, to recognize whether it is a horse. Of course, such capacities are never perfect. You may see a horse in the distance, or in bad light, and mistake it for a cow, or be uncertain whether it is a horse or a cow. Likewise, you may see a cow in the distance, or in bad light, and mistake it for a horse, or be uncertain whether it is a horse or a cow. Most people who understand the word 'horse', or a synonymous word in another language, have an associated recognitional capacity for horses, though that is not strictly needed for understanding the word. You could understand it by having heard others talk about 'horses', without ever having seen a horse, or even a picture of one.

Not all recognitional capacities are visual. A blind musician may have an auditory recognitional capacity for oboes. She has a capacity, when she hears something, to recognize whether it is an oboe.

One can also have a recognitional capacity for a particular individual. Many people have a capacity, when they see something, to recognize whether it is Donald Trump (or a picture of him), though such a capacity is not needed for understanding the name ‘Donald Trump’.

Recognitional capacities are typically *unreflective*. When you recognize a horse, or Donald Trump (perhaps in a picture), you do not normally engage in conscious calculation or inference. You just judge ‘It’s a horse’ or ‘It’s Donald Trump’. Of course, complex brain processes are at work behind the scenes, but they are not suited to appearing on the stage of consciousness.

I take the phrase ‘recognitional capacity’ from the work of Michael Dummett. However, I do not expect recognitional capacities to do the sort of work he introduced them to do. Dummett introduced recognitional capacities as candidates to be the more or less Fregean senses of terms that could not plausibly be treated as abbreviating verbal descriptions, at least for a given speaker (Dummett 1973). By contrast, I do not cast recognitional capacities in any strictly *semantic* role, although they may play a major *metasemantic* role in helping fix the referents of the associated terms, since an associated recognitional capacity can be central to how speakers use their terms, what they apply them to and what they do not. Arguably, dispensing recognitional capacities from any semantic duties enables us to give a more realistic account of their epistemology, free from ill-motivated and unrealistic constraints.

In the next section, I will explain how acknowledging recognitional capacities for moral properties enables us to give a more psychologically plausible account of moral epistemology. That account draws on a more general defence of moral realism that I have developed elsewhere. In the final section, I will explain how, by reflection on the dependence of our use of natural language on recognitional capacities, we can diagnose what is wrong with a widely held view of concepts, one which in particular underlies David Chalmers’ discussion of verbal disputes (Chalmers 2011).

Other philosophical uses of the category of recognitional capacities include cases where acquiring non-deductive evidence for a conclusion triggers a recognitional capacity by which one comes to know that very conclusion, thereby making it part of one’s evidence base; acknowledging such double-takes helps fit inferential evidence into the framework of knowledge-first epistemology (Williamson 2023, in response to Dunn 2014). Those more technical applications will not be discussed in this paper.

Moral recognitional capacities

One form of attack on moral knowledge consists in challenging moral realists to explain *how* they know moral truths. Their opponents suspect that sooner or later moral realists will have to fall back on a suspiciously convenient faculty of *moral intuition*.

An example: I judge ‘Poisoning Alexei Navalny was wrong’. Is that a moral intuition? I know that Alexei Navalny was poisoned, from various reliable news sources; it is a normal case of knowledge by testimony. I then make the moral judgment that poisoning him was wrong. Although I did not reach my conclusion by conscious step-by-step reasoning, I can provide non-deductive support for it, by citing known circumstances of the case and other considerations. In those respects, it is not so different from my non-moral judgment ‘Poisoning Alexei Navalny was premeditated’. In that case too, although I do not reach my conclusion by conscious step-by-step reasoning, I can provide non-deductive support for it, by citing known circumstances of the case and other considerations. In each case, I applied a new term (‘wrong’, ‘premeditated’) in the light of my evidence. Presumably, I *can* thereby come to know that poisoning him was premeditated. If my judgment ‘Poisoning him was premeditated’ thereby depends on ‘intuition’, presumably so too does my judgment ‘Poisoning him was wrong’, but since one can achieve knowledge by depending on intuition, as in the other case, there is no obvious cause for alarm.

More generally, we have *recognitional capacities* for properties and kinds of many types, enabling us to recognize whether a given case instantiates them. Most people have recognitional capacities for many species of plants and animals and for many types of artefact—types of clothing, types of furniture, types of household utensil, types of vehicle, types of building, and so on. They can recognize various types of weather, various types of art and music, and various types of behaviour: hasty or leisurely, careless or careful, confident or timid, rude or polite, cold or warm, serious or humorous, hostile or friendly, cruel or kind. Non-human animals too have recognitional capacities for various kinds of animal—their own species, predators, prey, males, females, young, mature—and for various kinds of food, building material for nests or dams, and so on.

Linguistic competence with an associated term is not necessary for having a recognitional capacity. A species of birds that nested only in elm trees would have a recognitional capacity for elms. Nor is linguistic

competence with the term sufficient for having the recognitional capacity: by ordinary standards, Hilary Putnam was linguistically competent with the word ‘elm’, but by his own admission he had no recognitional capacity for elms. However, a recognitional capacity associated by a core of users of a natural kind term may play a crucial role in fixing its reference to the natural kind (Brown 1993).

Although many recognitional capacities are fitness-enhancing from an evolutionary perspective, evolutionary constraints put no clear upper bound on what recognitional capacities we can acquire. A chess grandmaster has a recognitional capacity for positions on the board that are a win for black; spending time in Venice, you can develop a recognitional capacity for paintings by Tintoretto. No sound evolutionary argument excludes such recognitional capacities, despite the tenuous relation between evolutionary pressures and chess or art history. Similarly, to expect some evolutionary debunking argument to exclude recognitional capacities for moral properties would betray a remarkably superficial understanding of evolutionary theory. In particular, it is arguable that the properties expressed by moral predicates are not of some radically different metaphysical kind from those expressed by non-moral predicates such as ‘is a win for black’ or ‘is by Tintoretto’ (Williamson 2025, chapter 1).

As the examples above illustrate, typical recognitional capacities are non-reflective. In applying them, we do not use conscious step-by-step reasoning, but our judgment is still evidence-based. To describe recognitional capacities in all those cases as based on a faculty of ‘intuition’ merely obfuscates what sort of pattern recognition is going on.

A different view of particular moral judgments like ‘Poisoning Alexei Navalny was wrong’ is widespread in moral epistemology; its recent advocates include Paul Boghossian and Christopher Peacocke. On such a view, such judgments are *inferential*. They come from arguments like this, where ‘D’ is schematic for a qualitative *non-moral* description:

Premise 1 Poisoning Alexei Navalni was D.

Premise 2 Whatever is D is wrong.

Conclusion Poisoning Alexei Navalny was wrong.

The argument is deductively valid, with Premise 2 read as a universal generalization over all times. Premise 1 is particular, purely non-moral, and *a posteriori*; Premise 2 is general, moral, and supposedly *a priori*; the Conclusion is particular, moral, and *a posteriori*. On this view, the moral element comes from a capacity to assess general moral principles *a priori*. By contrast, on the view defended in this paper, the moral

element comes from a recognitional capacity to assess particular cases morally *a posteriori*.¹

The inferential view of moral judgment is implausible, both psychologically and epistemologically. Psychologically, we are not aware of making any such deduction, or even of entertaining any such universal generalization as Premise 2; nor have we any evidence that what goes on in us unconsciously takes any such artificial form. Epistemologically, we are typically *more* confident of the Conclusion than of Premise 2, because doubts as to whether ‘D’ excludes all potential counterexamples absent from the Navalny case pose a threat to Premise 2 but not to the Conclusion. Thus, even if some of our confidence in the Conclusion comes from Premise 2, not all of it does. Formulating universal moral principles like Premise 2 takes us into a realm of proto-philosophical speculation that we did not enter merely by making a moral judgment on a particular case.

The envisaged inferential alternative fails for a deeper reason too. For where does Premise 2 itself come from? If we are in a position to have ordinary pre-theoretic knowledge of Premise 1, the non-moral description ‘D’ will be couched in reasonably accessible, not too complicated terms. But such a description is likely to make Premise 2 exception-prone and so false, even in our own eyes—poisoning Hitler in 1939 would not have been wrong, and so on. Premise 2 is hardly innate, nor were we taught it as children.

Will moral realists invoke ‘intuition’ here, to explain our putative knowledge of Premise 2? That is what their opponents suspect. What is most dodgy about such reliance on ‘intuition’ is not that it sounds spooky but that it fails to engage with the semantic structure of the sentence. Premise 2 is a universal generalization of the form ‘Whatever is F is G’. Normally, to recognize such a sentence as true, without matching it to a list of pre-given truths, one must recognize some necessary or contingent connection between the two properties, being F and being G (here, being D and being wrong). One can articulate that connection as linking the simpler sentences ‘x is F’ and ‘x is G’, with the variable ‘x’ for an unspecified instance, prior to introducing the universal quantifier (‘whatever’). More specifically, by the informal analogue in natural language of a standard introduction rule for the universal quantifier in a system of natural deduction, having verified ‘x is G’ conditionally on the supposition ‘x is F’, one can then verify the universal generalization ‘Whatever is F is G’. That is logically the most basic way of verifying a universal generalization

1 For an account of such a form see Peacocke 2004: 198–231 and Boghossian’s discussion of moral intuition in Boghossian and Williamson 2020. For other contemporary invocations of moral intuition see Huemer 2005 and Wedgwood 2006.

of that form. To verify Premise 2, it means supposing an unspecified instance ‘*x* is *D*’ and verifying ‘*x* is wrong’ under that supposition. But ‘*x* is wrong’ is just an unspecified instance of the same pattern as the Conclusion, ‘Poisoning Andrei Navalny was wrong’. Moreover, making a judgment under a supposition, imaginatively, is just the offline analogue of making the judgment unconditionally, ‘live’ online, in a particular case. To make such a judgment we need just the sort of recognitional capacity (still for ‘wrong’) invoked by the simpler account above.²

In brief: assenting to a sentence on the basis of intuition was not intended to bypass semantic understanding of that sentence, but instead to work through it; the most natural way of applying that approach to Premise 2 makes the inferential apparatus redundant.

Might our assent to Premise 2 come from some more roundabout process of reasoning? One could postulate that we derived Premise 2 deductively from more basic moral principles, but that only postpones the general problem to those latter principles, on pain of an infinite regress: how do we know the more basic moral principles? Similarly, to postulate that we arrived at Premise 2 inductively or abductively returns us to the original problem, since the inductive or abductive evidence for Premise 2 would presumably itself include further particular moral claims like the Conclusion.

We could try replacing the universal generalization ‘Whatever is *D* is wrong’ as Premise 2 by a merely generic generalization or *ceteris paribus* ‘law’, ‘*D* actions are wrong’, making the argument non-deductive. But that modification does not address the underlying problems, which arise in similar ways for the new Premise 2. Without the recognitional capacity for particular cases, we still cannot verify the generic generalization.³

The upshot of these considerations is that inferential accounts fail to present a genuine alternative, since they will not work without a recognitional capacity, the need for which they merely obscure. Their extra complexity brings no extra explanatory power.⁴

2 For the suppositional assessment of universal generalizations see Williamson 2020: 142–6. For the connection between offline and online judgments see Boghossian and Williamson 2020: 121–2 and 179–81. Many of the epistemological and psychological misconceptions in the literature on ‘moral intuitions’ are analogous to those in the metaphilosophical literature on ‘philosophical intuitions’; I discuss them at length in Williamson 2021.

3 The account of the suppositional assessment of generalizations in Williamson 2020 is applied to generics.

4 Overestimates of the role general principles play in ordinary moral cognition can also have distorting effects on semantic theories of moral discourse. For example, according to Alex Silk’s Discourse Contextualism, ‘Moral uses of deontic modals

Of course, recognitional capacities themselves involve hidden complexity, like the neural processes underlying face recognition. Any recognitional capacity for a moral property, however limited, will involve such hidden complexity too. But the key objection to the inferential account is *not* that it postulates hidden complexity. Rather, it is that the inferential account is *regressive*. For it analyses the cognitive process by which we make a singular ascription of ‘wrong’ into an inference whose major premise (Premise 2) is a generalized ascription of ‘wrong’ (typically a novel one), our assent to which itself depends on a prior singular ascription of ‘wrong’, albeit offline. If every singular ascription of ‘wrong’ depended on a prior singular ascription of ‘wrong’, we could never get started.

By contrast, applying a recognitional capacity involves no such regress. No component of the neural processes underlying face recognition is itself another recognitional process as complex as the one we started with (a homunculus recognizing a face on the screen of a mental cinema). The same applies to our recognitional capacities for moral properties.

Once we have a (limited) recognitional capacity for ‘wrong’, we can use it to start formulating general moral principles and testing them against particular cases, real and imagined. We may adopt some of those principles, and even start using them to correct the deliverances of the original capacity, where we detect a bias or anomaly. We may indeed make some singular ascriptions of ‘wrong’ inferentially, using those principles as premises, as with the argument from Premise 1 and Premise 2 to the Conclusion. That would not be regressive. What triggers the regress is the assumption that *every* singular ascription of ‘wrong’ is inferential.

That we have recognitional capacities for moral properties is no mere postulate of moral realists. Much ordinary moral discourse takes it for granted, for example when we describe someone as recognizing that an action was wrong. Phenomenologically, much ordinary moral judgment (and misjudgement) feels like the exercise of recognitional capacities for moral properties.

[...] presuppose a body of moral norms endorsed in the context’ (2017: 231); a discourse-level parameter in the semantics ‘represents norms accepted for purposes of the conversation’ (ibid: 228). These general norms are needed because ‘The truth-conditional contents of deontic modal sentences are propositions about logical relations (e.g. implication, computability) between propositions and premise sets’ (ibid: 226). Silk takes such general norms to be shared (for conversational purposes) even when speakers have moral disagreement about particular cases. The picture of everyday moral disagreement as fundamentally about logic is a philosopher’s dream. Although one can probably think up some general moral claim or other so bland that all parties to the dispute would assent to it, that does not mean that its sparse logical relations were in dispute. There is no good evidence for the normal conversational presence of such general moral principles.

Naturally, a philosopher can still insist that what we have are *failed* recognitional capacities. For a moral anti-realist, there are no moral properties to be recognized. For an error theorist, the would-be recognitional capacity is like one for the property of being bewitched. But such views are irrelevant to the present dialectic. By hypothesis, there are moral truths, for the question is whether a distinctive problem of moral knowledge still arises under that supposition. It is a short step from moral truths to moral properties within a standard semantic framework, since moral truths require moral predicates to express moral properties in the relevant sense.⁵ Thus we may assume that there are moral properties like being wrong for the associated recognitional capacities to track.

Of course, not all recognitional capacities are good at their job. Someone could hold that, although the word ‘wrong’ does express the moral property of being wrong, the associated recognitional capacity is bad at tracking that property, indeed, is hopelessly unreliable. Some moral realists take our would-be recognitional capacities for moral properties to have been definitively superseded by the clear principles of their favoured moral theory—Kantianism, utilitarianism, whatever. However, such views pose little threat to the present strategy. They hold that one can do much better epistemologically by theoretical reflection on moral principles than by reliance on common sense recognitional capacities. Indeed, their scepticism about the latter is justified by the alleged epistemological superiority of the former when the two conflict. The suggestion is not that the moral domain is epistemically inaccessible, just that it is better accessed by theoretical reflection than by ordinary recognitional capacities. One may doubt that the theoretical reflections at issue are as cogent as claimed, but there is no need to press such doubts here, for either way the moral domain is epistemically accessible.

An alternative kind of argument for the unreliability of moral recognitional capacities invokes individual and cultural variation in their deliverances. Such reasoning grants the associated moral properties, at least

5 In epistemological contexts, we often need to make finer-grained distinctions amongst mental states than can be made simply by ascribing propositional attitudes such as knowledge and belief to coarse-grained intensions. We can do so by relativizing propositional attitudes to the guise under which the thinker entertains the proposition. For example, the guise may be a sentence of the speaker’s language, perhaps in a particular conversational and perceptual context. You can accept the same proposition under one guise without accepting it under another. In order to avoid bogus counterexamples, when cases are described in terms of propositional attitudes, care must be taken to distinguish guise-sensitive from guise-insensitive ascriptions; natural languages can be very misleading in this respect. See Williamson 2024a for detailed discussion and defence of the intensionalist approach. Anyway, none of this affects the connection between truth and properties.

for the sake of argument, and does not privilege some alternative form of epistemic access to the moral domain. The idea is just that if one person or group judges an action as wrong, while another judges it as right, one side is in error (since we are assuming moral realism, no relativistic doublethink is in play). Thus, the extent of moral disagreement puts an upper limit on moral reliability.

Such arguments from disagreement are not illegitimate in principle, but must be applied with care. In particular, we must remember that moral knowledge, like non-moral knowledge, requires only local reliability. For example, suppose that members of society S judge that actions of type A are wrong, while members of society S* judge that actions of type A are right. If members of each society restrict their judgment to actions performed within their own society, so far there is no disagreement, since whether an action of a given type is right or wrong may depend on the social setting: a physical movement may constitute an offensive gesture in S but not in S*. However, suppose that members of each society generalize to the other: members of S judge that instances of A are wrong, whatever society they are performed in, while members of S* judge that instances of A are right, whatever society they are performed in. Now there is genuine disagreement, since the judgments are made about the same domain of actions. At least one side is in error somewhere. Nevertheless, it may still be that members of each society are reliable about whether A-actions performed in their own society are wrong: instances of A performed in S are wrong, while instances of A performed in S* are not wrong. We should not be surprised when people are better at understanding the moral significance of actions performed in their own society than of actions performed in another society, with which they are less familiar. Thus, members of S may *know*, of particular A-actions performed in S, that they are wrong, while members of S* *know*, of particular A-actions performed in S*, that they are right.

That example is obviously very simple and schematic. Sometimes, the asymmetry may run the opposite way, with members of each society blind to its own faults but alert to the faults of its neighbours. The point is just that widespread moral disagreement is quite compatible with widespread moral knowledge of specific matters.

Gestures at the extent of moral disagreement are often quite perfunctory. In surveying the data, one danger to avoid is *cherry-picking*. For example, given any moral belief, one can probably find people somewhere with a contrary moral belief. But that shows nothing special about morality. Human individuals and human societies are very various. Given any *non-moral* belief, one can probably find people somewhere with a contrary non-moral belief. Disputes between proponents and opponents of

the theory of evolution are no more likely to be resolved than disputes between proponents and opponents of a right to abortion. On both moral and non-moral matters, disagreement is typically louder and more attention-grabbing than agreement, making us liable to underestimate the latter's prevalence.

Still, there *is* plenty of moral disagreement. Some philosophers may therefore be tempted to dismiss local patches of correlation between moral belief and moral truth as products of chance rather than cases of moral knowledge. Wouldn't one expect such patches, if moral belief and moral truth varied independently of each other over a large enough domain? As just explained, that attitude may well underestimate the extent of moral agreement. But it also ignores another crucial factor: the *metasemantics* of moral terms. Like other words, 'wrong' does not get its intension by magic. Speakers are willing to apply it in some cases and not in others; the patterns of their use help determine its reference. A principle of charity constrains correct interpretation. Arguably, the appropriate general principle maximizes the *knowledge* expressed by the relevant discourse: for 'right' and 'wrong' in their moral senses, that is moral knowledge.⁶ Although the correct interpretation is the resultant of several forces, correlations between moral belief and moral truth are far from random: they manifest in part a constitutive pressure towards knowledge.

Recognizing moral recognitional capacities also help us make sense of a category that has proved controversial: moral knowledge by *perception*. It has sometimes been regarded as a case where moral epistemology diverges from non-moral epistemology. I will suggest that it is much less exceptional, and much less problematic, than it may sound.

Perceptual language is often used in an extended or metaphorical sense applicable to both sensory and non-sensory cognition, including moral cognition. Someone can *perceive* a research programme as flourishing, or come to *see* that it is degenerating. Likewise, someone can *perceive* slavery as acceptable, or come to *see* that it is wrong. Such examples are no more contentious than moral knowledge; leave them aside. The question is whether there can be moral knowledge by genuinely sensory perception. For example, can one gain moral knowledge *by seeing*, in a sense which requires one to use one's eyes, though of course also one's brain?

Seeing-that must be distinguished from object-seeing. One can see an Anglo-Saxon coin without seeing *that* it is an Anglo-Saxon coin, but

6 For the role of charity in interpretation see Grandy 1973 and Davidson 1977. For a knowledge-maximizing principle of charity see Williamson 2007: 248–78. For a reinterpretation of that principle in terms of an account of mindreading in humans and other animals, see Williamson 2024b.

an expert numismatist can see *that* it is an Anglo-Saxon coin. If we have an encapsulated visual module, impervious to background information, it will not by itself generate judgments like ‘That’s an Anglo-Saxon coin.’ Nor, presumably, will it generate judgments like ‘That’s a giraffe’ or ‘That’s Novak Djoković’. Perhaps it confines itself to shapes and colours. Indeed, by itself, it will generate very little of the seeing-that which we need in science and in ordinary life—potentially including the ordinary moral life.

In the moral case, the issue is not whether one can see special moral objects, but whether one can see *that* an object has a moral property. Sometimes, one can see *that* a kicking was deliberate. Sometimes, one can also see *that* it was wrong. Some philosophers may insist that one only sees that the kicking had various non-moral properties, and then infers that it was wrong by an argument like that above from Premises 1 and 2 to the Conclusion. However, there is no good reason to impose such an inferential structure. Seeing-that is not restricted to contents internal to a primitive module for vision. Sometimes, an expert chess-player sees that a position is a win for black, and an expert boxer sees that his opponent is tiring. By other cues, sometimes a psychologically competent person sees that a kicking was deliberate, and a morally competent person sees that it was wrong.

Seeing-that is a way of knowing-that: when you see that something is so, you thereby know that it is so (Williamson 2000: 33–41). Thus, when you saw that the kicking was wrong, you thereby knew that it was wrong. Consequently, there is moral knowledge by perception.⁷

In those formulations, the words ‘way’, ‘thereby’, and ‘by’ must not be read instrumentally. Seeing that something is so is not a *means* to knowing that it is so. Rather, seeing that it is so is *already* knowing that it is so; seeing-that is a subtype of knowing-that. Seeing that the kicking was wrong is a form of knowing that it was wrong. One’s recognitional capacity for wrongness is triggered visually, but it can also be triggered in other ways, for example by hearing someone describe the kicking in words.

‘She saw that the kicking was wrong’ can be correctly paraphrased as ‘She saw the wrongness of the kicking’. The trouble is that someone may then misinterpret the paraphrase by misreading ‘saw’ in terms of object-seeing in place of seeing-that, casting the wrongness of the kicking as some strange kind of visible object. Such confusions can easily result in the impression that moral perception would be weird. But the confusion is not confined to the moral. Perhaps, in another case, she saw that the kicking was unintentional. ‘She saw that the kicking was unintentional’ can

7 See McGrath 2004 and 2019 for more on more moral knowledge by perception.

be correctly paraphrased as ‘She saw the unintentionality of the kicking’. Someone may then misinterpret the paraphrase by again misreading ‘saw’ in terms of object-seeing in place of seeing-that, casting the unintentionality of the kicking as another strange kind of visible object.

Bad theories of perception can be used to make moral perception look implausible. Fortunately, the badness of those theories emerges in non-moral perception too. Moral perception is unexceptional.

Recognitional capacities and graceful degradation

On a picture associated with the traditional programme of conceptual analysis, most concepts are complex and built up from simple concepts, into which they can be analysed. David Chalmers has offered a sophisticated defence of a similar view, for example in his work on verbal disputes (Chalmers 2011). He proposes to distinguish between verbal and non-verbal disputes by the method of *elimination*: if a word is used in a dispute, try eliminating that word from your vocabulary and see whether you can still have something like that dispute. If the dispute survives, so provisionally does the hypothesis that it is non-verbal: at least, it did not depend on the word at issue. But if eliminating the word eliminates the dispute, that usually indicates that the dispute was merely verbal, since it depended on the word. However, Chalmers allows for exceptions, cases of ‘vocabulary exhaustion’, when the word expresses what he calls a ‘bedrock concept’, and eliminating it reduces the overall expressive power of the language. For example, he suggests, the words ‘conscious’ and ‘ought’ may express bedrock concepts. But he insists that such cases are rare and exceptional. On his view, one can progressively eliminate word after word from the language without reducing its overall expressive power, until one is left with a tiny non-redundant core of exceptional words expressing bedrock concepts, from which all the others can be reconstructed, more or less. That core has the same expressive power as the original language.

How well does Chalmers’ picture fit natural languages? Suppose that you and I see an animal in the distance, which quickly disappears round a corner. I judge: ‘It’s a cat’; you judge ‘It’s a dog’. We are talking about the same animal. It is a prototype of a non-verbal dispute. But what happens if we eliminate the word ‘cat’, or the word ‘dog’, or both, from our vocabulary? We might happen to have the words ‘feline’ and ‘canine’ in our language, with the same meanings as ‘cat’ and ‘dog’ respectively, on which we could fall back to rephrase our dispute, though they are not needed for the original dispute, and in any case the same problem arises for the dispute between ‘It’s a feline’ and ‘It’s a canine’ once the words ‘cat’ and

‘dog’ have been eliminated. Normal speakers of a natural language do not have descriptive equivalents of ordinary natural kind terms in ‘more basic’ terms available to them, as Kripke and Putnam made clear in the 1970s. Of course, if we had pictures of a prototypical cat and a prototypical dog, we could point at them and argue about whether the animal we saw belonged to ‘this species’ or ‘that species’. But that is tantamount to the original dispute only if we recognize the first picture *as* a picture of a cat and the second picture *as* a picture of a dog, in which case we are still implicitly using our recognitional capacities associated with the words ‘cat’ and ‘dog’ in thought, and so undermining the conditions of Chalmers’ thought experiment, which depends on eliminating dependence on the terms at issue from thought as well as speech, otherwise any prototypically verbal dispute could survive the process, with each party tacitly relying on the ambiguous term in thought. In short, within Chalmers’ setup, to count the original dispute as non-verbal, we must count ‘cat’ and ‘dog’ as expressing bedrock concepts, quite subverting his picture of bedrock concepts as forming a tiny, fundamental core. Obviously, ‘cat’ and ‘dog’ are not special in this respect; analogous arguments could be made about thousands of other natural or social kind terms.

Nor is the point restricted to kind terms. Imagine a dispute as to whether a paragraph of prose is ‘subtle’ or ‘banal’. Only a philistine would think that such a dispute must be verbal. Yet, without the words ‘subtle’ or ‘banal’, the disputants may be unable to capture exactly what is at issue. Thus, ‘subtle’ and ‘banal’ come out expressing bedrock concepts too.

What has gone wrong with Chalmers’ approach? He has neglected the association of much ordinary vocabulary with recognitional concepts. Eliminating a word typically involves eliminating the recognitional capacity that goes with it, for cats or dogs, or for the subtle or the banal. Yet much of our cognition depends on such recognitional capacities, deployed online to real-life cases or offline to hypothetical cases. Thus, the method of elimination does not leave us with the same cognitive powers, just honed down to fundamentals. Instead, it turns us into idiots, depriving us of most of the distinctions we need to think about the world as adults.

Why is the destructive effect of Chalmers’ method of elimination not immediately obvious? The reason has to do with the phenomenon of *graceful degradation*, familiar in computer science. It is normally defined as something like the ability of a system to maintain limited functionality even when much of it has been destroyed or incapacitated. The point is to prevent catastrophic failure, or at least to postpone it for as long as possible. Graceful degradation is typical of systems that have evolved naturally or socially. A species of animals or plants which died as soon as they

incurred injury or damage would not be evolutionarily robust. Natural languages too tend to degrade gracefully. They do not collapse when a single word is removed. Speakers find all sorts of workarounds. Thus, one can easily have the illusion that no ‘essential’ cognitive damage has been done. But, in reality, during Chalmers’ process of iterated elimination, the language is being gradually eroded, bit by bit, in a sorites process. What is left at the end is just a stump. The cognitive power of a natural language consists in large part of a vast array of recognitional capacities associated with the terms of that language.⁸

References

- Boghossian, Paul, and Timothy Williamson. 2020: *Debating the A Priori*. Oxford: Oxford University Press.
- Brown, Jessica. 1993: ‘Natural kind terms and recognitional capacities’, *Mind*, 107: 275–303.
- Chalmers, David. 2011: ‘Verbal disputes’, *Philosophical Review*, 120: 515–566.
- Davidson, Donald. 1977: ‘The method of truth in metaphysics’, *Midwest Studies in Philosophy*, 2: 244–254.
- Dummett, Michael. 1973: *Frege: Philosophy of Language*. London: Duckworth.
- Dunn, Jeffrey. 2014: ‘Inferential evidence’, *American Philosophical Quarterly*, 51: 203–213.
- Grandy, Richard. 1973: ‘Reference, meaning, and belief’, *Journal of Philosophy*, 70: 439–452.
- Huemer, Michael. 2024: *Ethical Intuitionism*. New York: Palgrave Macmillan.
- McGrath, Sarah. 2004: ‘Moral knowledge by perception’, *Philosophical Perspectives*, 18: 209–228.
- McGrath, Sarah. 2019: *Moral Knowledge*. Oxford: Oxford University Press.
- Peacocke, Christopher. 2004: *The Realm of Reason*. Oxford: Clarendon Press.
- Silk, Alex. 2017: ‘Normative language in context’, *Oxford Studies in Metaethics*, 12: 206–243.
- Wedgwood, Ralph. 2006: ‘How we know what we ought to be’, *Proceedings of the Aristotelian Society*, 106: 63–86.
- Williamson, Timothy. 2000: *Knowledge and its Limits*. Oxford: Oxford University Press.
- Williamson, Timothy. 2007: *The Philosophy of Philosophy*. Oxford: Wiley Blackwell.

8 Thanks to the audience at the second Balkan Analytic Forum for comments on an early version of this paper. The section on moral recognitional capacities is mainly based on chapter 2 of Williamson 2025.

- Williamson, Timothy. 2020: *Suppose and Tell: The Semantics and Heuristics of Conditionals*. Oxford: Oxford University Press.
- Williamson, Timothy. 2021: *The Philosophy of Philosophy*, enlarged ed. Oxford: Wiley Blackwell.
- Williamson, Timothy. 2023: 'Knowledge-First Inferential Evidence', *The Monist*, 106: 441–445.
- Williamson, Timothy. 2024a: *Overfitting and Heuristics in Philosophy*. New York: Oxford University Press.
- Williamson, Timothy. 2024b: 'Where did it come from? Where will it go?', in Arturs Logins and Jacques-Henri Vollet (eds.), *Putting Knowledge to Work: New Directions for Knowledge-First Epistemology*, 21–70. Oxford: Oxford University Press.
- Williamson, Timothy. 2025: *Good as Usual: Anti-Exceptionalist Essays on Values, Norms, and Action*. Oxford: Oxford University Press.