

The Vehicle Indeterminacy Problem

Preprint

Caitlin Mace

cbm49@pitt.edu

University of Pittsburgh

Abstract: Proponents of some recent accounts of neural representations avoid classic issues with content indeterminacy by taking an antirealist or pragmatic stance toward representational content and a realist stance toward representational vehicles (e.g., Shea 2018; Egan 2020; Burnston 2020). To establish the reality of representational vehicles, such vehicle realists must rely on the fact that patterns of neural activity are putatively re-identified as the same pattern across systems. Re-identification shows that the pattern exists across means of detection, which shows that the pattern is robust and, therefore, real. But it is entirely unclear by what measure patterns should be determined as similar. Drawing on neuroscientific expertise, I show that determining the relevant details by which to assess similarity is done in practice by coarse-graining vehicles, i.e., compressing, abstracting, or suppressing variations in patterns. I conclude that this coarse-graining strategy vindicates neuroscientific practice but not realists' appeal to robust vehicles.

Keywords: neural representations, scientific realism, representational vehicles, indeterminacy

Neuroscientists often appeal to neural representations in the brain to explain experimental results. These posits are patterns of neural activity that carry information about the world. The past few decades have seen much philosophical debate about how to determine what precisely the information being carried by these patterns is about – e.g., is the neural activity in this frog that responds to flies about flies or food or small moving black dots? This is called the ‘content indeterminacy problem’. Proponents of some recent accounts of neural representations avoid classic issues with content indeterminacy by taking an antirealist or pragmatic stance toward representational content and claiming that only representational vehicles – which are the particular patterns that bear representational content – are real (e.g., Shea 2018; Egan 2020, 2025; Burnston 2020). Why be a vehicle realist? On one account, the neural patterns identified as representational vehicles are shown to be robust. Robustness indicates that the patterns exist

across different means of detection (Wimsatt 2007). Such re-identification suggests that those patterns are not due to error or bias and are, therefore, real. I argue that vehicle realists have their own indeterminacy issue to resolve to establish that such re-identification actually occurs. The vehicle indeterminacy problem in cognitive neuroscience is the problem of identifying robust representational vehicles across experimental contexts. The problem is that it remains underdetermined which properties are being re-identified and whether they are, indeed, similar in the relevant sense. Neuroscientists seem able to deal with this issue. In taking their cue, I show how researchers solve this problem in practice by coarse-graining vehicles to make them re-identifiable. I conclude, however, that these research strategies do not vindicate vehicle realism in philosophy.

In §1, I introduce and raise problems for the re-identification criterion for vehicle realism to motivate the vehicle indeterminacy problem. In §2, I consider three alternatives to robustness analysis for establishing real entities. But upon further analysis, these strategies themselves rely on re-identification. In §3, I turn to how neuroscientists identify vehicles in practice to see if neuroscience has the resources to support vehicle realism. In §4, I develop an account of the coarse-graining strategy that neuroscientists use to make vehicles re-identifiable. I conclude in §5 that the coarse-graining strategy vindicates neuroscientific practice but not vehicle realism.

1. Vehicle Indeterminacy

In scientific practice, entities are considered real to the extent that they can be observed over time, in different situations or contexts, with different measurements or tools, and by different people. The thought is that if one can re-identify the entity in various ways, then the entity is multiply determined and, therefore, robustly detected (Calcott 2011). With robust detection, there is reason to trust that observations of the entity are not due to error or bias, and that the entity

itself is unlikely to be an artifact of any particular means of detection. Robust detection, then, is a means of confirming that the entity is real.

In neuroscience, where the target is finding patterns of neural activity that perform various roles in perception, cognition, memory, and behavior, those patterns should be re-identifiable in other brains, in other species, and with different experimental paradigms (Cao 2022). We might even hope to find those patterns with different tools or in the same brain over time. Patterns of neural activity are whatever patterns are identified by various neuroscientific probes that are of interest to researchers; they may be active and inactive neurons at a particular time in a populations of neurons, active neurons in a neuronal ensemble, neural firing during a particular temporal block, electrophysiological properties of whole brain regions, among many other kinds of patterns. Vehicles, those patterns of neural activity that perform a representational role for the system, come in a similar variety; they may be identified as spatial particulars, such as neurons, or temporal particulars, such as neuronal firing.

Identifying robust representational vehicles across experimental contexts is a problem I will call ‘the vehicle indeterminacy problem’. As I argue below, indeterminacy arises in failures to determine what properties of vehicles are significant and how they can be re-identified. Perhaps we can localize a representation to some ensemble, as in the case of memory engrams, yet are not sure what properties of the ensemble are engram properties. Or perhaps the particular configuration of active and inactive cells in a neuronal ensemble is considered significant for that ensemble to perform its representational role, but we lack a theory about how that pattern of activity should be re-identified.

Re-identifying the pattern requires showing two token patterns to be the same pattern. Consider the identification of neuronal ensembles – populations of neurons thought to compose a

functional unit – with green fluorescent protein (GFP). The neuronal ensemble is a pattern of neural activity that is a good candidate for serving as a representational vehicle: neuronal ensembles can be activated to cause behavior (e.g., Liu et al. 2014), are thought to be capable of serving as computation units (Carrillo-Reid & Yuste 2020), and are often the target of inquiry involving optogenetics, chemogenetics and others. Researchers use GFP to identify neuronal ensembles by genetically modifying experimental animals so that active neurons express GFP and so, as the name suggests, shine fluorescent green when observed under a particular lighting. With such a tool, a neuronal ensemble for a particular behavior, memory, percept or other function is identified and putatively re-identified over time and in other brains. The patterns of neural activity that are supposedly re-identified are considered to be the same because they are similar in certain respects.

But on a cursory assessment, it is entirely unclear by what measure patterns should be determined as similar in the case of vehicles as concrete particulars. Is the relevant vehicle the spatial configuration of active neurons? The synaptic connections between active and inactive neurons? The downstream projections of the neurons? Not only is the referent of ‘pattern’ vague, but whether a token pattern is the same as another token pattern is a vacuous problem when one is unsure by what criterion to measure similarity. One possibility is that we should look for the same number of neurons, each positioned within the same structure with exactly the same connections. But this option is far too stringent given that neuronal connectivity varies drastically from brain to brain, and even within a brain from moment to moment. The problem exists for single neuron patterns as well – is the vehicle the neuron’s activity, dendritic growth, genetic material, or some other thing? In any of these cases, the relevant details by which to determine

similarity are unknown. This problem is compounded when neuronal vehicles are identified using other tools, as cross tool comparison will also rely on some similarity measure.

The problem exists with temporal particulars as well, such as with temporal patterns of neural activity. Consider spike trains, which are sequences of spikes – i.e., action potentials – produced by a neuron. Spike trains have long been used to show that a neuron is ‘tuned to’ or correlated with the presentation of some stimuli, such as the orientation of a bar (Hubel & Wiesel 1959) or the shape, texture or color of various objects (Tanaka et al. 1991). The idea is that if the neuron fires more with the presentation of a particular stimulus, then it seems to track that stimulus or perhaps carry information about it. Indeterminacy already exists in identifying the particular vehicle of interest: is the vehicle the individual spikes, spike trains, or firing rate compared to baseline (Facchin 2024)? What the vehicle is here makes a difference for all sorts of computational analyses that rely on spiking and spike rates. Once a vehicle is defined in any particular case, issues arise when re-identifying this temporal vehicle. For example, if the vehicle is individual spikes, should they occur at the exact temporal location relative to a stimulus, and if so, at what temporal scale? If the vehicle is a spike train, does it matter how many spikes are in the train and the particular shape of the train? What range for firing rate would be considered similar, and at what threshold is firing rate compared to baseline significant?

One might suggest at this point that any vagueness lies in the criteria for identifying vehicles. Once we have a theory that sets determinate criteria for vehicle individuation and re-identification, then these issues will be resolved. Determining what putative functional or computational units in the brain are representations is a major project of neuroscientific inquiry. Consider the numerous debates that persist over vehicle identification: debates about whether representations are local (single neurons) or distributed (populations of neurons), whether

representations are densely distributed over smaller populations or sparsely distributed throughout the brain, as well as debates about whether representations encode temporally (i.e., via the precise timing of spikes) or by firing rate (i.e., via the average spikes per unit of time). The existence of these debates not only demonstrates that vehicle indeterminacy is intolerable but that there are numerous, incompatible criteria available for vehicle individuation. Here, the indeterminacy is a disjunction of possible vehicles in which we are not sure whether representations have vehicle *A*, vehicle *B*, vehicle *C* and so on.

Going further, available criteria for individuating any particular unit remains an open empirical question. For any particular means of probing the brain in search of representations, the bounds of representational vehicles are inexact – which neurons are part of the relevant population? At what threshold does the firing rate of these cells track or carry information about some stimulus? Here, the indeterminacy is a blurriness about boundary setting. Such boundaries are set by convention (e.g., significance thresholds, standard anatomical templates for brain comparison, etc.), by the scale resolution of a particular tool (e.g., stimulus epochs at the millisecond scale, microstimulation at a millimeter scale, etc.) or by stipulation (e.g., within the timeframe of the experimental behavior, within a particular brain region, etc.). While some vagueness exists for these criteria, the problem is the lack of a principled way to formulate criteria that render the vehicles more exact, not the vagueness itself. Existing means of rendering vehicles exact and determinate are circumstantial at present. As such, not only is there debate about how to identify and individuate vehicles, but there is no account for how they might be re-identified when they are identified and individuated. The burden is on vehicle realists to present theories of representations that show how vehicles can be shown to be robust and, therefore, real.

2. Alternatives to Robustness Analysis

Perhaps re-identification is not on the right track and other criteria might establish that an entity is real. In this section, I further motivate appeals to re-identification by considering other methods for positing real entities. For instance, another way to ensure that the vehicle is real is to develop good tools for identifying real patterns. Such a tool should be effective at detecting or extracting real patterns precisely and within their functional role. Then, if that tool is used to identify patterns of neural activity, we have independent reasons for thinking that the identified patterns are real. But developing such a tool leads to an experimenter's regress (Collins 1985). We want to find out if the pattern we have identified is really a representational vehicle, so we develop a tool for identifying real vehicles. To ensure the tool really does this, we need to ensure that the patterns it identifies are real vehicles. Whether the pattern is a real vehicle, however, is determined by the efficacy of the tool used to identify the pattern, and so on. An independent criterion is needed to break into the circle. The proposed criterion is to show that the patterns that this tool picks out are those that are also picked out by other, reliable means of vehicle identification. However, this circles us back to the problem of determining cases in which the patterns are the same (or similar enough).

Another way of framing this circularity is in terms of measurement rather than tools. Consider the problem of nomic measurement (Chang 2004). Vehicles are not directly observable as such. (This is debated, see Thomson & Piccinini 2018; Facchin 2024; Drayson 2025. For the sake of argument, I assume the position here. If vehicles are directly observable, then we face re-identifiability issues directly in the ways discussed above.) So, a measurement of vehicles is needed via some proxy, or some other quality that can serve as a proxy. To identify the right quality, we need a theory that this quality is related to vehicles in a law-like fashion. But such a law between the vehicle and quality cannot be empirically tested, as testing the law would

require already knowing the vehicle (and the quality). On Chang's account, researchers bootstrap their way into a self-consistent law-like relation. Neuroscientists might do this by showing that the pattern being identified has a law-like relation to corresponding behavior. But because we want to test this theory with real vehicles, we must already know that the vehicles are real lest we bootstrap our theories to artifactual vehicles. Here, again, our best bet is re-identification of those vehicles.

One final method to establish vehicles as real entities is to demonstrate that manipulation of those vehicles directly leads to observable effects (Hacking 1983; Cartwright 1983; Nanay 2022). The idea is that, the more evidence we have that some vehicle has causal powers, the more reason we have to think that the vehicle is real. The problem, as I have been emphasizing, is determining *what* has these causal powers. While fascinating neuroscientific innovations allow researchers to manipulate various structures and processes in the brain to observe the effects, such interventions remain coarse-grained and inexact. This means the causal variables are themselves underdetermined.

3. Indeterminacy in Neuroscientific Practice

In neuroscientific practice, patterns are identified as vehicles partially by identification of those patterns' functional profiles. That is, patterns are recorded, activated, inactivated or otherwise manipulated to analyze their relation to behavior in some experimental paradigm, thereby providing information about the pattern's functional role in the system. Such analyses are a way to get precise about the particular unit performing the relevant function across contexts and with various tools given brain complexity, neural dynamics and remapping. In this way, the pattern is whatever activity (e.g., spikes, synaptic growth, etc.) covaries with the experimental stimulus variable and is involved in producing experimental behavior. Similar patterns at similar locations

in the brain that have the same functional profile are considered to be the same kind of vehicle, whether identified at various times or in different brains in the same experimental context or identified across experiments and labs. (Shea (2024) briefly makes a similar point that function partly determines vehicle individuation.) This is evident in any neuroscientific experiment that gathers evidence about neural representations from multiple animals. Finding a neuron in the visual cortex that responds to Jennifer Aniston's face is suggestive of single neurons as vehicles (see Quiroga et al. 2005). We have better evidence if that neuron has the same response profile the next day. Even better evidence is had if a neuron in the visual cortex of another animal has that same response profile. Importantly, then, re-identification is achieved not only via functional profiles but also via measures of similarity in some respect, whether that be their structure, composition or location in the brain. The point is this: functional individuation remains a key tool for vehicle identification and re-identification.

The strategy is to look in a particular brain region for a pattern that has a similar structure and functional profile as another pattern identified in that brain region at another time or in another brain. But what should we make of the fact that function individuation relies heavily on similarity in those other respects? Ideally – that is, with the right tools at our disposal – the boundaries of an identified pattern can be sharpened by investigating the components and processes necessary and sufficient for that pattern to perform its functional role. Our current tools lack the precision needed to do this with neural entities, and there is good reason to be suspicious that we could ever control for all other variables in the brain to establish sufficiency for any component part. As for spiking, our current theories lack the specification needed to set such bounds as it is unclear whether the rate or timing of spiking matters, and in either case, at what threshold meaningful activity occurs across contexts. And as we saw previously, it remains

underdetermined which patterns are being re-identified in either case, if they are, and whether those patterns are, indeed, similar in the relevant sense. Re-identification is what matters for establishing that the vehicles are robust.

While both vehicle and content indeterminacies are partly due to a disjunction of possible vehicles or contents, vehicle indeterminacy is contrasted with content indeterminacy by the emphasis on re-identification. Re-identification challenges are what makes vehicle indeterminacy intractable, as both biological complexity and neural dynamics suggest identified patterns are likely to change over time. (This does not necessarily imply, however, that the vehicle changes, especially since the relevant properties have not been determined. Furthermore, changes in vehicles do not imply that the content changes, see Robins 2020.) In contrast, re-identification of content resolves *some* content indeterminacy issues. Consider a case in which neuroscientists identify a vehicle in the toad's retina for fly representations that does not respond to small moving black dots. In this case, <fly> content is being re-identified by researchers but not <small moving black dot> content. As for vehicles, the only way to determine that they are vehicles is by their function and similarity to other vehicles of the same type. Similarity criteria, while straightforward enough for content, are challenged by blurry vehicle boundaries that result from complexity and dynamics. In the next section, I investigate whether typing might resolve this issue.

4. *Vehicle Typing*

Given a lack of alternatives to re-identification, how can re-identification practices be improved to resolve vehicle indeterminacy issues? Here, we might take a cue from neuroscientists – the experts at vehicle identification – and consider vehicle typing. Looking at the ways that neuroscientists abstract from the particularities of individual vehicles might tell us how such

vehicles are typified and made generalizable or projectable. Coarse-graining practices in neuroscience – those that abstract from complexity and detail – type vehicles by compressing, abstracting or suppressing variations in patterns. Coarse-graining allows researchers to identify activity patterns covarying with stimulus or behavior onset that occur in the same general area or have the same general shape. Examples for anatomical localizers include smoothing spike trains, getting spiking averages from a microelectrode array or identifying an active cluster of neurons in a particular brain region. For functional localizers, an example of coarse graining is taking the average of an experimental group. These coarse-graining practices are required even for re-identification in the same paradigm, with the same probe, in the same animal – that is, such coarse-graining is needed to establish the pattern even over time in the same experimental context. Coarse-graining is also required for the robustness achieved by getting the same results using different techniques. So long as this coarse-graining works to re-identify patterns that perform some function in different experimental contexts, the practice establishes vehicle types. In this section, I offer an account of vehicle determination in neuroscience that draws on Shea's (2007) desiderata for typing vehicles.

Shea (2007) offers four desiderata for typing representational vehicles in order to determine the vehicles of connectionist systems. His account is focused on representations in artificial connectionist systems, which are systems that putatively represent via associations between units in a network. At least some of the vehicles identified in neuroscience may be treated as parts in a connectionist system, such as individual neurons or neuronal ensembles. But the desiderata should be expanded to encompass the many kinds of vehicles that neuroscientists identify. Shea's desiderata for typing vehicles in connectionist systems is as follows: typing should

- (i) capture some underlying property of the network's mechanism of operation by which it performs its task;
- (ii) abstract away from individual weight matrices and particular patterns of activation;
- (iii) be such that it may be shared by different networks trained on the same task; and
- (iv) form part of an explanation of the network's ability to project its correct performance to new samples outside the training set. (Shea 2007, 250)

The first desideratum seems unobjectionable for vehicles generally. The vehicle type should maintain the properties by which the pattern performs its role as a vehicle. One might do this by suppressing patterns caused by irrelevant factors to get a description of phenomenon from data (Wright 2018). Whichever factors are irrelevant will be determined by a particular theory on vehicles and the particular task. The second desideratum is a necessary aspect of typing: abstraction. The particularities of individual patterns should be abstracted from to remove those details that do not matter for serving as a vehicle. Consider the many vehicles for the word 'cat' here: . Much variation between these patterns is permitted while still serving as the same vehicle kind. The third desideratum is, generally speaking, that the vehicle exists in different systems, especially those systems that perform some task that makes use of that vehicle type. This desideratum lends to re-identifiability. Finally, the vehicle should be explanatory. For general purposes, explanatory vehicles contribute to making sense of the system's ability to perceive, cognize, remember or behave. While I think the tie to explanation is pragmatic, I do not intend to beg the question against vehicle realists by including it here. This desideratum can be understood as the kind of sense-making that is required to know that the system was approximated as it really is.

Here are the general desiderata I propose for typing patterns of neural activity in neuroscience: typing should

- (1) Identify some property of the mechanism by which it performs a functional role in an experimental task;
- (2) Abstract away from particular patterns of activation and the morphological details of individual units;
- (3) Permit re-identification of the vehicles in other systems that perform the same functional role; and
- (4) Form part of an explanation of the system's ability to perform that task.

Because I formulate these desiderata for patterns of neural activity generally, theoretical disagreements will remain about when such typed patterns are representational vehicles. That is, many theories will require more from these desiderata for the identification of representational vehicles in the brain. For example, a common view is that the experimental task used to identify representations should be 'representation-hungry' (Clark & Toribio 1994), meaning that some experimental behavior was performed that was sensitive to entities or features not available to the organism, such as those that are absent or even non-existent or abstract. Or one might require that the representation be manipulated to demonstrate its causal powers (Nanay 2022). One can modify the first desideratum to accommodate these views. Another view requires that some computation be performed in cases of representation. In these cases, the functional role in the first and third desiderata can specify that the vehicle performs a particular computation. I am open to such modifications to accommodate particular theories of neural representations. Here, these desiderata are broad as their purpose is only to provide a means to coarse-grain vehicles such that they are re-identifiable over time, in different brains and across labs.

The hope is that the coarse-grained pattern retains the significant features of the activity pattern that can be re-identified across contexts by matching token patterns to the type. Importantly, however, the second desideratum might very well permit abstraction away from real vehicles in the brain. That is, coarse-graining techniques might gloss over numerous disparate vehicles for representations, inappropriately lump patterns when it would have been better to split them, or, critically, remove or ignore morphological details that are crucial for serving as a representation in the system. To improve typing practices, token patterns should share characteristics that are weighty and relevant for consideration as a representational vehicle. These characteristics cannot be determined a priori and should not be determined arbitrarily. The project of neuroscientific practice and theorizing is to determine which characteristics are important, which is ongoing. This is important, as even if a representation-hungry task is used, causal powers are demonstrated or a computational task is performed, vehicle indeterminacy remains.

The concern here is that patterns of neural activity that are of interest to researchers or that are accessible with current neuroscientific tools may not be the same patterns that the brain uses to carry information about the world. This is fine as far as neuroscientific practice goes. Sometimes it is interesting enough that one can reactivate a population of neurons to cause a behavior, even if that population is engineered by neuroscientific tools. If a similar population of neurons meets all the desiderata above, neuroscientists will have then discovered important properties of brain function, such as that populations of neurons are possible representational units. And in many cases, the third desideratum will vindicate decisions made in meeting the second desideratum. But it may very well be the case that non-representational patterns are being generalized, which means that the vehicle indeterminacy problem puts generalizations of experimental results on

shaky epistemic grounds. But this is just the status quo in neuroscience and a place where theorizing plays a key role.

5. Conclusion

I have worked through possible strategies for identifying real vehicles, all of which rely on re-identifying those vehicles over time, with other tools, and across contexts to varying degrees.

The best strategy in neuroscience involves coarse-graining vehicles to type them, making those vehicles re-identifiable. Coarse-graining does not, however, ensure that the patterns are those naturally occurring in the brain. Theorizing about neural representations will determine legitimate ways to coarse-grain those natural patterns. But while the concerns raised here are status quo in neuroscience, they remain a serious obstacle for vehicle realism to succeed.

Because coarse-graining is how re-identification is achieved, and coarse-graining may be done on non-representational patterns or involve manufacturing of patterns by researchers, the strategy is insufficient for establishing vehicle realism. I conclude that this coarse-graining strategy for re-identification vindicates neuroscientific practice but not vehicle realism.

Vehicle realists need a response to the vehicle indeterminacy problem, which means they need an account of vehicle individuation and reidentification. Neuroscientific theories will specify the properties of vehicles that are important for performing particular tasks, such as producing action potentials at a particular rate in correlation with some stimulus or performing a causal role in producing behavior. Insofar as such properties can be re-identified performing their functional role over time, in other brains, across labs, and with other tools, there is good reason to trust such theories. Because vehicles are re-identified according to theoretical parameters – i.e., the vehicle's functional profile and properties thought to be relevant for carrying information or content – vehicle realists must show that the individual theories and explanations that lend to re-

identification are sufficient for showing that the vehicles they postulate are real beyond the mere fact that they have been re-identified. Re-identification is always grounded in similarity, and as Goodman (1972) warned, similarity is always in a certain respect, according to certain characteristics or relative to some theory in order to make comparisons according to important properties.¹

¹ **Acknowledgements** I am grateful to Frances Egan, Jordan Dopkins, David Barack, Adina Roskies, and Edouard Machery for many helpful comments on this work as well as to Michael De Vivo, Sarah Robins, Andrew Rubner, and David Chalmers for helpful discussion. This work moreover benefitted from audience feedback at the 50th Annual Meeting of the Society for Philosophy and Psychology and the British Society for the Philosophy of Science 2025 Annual Conference.

Citations

- Bechtel, William, and Jennifer Mundale. 1999. "Multiple Realizability Revisited: Linking Cognitive and Neural States." *Philosophy of Science* 66 (2): 175–207.
- Burnston, Daniel. 2020. "Contents, Vehicles, and Complex Data Analysis in Neuroscience." *Synthese* 199 (1): 1617–1639.
- Calcott, Brett. 2011. "Wimsatt and the robustness family: Review of Wimsatt's Re-engineering Philosophy for Limited Beings." *Biological Philosophy* 26: 281–293.
- Cao, Rosa. 2022a. "Putting Representations to Use." *Synthese* 200 (151): 1–24.
- Carrillo-Reid, Luis, and Rafael Yuste. 2020. "What is a Neuronal Ensemble?" *Oxford Research Encyclopedia of Neuroscience*. Oxford University Press.
- Chang, Hasok. 2004. *Inventing Temperature: Measurement and Scientific Progress*. Oxford University Press.
- Clark, Andy, and Josefa Toribio. 1994. "Doing Without Representing?" *Synthese* 101 (1): 401–431.
- Collins, Harry. 1985. *Changing Order: Replication and Induction in Scientific Practice*. Sage Publications, Ltd.
- Drayson, Zoe. (forthcoming) "Representations are (still) theoretical posits." *Theoria: An International Journal for Theory, History and Foundations of Science*.
- Egan, Frances. 2020. "A Deflationary Account of Mental Representation." In *What Are Mental Representations?*, ed. Joulia Smortchkova, Krzysztof Dołęga, and Tobias Schlicht, 26–53. Oxford: Oxford University Press.
- Egan, Frances. 2025. *Deflating Mental Representation*. The MIT Press.
- Facchin, Marco. 2024. "Maps, Simulations, Spaces and Dynamics: On Distinguishing Types of Structural Representations." *Erkenntnis*
- Facchin, Marco. 2024. "Neural representations unobserved—or: a dilemma for the cognitive neuroscience revolution." *Synthese* 203 (7): 1–42.
- Goodman, Nelson. 1972. "Seven Strictures on Similarity." In *Problems and Projects*, 437–447. The Bobbs-Merrill Company, Inc.
- Hubel, David, and Torsten Wiesel. 1959. "Receptive Fields of Single Neurones in the Cat's Striate Cortex." *Journal of Physiology* 148 (3): 574–591.
- Liu, Xu, Steve Ramirez, Roger Redondo, and Susumu Tonegawa. 2014. "Identification and Manipulation of Memory Engram Cells." *Cold Spring Harbor Symposia on Quantitative Biology* 79 (1): 59–65.
- Nanay, Bence. 2022. "Entity Realism About Mental Representations." *Erkenntnis* 87: 75–91.

- Quiroga, Rodrigo, Leila Reddy, Gabriel Kreiman, Christof Koch, and Itzhak Fried. 2005. "Invariant Visual Representation by Single Neurons in the Human Brain." *Nature* 435 (7045): 1102–1107.
- Robins, Sarah. "Stable Engrams and Neural Dynamics." *Philosophy of Science* 87 (5): 1130–1139. doi:10.1086/710624
- Shea, Nicholas. 2007. "Content and Its Vehicles in Connectionist Systems." *Mind & Language* 22 (3): 246–269.
- Shea, Nicholas. 2018. *Representation in Cognitive Science*. Oxford University Press.
- Shea, Nicholas. 2024. *Concepts at the Interface*. Oxford University Press.
- Tanaka, K., H. Saito, Y. Fukada, M. Moriya. 1991. "Coding visual images of objects in the inferotemporal cortex of the macaque monkey." *J Neurophysiology* 66 (1): 170–189. doi: 10.1152/jn.1991.66.1.170. PMID: 1919665.
- Thomson, Eric, and Gualtiero Piccinini. 2018. "Neural Representations Observed." *Minds & Machines* 28 (1): 191–235.
- Wimsatt, William. 2007. *Re-engineering Philosophy for Limited Beings: Piecewise Approximations to Reality*. Harvard University Press.
- Wright, Jesse. 2018. "The Analysis of Data and the Evidential Scope of Neuroimaging Results." *British Journal for the Philosophy of Science* 69: 1179–1203.