

Calibrating and Bootstrapping Modal Judgment

Felipe Morales Carbonell

University of Chile

doi: 10.2478/disp-2023-0011

BIBLID: [0873-626X (2023) 69; pp. 250–66]

Abstract

In this paper, I consider the question of whether calibration is required for modalizing mechanisms to be reliable, that is, whether it is necessary for modalizing mechanisms to be adjusted to prevent overgeneration and undergeneration of modal beliefs. I first argue that the calibration requirement affects differently what I call bootstrapping and ordinary cases. Identifying different ways in which a modalizing mechanism could be calibrated, I argue that not all of them are effective or even viable in bootstrapping cases. Then, by taking a diachronic perspective, I offer a simplified account of how different calibration mechanisms can be bootstrapped, with an emphasis on the social dimension of modal judgment.

Keywords

bootstrapping; calibration; modalizing mechanisms; naturalistic modal epistemology; social modal epistemology

1 Introduction

What sort of requirements must a modal belief formation mechanism satisfy, so that we can endorse it as a potential source of modal knowledge or reasonable belief? A common thought is that a viable candidate source has to be reliable. In turn, this will require the satisfaction of further conditions. One plausible requirement for reliability is

Calibration

A reliable modal belief-formation mechanism must be calibrated so that overgeneration and undergeneration of modal beliefs are prevented.

Calibration, in the sense meant here, is the adjustment of some belief forming mechanism to produce beliefs that more closely meet certain standards. The calibration requirement is implicit in the idea that admissible modal judgment requires so-called *epistemic friction*.¹ If we should say that some possible judgments concerning possibilities are incorrect (for example, in cases where the relevant possibilities simply do not hold), a mechanism that vindicated them would be unsuitable, because it would overgenerate modal beliefs. This mechanism would lack enough friction. In different cases, where we should allow for some possibility, a mechanism that prevented us to vindicate it would be adding excessive friction. In both cases, what we are concerned with is the calibration of our modal belief formation mechanisms against some standard of friction; that is, calibration just is a way to make our modal opinions to be delivered under the adequate friction. The question naturally rises: where do we get such standards from?²

One of my goals here is to show that we need to take a step back and reconsider whether Calibration is in fact a requirement for modal judgment; in fact, that will be my main concern in this paper. I will argue that calibration is not required—indeed, that it cannot be required—in an important class of cases. Later, I will adopt a perspective that should also allow us to push back against the need for calibration more generally. The need for calibration has been used to defend modal epistemologies where judgments of a given kind are taken as fundamental because they are arguably indispensable for calibration (for example, Vaidya & Wallner [2021] have argued that concerns of this kind favor an essence-based epistemology). If my argument holds, this strategy is blocked (or at the very least, it is restricted).

¹ In the context of modal epistemology, the epistemic friction requirement has been defended by Vaidya & Wallner [2021]. A more general version of the requirement of epistemic friction is raised by Sher [2010, 2016].

² A reviewer rightly notes that this topic has a strong connection to the so-called frame problem in cognitive science (cf. Dennett [1987]). The form of the problem that is relevant here is what Chow [2013] calls the *Generalized Relevance Problem*, “the problem of determining, from all the information that can in principle bear, what is relevant to the task at hand” ([2013: 312]). Briefly put, modal calibration works by generating solutions to certain kinds of frame problems. This will hopefully be made clearer later, but it is important to keep this observation in mind from the start.

2 Kinds of calibration

To make this point more precise, we need to introduce a distinction between various ways in which belief-formation mechanisms can be calibrated. As a background, I will use a functionalist model of modalizing mechanisms, where such mechanisms are information processing systems/computational methods³ that may take a potentially non-zero number of sources of information as input, and that generate modal beliefs as their output.⁴

Given this basic functional scheme, it is possible to identify two places where calibration can happen. First, it can happen during the processing of information itself. Call this *internal* calibration. Second, it can happen before processing, through restrictions on the inputs of modalizing mechanisms. Call this *external* calibration. The same mechanism can be calibrated both internally and externally.

From this initial typology, we can start to further elaborate on what kind of processes could serve a calibrating role. An important distinction will be whether the relevant processes are under personal control or not, but for our purposes here a perhaps more important difference will be whether calibration is executed conditionally or unconditionally, that is, whether calibration happens given certain triggers or if it happens even in the absence of them. The two distinctions are orthogonal: whether calibration is executed conditionally or not does not necessarily depend on personal control (calibration could happen conditionally at the sub-personal level), although executing personal control can be a way to impose conditional calibration. In the conditional case, we can talk of *dynamic* calibration, whereas in the unconditional case, we can talk of *static* calibration.

Internal calibration can happen both statically or dynamically. In fact, if we take reliability and thus calibration to at least supervene on the likelihood of producing correct beliefs when truths are available and to not produce false beliefs when truths are not available, modalizing mechanisms will always have a degree of calibration just because of their constitution. Call this the *ab ovo* calibration of a modalizing mechanism. This is the baseline calibration of the mechanism. Different mechanisms can have different degrees of *ab ovo* calibration; in order for some mechanisms to be reliable, some additional

³ I use ‘computational method’ along the lines of Knuth [1997: 5], in opposition to ‘algorithm’, because for the methods we are dealing with there is no guarantee that they will terminate in a finite number of steps.

⁴ Here I use this model only for expedience; I believe what I will say here can be reconstructed in other terms (I will sketch one way to do so later, in fact), and I don’t want to commit to a computational model of the mind.

way of calibrating the mechanism will be necessary, as we will see. In the case of human modalizing mechanisms, it is implausible that *ab ovo* calibration will suffice in every case. Surely, ordinarily our modalizing mechanisms are not infallible.⁵ Nevertheless, calibration *ab ovo* will become important later.

Other internal methods of calibration are dynamic. In this case we will speak of *in situ* calibration. If we think of the belief-formation mechanism as a sub-personal process, *in situ* calibration happens when an instance of calibration is a stage of that process—that is, if it happens ‘in the black box’. For example, it could be that in order to calibrate the operation of the mechanism, previous modal knowledge needs to be recalled at some point during the operation of the mechanism, before it results in the formation of beliefs. The case of pulling in the knowledge that one cannot climb a tree while sleep while one considers a question of how one can climb a tree is an example of *in situ* calibration. As part of *in situ* calibration, an instance of the belief-formation mechanism in question could be executed recursively: in absence of previous judgment on a given modal matter that is relevant to calibration, an opinion on that matter might need to be formed. This process can iterate arbitrarily, although it is important to note that this by itself does not lead to regress.

On the other hand, calibration can happen outside the operation of the belief-formation mechanism properly speaking. Call this *ex situ* calibration.⁶ *Ex situ* calibration is typically temporally prior to instances of the relevant belief formation mechanism. In these cases, calibration happens indirectly: rather than constraining the range of possible outputs of the mechanisms by altering the internal operation of the mechanisms, in these cases the range of *inputs* is constrained by not making contextually determined irrelevancies available in the

⁵ It is unclear if we have any general belief-formation mechanism that is well-calibrated *ab ovo*. In the case of perception, Millar [2019] makes use of the notion of a perceptual-recognitional ability, and argues that an ability of this kind is exercised only when the subject succeeds in recognizing its target; this seems to entail that the mechanism of a proper exercise of these abilities is calibrated *ab ovo*. Accordingly, in this view, to fail to recognize perceptually something as what it is is not a failure of calibration, but an instance of having exercised a different mechanism. This could be important for empiricist models of modalizing, like in Strohming [2015].

⁶ The distinction *in situ/ex situ* is related to the internalism/externalism debate concerning modal epistemic friction. Cf Vaidya & Wallner [2021: 1920–5].

situations where the mechanisms operate.⁷ *Ex situ* calibration also includes the uptake of assorted knowledge from means like testimony, which in turn serves as a background for the evaluation of modal judgments in other contexts. In these cases, we defer to the calibration of other epistemic subjects in order to adjust our own belief-formation mechanisms. This reliance might or not be something we are aware of; I will return to the point later. In the case of *ex situ* calibration, the distinction between static and dynamic modes of calibration is less straightforward to maintain, so I will not push the point.⁸

In summary, across these two dimensions we identify three types of calibration modes, as in the following table:

	Internal	External
Static	<i>ab ovo</i>	
Dynamic	<i>in situ</i>	<i>ex situ</i>

3 Calibration in bootstrapping and ordinary scenarios

The manner in which calibration takes place bears on the structure of our modal epistemological accounts. In ordinary cases both *in situ* and *ex situ* methods of calibration are viable, and they complement each other. For example, for much of our ordinary modal knowledge concerning how tools can be used, we rely on both the testimony of others (which provides *ex situ* calibration as well as a source of modal knowledge on its own right) and our own ability to recognize the ways in which we can or cannot use them in a given circumstance (for *in situ* calibration). In bootstrapping cases, however, the situation is

⁷ This description makes the connection to the frame problem quite evident in this case (*cf* fn 2), but it is important to see that both *in situ* and *ex situ* calibration modes offer alternative ways to solve the frame problem. In the case of *in situ* calibration, irrelevancies are filtered as part of the modalizing mechanism itself—in fact, this is the canonical context where the frame problem appears (where it typically takes the form: how does the cognitive agent manage to filter out irrelevances given that it is exposed to an environment that makes them available?).

⁸ One way in which we could try to make the distinction is as follows. Static *ex situ* calibration is purely environmental, without influence from the subject (for example, a more regular environment could provide more constraints for calibration). Dynamic *ex situ* calibration could then happen with influence from the subject (for example, the subject could probe and alter its environment).

different. In genuine bootstrapping cases calibration cannot rely on previous modal knowledge since such knowledge will be *ex hipotesi* scarce. It is possible, however, that because of the way modalizing mechanisms are constituted, they could only be well-calibrated *in situ* with the use of modal-bearing input. In those cases *in situ* calibration wouldn't be viable. Since the scarcity of modal-bearing knowledge is orthogonal to the mode of calibration, in bootstrapping scenarios *ex situ* calibration cannot work on the basis of a background of available modal opinion either (for example, the mechanism cannot be calibrated merely by putting the subject in an environment where suitable modal opinions are available, since the environment does not contain sufficient subjects with appropriate modal opinions).⁹ So the only alternative in these cases is to calibrate in relation to non-modal opinion.

Now, *in situ* calibration requires the measure of some intermediate product of the belief formation mechanism against some standard. A natural way to put the point is that in the case of *in situ* calibration the question is considered whether certain *constraints* are respected.¹⁰ But if calibration requires judgments concerning constraints in this way, *in situ* calibration might by itself require previous modalizing capacities in all cases: the very notion of a constraint is plausibly modal.¹¹ Something satisfies a constraint only if it is compatible with the constraint not being broken. But judgments about compatibility are modal: that A is compatible with B means that it is possible for A and B to be the case together. Similar points can be made in case we spell the relevant notions in terms of incompatibility. If this is the right way to think about *in situ* calibration, then it simply cannot be used in bootstrapping scenarios, even if the base of opinion that it operates on is not-modal. To put the point more directly: *in situ* calibration is itself a modalizing capacity

⁹ Note that consulting others for their opinions during a modalizing task is a form of *in situ* calibration, albeit requiring an external resource.

¹⁰ When I say that the question is considered, I only mean that information concerning the question is used in the operation of the mechanism, that is, that its implementation requires information on the question. I don't mean to suggest that a special (conscious) mental act of raising the question is necessary.

¹¹ This would be particularly important for approaches where knowledge of compatibility and incompatibility is foundational for modal knowledge more generally.

that needs to be bootstrapped.¹² The problem does not arise for *ex situ* calibration, since it can happen purely environmentally.¹³

We can summarize the preceding observations concerning the viability of ways to calibrate modal belief-formation mechanisms in different scenarios with the following table:

	In situ	Ex situ	Ab ovo
Bootstrapping	✗	✓	✓
Ordinary	✓	✓	✓

There is a substantive difference in terms of the possibilities for calibration in ordinary and bootstrapping scenarios. *In situ* calibration cannot be used in bootstrapping scenarios. This is the key to our problem.

4 Calibration through time

The question that remains then is how modal judgment (including *in situ* calibration) can be bootstrapped merely on the basis of *ex situ* and *ab ovo* calibration. To tackle this question, we have to move from what is a fundamentally *synchronic* picture of calibration (one where we only consider what methods of calibration are available to a subject at a given point in time) to one that is *diachronic* (where we also consider progressive changes to the methods).

Here is an sketch of an account of how *ex situ* and *ab ovo* calibration can bootstrap modal judgment. Imagine a population of agents with some minimal cognitive capacities, who are able to track objects in their environment and some of their properties (size and color). The objects in their environment have some modal properties that might be of interest

¹² If *in situ* calibration is bootstrapped for a (type of) modalizing agent, does that also mean that they have bootstrapped other modalizing capacities? In principle, there is no reason to think that what bootstrapping *in situ* calibration requires is any different to what any other modalizing capacity requires—it would need to be shown that this capacity is somehow exceptional (cf Williamson’s argument against modal exceptionalism). This can go both ways: by acquiring the capacity to modalize in some other way, an agent might acquire the capacity to calibrate *in situ*.

¹³ Going back to the point about the frame problem, what these observations about *in situ* calibration suggest is that we are dealing with two frame problems: one is the general one that concerns any form of modal calibration, and the other is the narrow problem of how *in situ* calibration can be done in a relevance-sensitive manner.

to the agents; for example, they can be arranged in certain patterns that can be useful. Suppose that the agents engage in a practice of arranging the objects. Suppose that there is some cost to engaging in this practice. It makes sense for these agents to maximize the usefulness of the produced arrangements, in order to compensate for the cost. So there is some pressure for the mechanism that plans the arrangement of the pieces to maximize usefulness. Suppose that this mechanism involves the evaluation of possible arrangements, that is, a form of modalizing. For the same usefulness maximizing reasons, it is desirable that the modal-belief formation mechanism involved is calibrated. Now, suppose that the agents themselves can exert no control on their belief-formation mechanism: once they are triggered to form a modal belief given some sensory input, their opinion on the matter will not change. How can the mechanism be calibrated? This cannot happen at the level of the individual agents, but it can happen at the level of the population by the selection of better or worse modalizers through a process of Darwinian selection, or through a process of niche construction.¹⁴ These methods constitute purely *ex situ* ways to calibrate the agents' modal belief-formation mechanisms. More precisely, the product of this kind of calibration is an improvement of *ab ovo* calibration, which in turn affects the conditions for subjects to apply different *ex situ* and *in situ* calibration strategies.

It is important to note that the result of this calibration process can turn out to be not exceedingly reliable when tracking modal truth in a general sense; for this reason, it might not be a way to produce a way to acquire the capacity for modal *knowledge* properly speaking.¹⁵ However, for our purposes here it is sufficient to see that *ex situ* calibration can apply to a mechanism that tracks modal features of the world and that more importantly can generate representations of modal matters. Once something like this is in place, we can move on to an account of how from this basis something like reliable modal judgment can evolve. It is important to realize that the former does not immediately entail the latter.

As suggested above, this is at least part of the story when it comes to explaining how come subjects like us have the capacity to calibrate *in situ*. We may hypothesize that at some point having the trait of filtering out irrelevancies as part of some belief formation mechanism that bore on modal matters was selected because it made some positive difference

¹⁴ Cf. Martinez [2015] for an account along these lines of how modal sensitivity can emerge in a population. On evolutionary strategies to account for modal knowledge, see also Kroedel [2017] and Fischer [2017: 271]. On how to conceptualize niche construction, see Aaby & Ramsey [2022].

¹⁵ Cf. Dohrn [2017], who raises this as an objection to Martinez.

on fitness. For example, it could be that the environment changed quickly enough that it outran the capacity of members of the relevant population to calibrate merely *ex situ*. Then, the part of the population that had the ability to calibrate conditionally *in situ* gained an advantage, and the trait got reproduced.

To reiterate, the general idea is that a richer account of the capacities of agents that are able to modalize in a minimal manner can provide an explanation of the emergence of *in situ* calibration and more complicated modalizing mechanisms over time.

It may be helpful to describe a different way in which this could have gone, that will help us highlight some other aspects of what is at stake here. One may think that, once some minimal capacities for modalizing are available, the cognitive environment in which these agents live may get enriched with information (and misinformation) that bears on modal matters. Assuming that agents in the population can interact with each other, they can start acting in ways that calibrate the modalizing capacities of the population—for example, by sharing the result of their modalizing with other members of their communities. More complex types of agents might be able to update their modal opinions by using general capacities to revise their actionable information about matters of interest. The resulting acquired capacities to modalize would be similar, in some respects, to the operation of *in situ* calibration. In fact, we can describe what happens in these scenarios in two ways, not necessarily exclusive of each other: on the one hand, we may think that in this case individuals engaged in these social practices acquire the capacity for *in situ* calibration by internalizing directly some of the social work (for example, by imagining the execution of social calibration of modal opinion), and on the other, we may think that in these cases, something like group-level *in situ* calibration emerges (the execution of social calibration just is the execution of a parallel to a modalizing process that has the group as a whole as its subject), which then individuals are able to replicate. In either case, this shows that at least two things are true about the emergence of *in situ* calibration: (a) it could begin as the application of mechanisms that are not specific for modalizing purposes for limited modalizing purposes, and (b) it could arise from internalizing what initially was a merely external mechanism.

Such processes could go in tandem with the development of more sophisticated communication methods, allowing for yet more complex forms of calibration. For example, an agent could convince another to calibrate *in situ* assuming that a modal fact holds even though the other agent has no means of their own to check this fact. Thus, depending on the capacities and information available, there will be a more or less gradual passing from

bootstrapping scenarios to ordinary ones.¹⁶

Through processes like these, belief-formation mechanisms that might have initially been one-stage (that is, that goes from input to output directly in one step) give way to mechanisms of a higher complexity. In particular, one-stage modal belief-formation mechanisms that are *ex situ* calibrated give room to multi-stage modal belief-formation mechanisms that *can* be *in situ* calibrated (multi-stage mechanisms where some stages fulfill a calibration function). Adapting some terminology from Johnson-Laird [1987, 2002, 2009], we can identify three kinds of mechanisms that are of special interest here: *neo-Darwinian*, *neo-Lamarckian*, and *hybrid*. A neo-Darwinian mechanism of the sort we are interested in operates in two stages: first, it generates belief-like contents indiscriminately, and then these are filtered by some selection mechanism that responds to constraints. The selection process plays the role of calibrating the mechanism, and it can work either *in situ* or *ex situ*.¹⁷ A neo-Lamarckian mechanism applies criteria or constraints during the generation process itself (that is, it calibrates the output of the generation sub-process *in situ*), and selects arbitrarily among the viable results. In contrast to both, the hybrid model is a multi-stage process: it iterates on a two-stage loop of generating belief-like contents and filtering them according to various criteria, to then select some generated belief-like contents either arbitrarily or according to some criteria. Neo-Darwinian mechanisms tend to be inefficient, but are the only viable method to bootstrap modal judgment on the absence of opinion that could be used to inform the calibration of our modal belief-formation mechanisms or more generally of the conditions to make use of a more guided calibration mechanism.¹⁸ The more advanced kinds of modalizing that we use in ordinary cases is more deliberate

¹⁶ Evolutionary processes might be able to account for the existence of modal belief-formation mechanisms that are well-calibrated *ab ovo*, as traits of the members of a population that evolved to make some class of modal judgments infallibly or much better than the baseline in the sense described above. It is likely that if such mechanisms exist, they are highly specialized rather than general (for example, mechanisms to pickup on the affordances of middle-sized objects in ‘normal’ environments could be close to infallible in this sense and match the description of mechanisms that are calibrated *ab ovo*). Since the development of these mechanisms could be a phylogenetically early adaptation, agents like us could have made use of mechanisms like these (not necessarily ones producing modal beliefs) in bootstrapping scenarios for more sophisticated forms of modalizing.

¹⁷ In fact, the bootstrapping scenario described above corresponds to an instance of this neo-Darwinian model, with the *ex situ* calibration being simply the selection of agents that are fitter because of the characteristics of their particular modal belief-formation mechanisms.

¹⁸ Cf Johnson-Laird [2002: 421].

and incorporates elements of the neo-Lamarckian or hybrid models.

5 Social calibration, anti-exceptionalism and inefficiency

An interesting feature of these kinds of models is that they allow for modalizing mechanisms to be distributed over a population. On the condition that communication between the members of the relevant populations is possible, different members can take on different roles in the generation of modal and modal-bearing opinion and calibration, and thus *modalize together*. Certainly, *we* (namely, people who are interested in providing philosophical accounts of a range of subject matters) are part of a population where modalizing work is not only individual but also social.¹⁹ This is something we should expect, because modal matters are not only a concern for individuals, but also for collectives. Action is guided by considerations about what is possible for agents to do; if we allow groups to be agents on their own right, we should expect group action to involve some form of modalizing.²⁰ This goes beyond the point that individuals can modalize by pooling and coordinating resources from other members of their communities. However, in both cases we should see that modalizing is guided by practical concerns that go beyond what is purely epistemic. From this more pragmatic point of view (and on this I draw on Haugeland [2017]), we need to be mindful of why we care about modal judgment in the first place, that is, why we think we *should* care. Van Inwagen [1998] chided philosophers for making modal judgments that had their source in “professional socialization, in ‘what his peers will let him get away with saying’” ([1998: 73]). Of course! But those allowances serve certain purposes; the skeptical argument seems to hinge on whether certain goals should be pursued or not. In this, the skeptic also plays an important role in the collective modalizing task, which is to keep things honest. In doing this they calibrate our collective modalizing.

If what I have suggested here is right, we should notice that the cost for bootstrapping has, in some sense, been already paid for us in advance. We reap the rewards by being able to modalize fruitfully in relation to many subject matters and for many different purposes. However, the rational demands for further inquiry into ever more complex modalizing

¹⁹ This is another point that most modal epistemologists seem to me to have under-appreciated. One exception to this trend can be found in modal normativism (*cf* Thomasson [2020]), but I think Haugeland [2017] hints at a more satisfactory approach.

²⁰ On the possibility of treating groups as agents, see Tollefsen [2015]. For an approach that also places the source of our capacity to modalize on broadly practical concerns, see Vetter [2023].

puts a strain on the baseline capacities, for some of the modal questions we raise we might need capacities that are presently out of reach.²¹ On the supposition that we are ignorant of at least some things that we could treat as a subject matter for our inquiries, there is no reason to expect the need for modal recalibration to stop.

This last point has consequences for the exceptionalism/anti-exceptionalism debate. Williamson [2007] and others have made a strong case for anti-exceptionalism on the grounds that our ordinary capacities can in great part account for our capacity to engage in advanced modalizing (for example, by pointing out that we can appeal to the same mechanisms of belief-formation when we are doing metaphysics and when we try to haul a piano).²² The motivation behind this form of anti-exceptionalism was to counter certain unpalatable forms of modal skepticism, that separated too strictly the capacities required for everyday modalizing and advanced modalizing. What I am suggesting here, however, is that we cannot hold on to anti-exceptionalism definitely as a substantive thesis about modal knowledge (namely, the thesis that the sources of reasonable modal opinion must be non-exceptional). We must leave the possibility open that there are modal questions that are outside the epistemic reach of our current ordinary capacities, but not absolutely out of reach. It would be surprising if, like the sailors of Neurath's raft, we were in a position where we needed to rebuild our ship at sea, so to say, but where we had the best materials to do so from the very beginning, not knowing in advance what the challenges will be.²³ Of course, we can turn this issue around and ask if we really could offer an answer to every challenge that may present itself, and to *that* question it may be legitimate to give a negative answer.

The increased complexity comes with advantages, but also costs. One of the charges that Roca-Royes [2011] raises against Williamson [2007] is that a model where fixing constitutive facts is required is more cognitively expensive than one where this is not the case, and consequently less naturalistically plausible. This line of thought can be construed as an argument against the need for *in situ* calibration, or against certain forms of *in situ*

²¹ Cf the work by the late Paul Humphreys on "epistemic enhancers" (Humphreys [2004]) and Sterelny's [2003, 2010] ideas on epistemic engineering and scaffolding. If there existed hitherto undiscovered or undeveloped modalizing mechanisms, coming to possess them would extend our epistemic modal reach.

²² Cf Strohming & Yli-Vakkuri [2018].

²³ Neurath's [1956: 201] own introduction of the metaphor strikes me as a more apt description of our epistemic situation: 'we are like sailors who must rebuild their ship on the open sea, never able to dismantle it in dry-dock and to reconstruct it there *out of the best materials*' (my emphasis).

calibration. However appealing, the point can be resisted. Inefficient mechanisms can be selected if they do not reduce fitness, or as by-products of other traits. Initial cost can also be recovered in some sense with later gains. This also undercuts an argument against the viability of relying on *ex situ* calibration for bootstrapping. Reliance on purely *ex situ* calibration is inefficient: it takes generations for mechanisms to be calibrated, so the cost for early adopters is high, and there is plenty of risk. But on the long run, the cost is recovered because the adoption of those mechanisms paves the way to other potential adaptations that offer more efficient calibration methods. To make the appropriate calculation, as I have argued above, we have to look beyond the costs and benefits for individual agents, and focus instead on populations across generations.²⁴

6 No royal way from calibration to the necessity of sources?

Someone could be tempted by an argument like this:

- 1) We have reliable capacities to form epistemically justified modal opinions,
- 2) It is part of the reliability of those capacities that they are well calibrated,
- 3) In order for those capacities to be well calibrated, their application needs to make use of knowledge of type X,
- 4) If knowledge of type X is necessary for knowledge of type Y, knowledge of type X is foundational to knowledge of type Y,
- 5) We must have foundational knowledge of type X (from 1, 2, 3, 4)

Premise 1) is a statement of modal non-ignorance, which we may grant (or we may simply suppose it and conclude that (5*) if we have reliable capacities to form epistemically justified modal opinions, we must have foundational knowledge of type X). Premise 2) is the supposition that calibration is a requirement for reliability. Premise 3) asserts that the calibration of the capacities that we rely on to form modal opinions must be *in situ* with relation to some specific type of knowledge X. The idea could be that calibration requires sensitivity to a certain class of fact, so that invoking knowledge about facts of

²⁴ Note that this does not undermine Roca-Royes' arguments against the necessity of constitutive thought for the calibration of counterfactual evaluation. Deploying constitutive thought might be a more efficient way to calibrate modal belief-formation than not using it, at least by some standard. But this does not mean that less efficient mechanisms are not viable. Nature is under no obligation to prefer efficiency if it does not increase fitness, and support for a positive correlation in this case is tenuous. Cf also Roca-Royes [2012].

that type is a reliable way for the modalizing mechanism to be sensitive to those facts.²⁵ Premise 4) derives from a plausible definition of foundational knowledge. The conclusion is a substantive thesis about the structure of modal knowledge. The argument shows how considerations concerning calibration can be used to motivate architectural claims in modal epistemology.

The possibility to bootstrap modal judgment purely on the basis of *ex situ* and *ab ovo* calibration, along with the observation that reliability should be assessed at the population level, suggest that arguments like this are in general unsound. First, because calibration does not require of the application of knowledge of any kind, since it can happen purely *ex situ*. This knocks down premise 3). In fact, it seems to knock down the idea that modalizing mechanisms need to be under explicit rational control, at least in this sense. Second, adopting a population-centred perspective offers a notion of reliability that does not require the calibration of the modalizing mechanisms of individuals. That a population is reliable in forming modal opinions can differ from facts about how reliable any of its members is to form them. So *we* (collectively) can have reliable ways to modalize without it being true of anybody that their individual modalizing mechanisms are reliable. This puts a restriction on the scope of premise 2). The calibration of modal judgment in a population may not be a matter of individuals having the means to generate knowledgeable modal opinions, but of the whole population having the means to check the modal opinions made public by their members. Rather than constraining the generation of modal contents, the focus will be in this case on selecting the opinions that are deemed worthy for the population to hold on. As the reader can tell, this is just a population-level neo-Darwinian model of modalizing mechanisms. If it is plausible as an alternative, then we cannot take the royal way from the requirements for individual *in situ* calibration to substantive claims about the architecture of modal judgment.²⁶

²⁵ Some candidates for X would be constitutive knowledge, as in Williamson [2017] or essentialist knowledge, as in Vaidya & Wallner [2021]. In Williamson's model, the reliability of the counterfactual-based modalizing mechanism relies on the availability of constitutive knowledge that can be used to ground the development of counterfactual suppositions. In Vaidya & Wallner's model, the reliability of modalizing mechanisms depends on the sensitivity of those mechanisms to the nature of things, which serves as a source of modal epistemic friction.

²⁶ There is no great prospect to trying to come up with an argument of this kind at the population level, I feel. The argument against the second premise is equally strong for the collective and individual cases.

7 Conclusion

Let us take stock. We have raised the question of whether calibration is in some necessary for modal judgment. I have argued that if we distinguish between ordinary and bootstrapping cases, we should expect the requirement of calibration to hold differently depending on what case we are dealing with. For bootstrapping cases, in particular, *in situ* calibration is not viable. This leaves us with the question of how *in situ* calibration can be itself bootstrapped. I have suggested that in order to tackle this question, we need to adopt a diachronic perspective, and think of calibration as something that evolves over time. I think the lesson can be generalized: modal epistemology needs to pay more attention to the resources that evolutionary theory provides. I have also argued that it needs to consider the idea that modalizing can be seen as a collective process. I believe that in doing so, we will advance towards a more richly naturalized modal epistemology.²⁷ Perhaps such approach could also offer insight in the ways in which our modalizing could not be only understood, but also improved.²⁸

Felipe Morales Carbonell

University of Chile

Facultad de Filosofía y Humanidades

Av. Capitán Ignacio Carrera Pinto 1025,

7800284 Ñuñoa, Santiago, Chile

ef.em.carbonell@gmail.com

References

- Aaby, Bendik Hellem & Ramsey, Grant [2022]. “Three kinds of niche construction”. *The British Journal for the Philosophy of Science* **73**: 351–72.
- Chow, J. Sheldon [2013]. “What’s the problem with the frame problem?”. *Review of Philosophy and Psychology* **4**: 309–31.

²⁷ Cf Nolan [2016].

²⁸ This paper was funded by the project ANID Fondecyt de Postdoctorado 3220017. Many thanks to Sonia Roca-Royes, Michael Wallner, Tom Schoonen, Giulia Lorenzi and Rodrigo Gonzalez for their feedback, as well as to two anonymous reviewers for their comments.

- Dennett, Daniel [1987]. "Cognitive wheels: the frame problem of AI". In *The Robot's Dilemma: The Frame Problem in Artificial Intelligence*, edited by Zenon Pylyshyn. Norwood, NJ: Ablex Publishing.
- Dohrn, Daniel [2017]. "Nobody bodily knows possibility". *Journal of Philosophy* **114**: 678–86.
- Fischer, Bob [2017]. "Modal empiricism: Objection, reply, proposal". in *Modal Epistemology After Rationalism*, edited by Bob Fischer & Felipe Leon. Cham: Springer: 263–80.
- Haugeland, John [2017]. "Two dogmas of rationalism". In *Giving a Damn: Essays in Dialogue with John Haugeland*, edited by Zed Adams & Jacob Browning. Cambridge, MA: The MIT Press: 293–310.
- Humphreys, Paul [2004]. *Extending Ourselves: Computational Science, Empiricism and Scientific Method*. Oxford: Oxford University Press.
- Johnson-Laird, N. Philip [1987]. "Reasoning, imagining, and creating". *Bulletin of the British Psychological Society* **40**: 121–9.
- Johnson-Laird, N. Philip [2002]. "How Jazz musicians improvise. *Music Perception: An Interdisciplinary Journal* **19**: 415–42.
- Johnson-Laird, N. Philip [2009]. *How We Reason*. Oxford: Oxford University Press.
- Knuth, Donald [1997]. *The Art of Computer Programming, Vol I: Fundamental Algorithms*, 3rd edition. Reading, MA: Addison Wesley.
- Kocurek, W. Alexander [2021]. "Counterpossibles". *Philosophy Compass* **16**: e12787.
- Kroedel, Thomas [2017]. "Modal knowledge, evolution, and counterfactuals". In *Modal Epistemology After Rationalism*, edited by Bob Fischer & Felipe Leon. Cham: Springer: 179–196.
- Martinez, Manolo [2015]. "Modalizing mechanisms". *Journal of Philosophy* **112**: 658–70.
- Millar, Alan [2019]. *Knowing By Perceiving*. Oxford: Oxford University Press.
- Neurath, Otto [1959]. "Protocol sentences". In *Logical Positivism*, edited by A.J. Ayer. New York: Free Press; 199–208.
- Nolan, Daniel [2017]. "Naturalised modal epistemology". In *Modal Epistemology After Rationalism*, edited by Bob Fischer & Felipe Leon. Cham: Springer: 7–28.
- Roca-Royes, Sonia [2011]. "Modal knowledge and counterfactual knowledge". *Logique et Analyse* **54**: 537–52.
- Roca-Royes, Sonia [2012]. "Essentialist blindness would not preclude counterfactual knowledge". *Philosophia Scientiae* **16**: 149–72.
- Sher, Gila [2010]. "Epistemic friction: Reflections on knowledge, truth and knowledge". *Erkenntnis* **72**: 151–76.
- Sher, Gila [2016]. *Epistemic Friction: An Essay on Knowledge, Truth, and Logic*. Oxford:

- Oxford University Press.
- Sterelny, Kim [2003]. *Thought in a Hostile World: The Evolution of Human Cognition*. Oxford: Routledge.
- Sterelny, Kim [2010]. "Minds: Extended or scaffolded?". *Phenomenology and the Cognitive Sciences* **9**: 465–81.
- Strohming, Margot. [2015]. "Perceptual knowledge of nonactual possibilities". *Philosophical Perspectives, Epistemology* **29**: 363–75.
- Strohming, Margot & Yli-Vakkuri, Juhani [2019]. "Knowledge of objective modality". *Philosophical Studies* **176**: 1155–75.
- Tahko, Tuomas [2012]. "Counterfactuals and modal epistemology". *Grazer Philosophische Studien* **86**: 93–115.
- Thomasson, L. Amie [2020]. *Norms and Necessity*. New York, NY: Oxford University Press.
- Tollefsen, Deborah [2015]. *Groups as Agents*. Cambridge: Polity.
- Vaidya, Anand and Wallner, Michael [2021]. "The epistemology of modality and the problem of modal epistemic friction". *Synthese* **198**: S1909–35.
- van Inwagen, Peter [1997]. "Modal epistemology". *Philosophical Studies* **92**: 67–84.
- Vetter, Barbara [2023]. "An Agency-Based Approach to Modal Epistemology". In *The Epistemology of Modality and Philosophical Methodology*, edited by D. Prelević & A. Vaidya. New York, NY: Routledge: 44–69.
- Wieland, Jan Willem [2014]. *Infinite Regress Arguments*. Dordrecht: Springer.
- Williamson, Timothy [2007]. *The Philosophy of Philosophy*, Oxford: Oxford University Press.