



Tracers in neuroscience: Causation, constraints, and connectivity

Lauren N. Ross¹ 

Received: 30 January 2020 / Accepted: 21 November 2020
© Springer Nature B.V. 2021

Abstract

This paper examines tracer techniques in neuroscience, which are used to identify neural connections in the brain and nervous system. These connections capture a type of “structural connectivity” that is expected to inform our understanding of the functional nature of these tissues (Sporns in *Scholarpedia*, 2007). This is due to the fact that neural connectivity constrains the flow of signal propagation, which is a type of causal process in neurons. This work explores how tracers are used to identify causal information, what standards they are expected to meet, the forms of causal information they provide, and how an analysis of these techniques contributes to the philosophical literature, in particular, the literature on mark transmission and mechanistic accounts of causation.

Keywords Causation · Philosophy of neuroscience · Causal reasoning · Explanation

1 Introduction

In efforts to better understand the functioning of the human brain, many projects in neuroscience examine the causal process of signal propagation along neurons. In studying this causal process, a large amount of research investigates anatomical neural connections, which constrain and track the flow of these signals (Sporns 2007). These anatomical connections capture a form of “structural connectivity” that is thought to provide us with information about function, as structure informs function.¹ In fact, this is a significant motivation behind the human connectome project, which aims to map all neurons and neuronal connections in the human brain (Lichtman and Sanes

¹ For these structure to function claims, see Frick et al. (2013), Saleeba et al. (2019) and Sporns and Bassett (2018). Also, note that structural connectivity differs from both functional and effective connectivity (Friston 2011).

✉ Lauren N. Ross
rossl@uci.edu

¹ Department of Logic and Philosophy of Science, University of California, Irvine, 3151 Social Science Plaza A, Irvine 92697-5100, USA

2008; Sporns 2012). In some sense, this project is similar to how we might study an electronic device. If we are interested in how this device works we might start by identifying the circuit along which the electricity flows. Similarly, we hope that identifying the “circuit architecture” of the brain, will allow us to better understand how it “works” and how it enables higher-level cognitive processes.²

One set of methods that have been used to study structural connectivity of neural systems are tracer techniques. Many of these techniques exploit the fact that cellular materials are transported along a neuron’s cell body and even transynaptically, across neurons that are connected in series. When these transported materials are marked with an identifiable tag, scientists can follow the tag as it flows along the causal process, “discovering” steps of the process that may have been unknown or poorly understood. Successful application of these methods has led neuroscientists to view tracers as a “fundamental technique” in this field and “the primary method for visualizing brain networks in all areas of neuroscience” (Levy et al. 2015, p. 57). This paper examines tracer techniques in neuroscience with respect to (1) the standards they are expected to meet, (2) their role in identifying causal information, and (3) the particular types of causal information that they provide. In addressing these points, I clarify how an analysis of neuroscientific tracer techniques contributes to the philosophical literature, in particular, the literature on mark transmission and mechanistic accounts of causation.

2 Background

Neural connections are often described as “circuits” of causal pathways, which guide the flow of signals and information. This signaling is a type of causal process, which propagates along individual nerve cells and chains of neurons in series. In these cases, an upstream signal causes a downstream signal and so on, in a sequence that flows along the pathway. Whether a signal is present at an upstream location or not “makes a difference” to whether it will occur at a downstream location. This basic causal structure is accommodated by a minimal interventionist framework (Woodward 2003). Within this framework, to say that X is a cause of Y means that an ideal intervention (in background circumstances B) that changes the values of X, produces a change in the values of Y. In this case, changing the presence or absence of an upstream signal, causes (or is one causally relevant factor for) the presence or absence of a downstream signal. Given three locations along an individual neuron, in sequential order (A to B to C), this captures a basic difference-making relationship among them—the presence or absence of a signal at upstream location A makes a difference to its later presence at downstream location B, and so on. The same can be said of signal flow across three neurons connected in series (A to B to C), or more complicated circuits of connections.

While neural signaling is viewed as straightforwardly causal, it has some unique features that are not present in other types of causal systems. First, the causal process of neural signaling involves the unique feature of having a (i) fixed, physical structure—this is the nerve cell and its cell membrane along which signals and information flow.

² As Sporns states “Neural activity, and by extension neural codes, are constrained by connectivity. Brain connectivity is thus crucial to elucidating how neurons and neural networks process information” (Sporns 2007).

Even when signals are not being transmitted down neurons, this physical structure is still present and it can be studied to get information about the pathways along which signals move. This is similar to how we might study roadways in order to get information about the flow of traffic through a city. Alternatively, other causal processes lack such a fixed, physical structure. These include DNA replication, signal transduction, and drug receptor interactions. In these cases, causal parts interact, but there is no large-scale physical barrier that guides these interactions. In the case of neurons, this relatively fixed, physical structure serves as a constraint that guides the flow of neural signals and information.

This fixed, physical structure has been studied with various techniques in order to better understand signal propagation in neurons. One of these techniques is tract tracing in gross anatomical dissection. This involves following nerve tracts through animal models or human cadavers in order to identify which areas of the body they connect up. This work originated in the seventeenth century with Steno's "teasing methods," which isolated nerves by distinguishing white matter from grey matter (Heimer 2005). While these methods have led to various advances, they also face a number of challenges. A first challenge is that nerve tracts can be exceedingly difficult to identify and distinguish from other structures in the body, such as fascia and connective tissue.³ Second, even when these tracts are accurately identified, they are a larger structure that is comprised of many single neurons. If we are interested in single neurons and their connections, we will need other tools to identify these smaller-scale structures. Further challenges include the fact that these single neurons are often highly tangled and that gross anatomical dissection does not reliably indicate the directionality of signal flow.⁴

In order to overcome these challenges, neuroscientists have developed further techniques that exploit a second unique feature of these causal systems. This second feature is that neural systems involve the (ii) flow of material down each step of the causal process—namely, down each neural cell and along chains of neurons connected in sequence. Weiss and Hiscoe's (1948) identification of axonal flow led to the discovery that neurons contain cellular materials that are shuttled along their axons and from one neuron to another, through axonal and transynaptic processes, respectively. Kristenson (1970) built on this by showing that cellular materials moving via axonal transport could be tagged with horseradish peroxidase (HRP), in order to provide "direct observation" of this causal process. This led to the development of tracer and tagging techniques, which involve tagging these and other materials, in order to watch their flow along single neurons and neurons connected in series (Martin and Dolivo 1983). By following an visualizable tracer, scientists can identify the projections, length, and directionality of material flow along the neuron and along chains of neurons connected in sequence. Tracing this flow of material is important, because it provides information about the nerve membrane, which constrains and guides the flow of electrical signals. As the flow of material and electrical signals are both constrained by the same

³ This is especially the case when compared to identifying blood vessels and other more distinct anatomical structures.

⁴ Additional studies would need to be performed (likely in living specimens), to determine the direction of signal flow.

cell membrane, studying one of these processes (material flow) provides information about the other (signal flow).

The tracer experiments described above exploit transport processes in living neurons in order to study their connectivity. While these techniques “continue to be the most reliable way of inferring axonal connection in mammalian brains,” they are just one type of tool that neuroscientists use to trace these connections (Shi and Toga 2017). Other neuroanatomical tract tracing tools include: (1) gross anatomical dissection, (2) neural degeneration, (3) diffusion in fixed and post-mortem tissue, and (4) in vivo diffusion MRI (Lanciego and Wouterlood 2011). While these techniques rely on different types of information, they are all used to study neural connectivity and they contribute to our understanding of brain and nervous tissue functioning.

This analysis focuses on neuroanatomical tracers that target axonal transport and transynaptic processes in living neurons. In order to best understand how these tracers work we should examine how they are used and what types of causal information they provide. I examine these points in the next two sections and explore the implications of this analysis for the philosophical literature on causation.

3 Neuroanatomical tracers: How and when are they used?

Understanding how and when tracers are used requires a brief description of what they are used for (this will be examined in more detail later). In the context of neuroscience, tracers are used to identify neurons and their connections—in particular, where a neuron is spatially located, what its most immediate connections are, and where it resides in a larger network of connections. In this and other scientific contexts, tracers are used to get information about the basic steps of a causal process—they trace the process without providing detailed information about how it works. How exactly are scientific tracers used to study causal systems? How do they differ from other causal investigative strategies and what standards are they expected to meet?

3.1 Tracers: scientific standards

A first important feature of scientific tracers is that they need to trace something that is reliably moving through a causal process. What exactly this something is depends on the causal process of interest. Consider causal processes that exhibit *material continuity* in the sense that they involve some material that reliably flows through all steps of the causal process. Neurons fall into this category because they involve cellular materials that are shuttled along individual neurons and across neurons connected in sequence. Neurons contain an axonal transport system with small molecular motors that shuttle materials along a microtubule cytoskeleton, positioned along the length of the neuron. These molecular motors run *anterograde* from the cell body to the axon terminal and *retrograde* from the axon terminal to the cell body. Scientists explicitly analogize these transport processes to a freeway, stating that they are a “rail like” system, in which “cargo [is] delivered” via “transport mechanisms” along “highways within the axon” (Morecraft et al. 2014, p. 369).

This material continuity feature of neural transport processes has a number of advantages—in particular, it can be exploited to study neural projections and neural connectivity. If this material is tagged with a visualizable tracer, the tracer can be followed to reveal a single neuron’s spatial location and its downstream or upstream connections. These tracers are viewed as “pathfinders in the nervous system” as they chart individual neural projections and sequences of neural connections along which this material flows (Ekstrand et al. 2008). In the context of neuroscience, these physical tracers often include dyes, radioactive substances, and even viruses (Morecraft et al. 2014). These tracers are distinguished on the basis of whether they flow in an antero-grade or retrograde direction, and depending on whether they move transynaptically or not. A transynaptic tracer is capable of moving past the synapse and into downstream neurons connected in series. Interestingly, this basic tracer methodology is not just found in neuroscience, as similar methods exist in various subfields of biology (Ross Forthcoming). For example, biochemists attach radioactive tracers to metabolites in order to identify the steps of metabolic pathways. Similarly, ecologists attach radioactive markers to caloric energy in order to identify the sequence of species in a food chain. In these examples, tracers are attached to various materials that outline the steps of a causal process.

Stating that a tracer needs to “tag” something that reliably moves through a causal process is a fairly obvious and underspecified requirement. There is much more to say about how this actually works. I am going to suggest that, in order for a tracer to provide reliable information about a causal system, it needs to meet particular criteria. These criteria are often implicitly assumed in scientific work, and sometimes explicitly formulated, with significant variation.⁵ I identify three of these criteria, which concern the (1) influence of the tracer on the causal system under study, (2) the identifiability of the tracer, and (3) the ability of the tracer to tag exactly and only the causal process in question.

According to the first criterion, the (1) tracer should not disrupt the normal causal process under study. The tracer should provide us with information about how a system normally functions in the world, as opposed to a system that is perturbed with the introduction of a foreign element. This criterion is used to rule out toxic tracers, which can damage the nerve membrane and then extravasate into extracellular regions. (This relates to the specificity of the tracer, which will be discussed soon.) For these reasons, it is claimed that an “optimal” tracer should be “relatively harmless” and have “minimal neurotoxicity” (Hu et al. 2013; Huh et al. 2010; Mi et al. 2019).

A second criterion ensures that the tracer is sufficiently identifiable so that it can be successfully tracked throughout the causal process. This involves ensuring that the (2a) tracer is sufficiently unique so that it is not confused with nearby entities in the system and that the (2b) tracer is present in large enough concentration so that it does not dilute to the point of being unidentifiable in later steps of the causal process. In the context of neuroscience, viral tracers are well suited to meet (2b) as their ability to amplify and self-replicate increases tracer concentration.⁶ This allows for “intense

⁵ For helpful discussions of “ideal” tracers and the criteria they should meet see Nassi et al. (2015), Hu et al. (2013), Huh et al. (2010) and Mi et al. (2019).

⁶ The self-replicating feature of viral tracers may seem to conflict with material continuity. While the particular material associated with the original, introduced viruses is unlikely to reliably flow across numerous

labeling” and “minimal fading,” which both support easier tracer identification (Hu et al. 2013; Nassi et al. 2015).

A third criterion pertains to the tracer’s ability to tag exactly and only the causal process in question. According to this standard (3a) the tracer should bind to material that reliably moves through the causal process and it should (3b) bind to it firmly without falling off. There are many ways that a tracer can fail to meet these standards. A tracer might firmly bind to some material, but the portion it binds to could get spliced off at some point during the causal process.⁷ Alternatively, even if material is reliably moving through a causal process, a tracer can fail by unsuccessfully binding to it.⁸ Tracers will not provide reliable information if they detach from their target or if they indiscriminately mark other entities in close proximity, which are causally irrelevant to the system. The ability of a tracer to meet these requirements and tag exactly and only those neurons that are synaptically connected is often referred to as “specificity” and the “specificity of tracing” (Nassi et al. 2015).⁹ In this sense, a specific tracer “propagates exclusively between connected neurons...allowing for the stepwise identification of neuronal connections of progressively higher order” (Ugolini 2011). It is important that the tracer spread only through synaptic connections, because information flow is restricted to these connections, and tracers are intended to reveal these routes of information flow (Nassi et al. 2015). If a tracer were to indiscriminately tag nearby neurons by diffusing through the extracellular space, this would no longer provide useful information about signal flow. As Callaway states:

One of the most important characteristics to consider for any transneuronal tracer, including neurotropic viruses, is whether spread is restricted to neurons that are connected by synaptic contacts. Because the transfer of information between neurons is primarily dependent on synapses (and gap junctions), the most relevant circuit diagram is based exclusively on connectivity, not proximity. (Callaway 2008, p. 617)

The meaning and importance of having “specific” tracers is further clarified by Martin and Dolivo who state that “an ideal tracer [should be] selective and...stay in the considered pathways” (Martin and Dolivo 1983, p. 268).

These three tracer requirements are articulated with a focus on causal systems with material continuity. However, not all causal systems have this feature, as some lack

Footnote 6 continued

neural connections, what matters is that it does flow across the synaptic junction (from one neuron to another). Material continuity is preserved at each causal link. A related point is discussed further in footnote 10.

⁷ This would occur, for example, if a radioactive molecule was attached to a phosphate group, which is spliced off a metabolite as it makes its way down a biochemical pathway. This is similar to attaching a bug to the coat of a spy, in order to trace them, while realizing at some later point in time, that they have removed their coat. Although the bug is still attached to the original material, this no longer follows the causal process of interest.

⁸ For example, if the spy keeps their coat on but the bug falls off, tracing efforts will again fail.

⁹ One exception to this third criterion is when the tracer reliably flows with material moving along a causal process, without directly tagging it. Examples of this are viral tracers in neural pathways and radioactive tracers in blood vessels and the gastrointestinal tract. Although some of these tracers do not directly tag material moving along the causal pathway, we have good reason to believe that they flow along with it.

the flow of material from one causal step to the next. A biological example of this is a hormone cascade. In this case, a molecule of some hormone binds to an extracellular receptor—this causes the intracellular portion of the receptor to change conformation, which activates an intracellular enzyme, which results in a downstream cascade of intracellular changes. Scientists view this process as legitimately causal despite the fact there is no material that moves from the initial hormone through the downstream sequence of steps.¹⁰ An ordinary life example of this is a sequence of dominos that fall over in succession. In this case, toppling the first domino clearly causes the rest to fall over in series, but again, there is not any material that reliably moves from the first domino down to the tenth or twentieth domino. Kinetic energy or momentum might flow through this causal system, but material does not. Instead of using a physical tracer that tags material, an alternative way to study this system might involve tagging something that does reliably flow through it, such as kinetic energy or momentum.¹¹ It is worth considering whether the tracer criteria above could be rephrased with respect to tracing these other entities, in order to better accommodate these causal systems.

At this point, we can draw four main lessons from this analysis of tracer methodology. First, there are important differences across types of causal systems that matter for tracer experiments. Some causal relationships involve the flow of material, while others do not. This difference matters for the types of tracers that can (and cannot) be used to study a causal system. Tracers that tag material are best used for systems with material continuity, as opposed to systems that lack this feature. Any account of causation that is used to understand tracer methodology should capture these differences across causal systems and how they matter for tracer selection and use.

Second, developing and using tracers requires already knowing something about the causal process of interest. For example, you need to know what property moves through the causal system in order to know what you should tag and what will successfully mark this property. In some sense, this is apparent in the three criteria for tracers. How are you going to know whether the tracer meets the specified requirements unless you know what it is supposed to trace and whether it succeeds in tracing this? This reveals a clear difference between tracers and causal discovery methods involving statistical analysis of observational data and interventionist-type experiments. With these latter methods the main question is *whether* a causal relationship exists between various properties or not. When scientists use tracers they have already accepted that

¹⁰ Notice that signal transmission in neurons also involves receptors, in a way that is similar to this hormone example. Pre-synaptic neurons release neurotransmitters, these bind to receptors on postsynaptic neurons, and this binding triggers downstream effects. As material is not reliably moving from the neurotransmitter, to the receptor, to the downstream effects, this causal process appears to lack material continuity. If signal transmission in neurons lacks material continuity, is this a problem for my analysis, which claims that many neuroscientific tracers work by exploiting material continuity? No. The physical tracers discussed in this paper (viruses, dyes, radioactive markers, etc.) are not directly tracing signal transmission in neurons. They are tracing axonal transport and transynaptic processes that reliably move material (lysosomes and viruses, for example) along and across neurons. Tracing (a) axonal transport and (b) transsynaptic processes provides information about (c) signal transmission as all of these processes are constrained by the physical contours of neurons and their sequences of connections. Furthermore, (a) and (b) are far easier to tag and trace than (c), which helps explain why they are targeted with this methods. This is discussed in more detail shortly.

¹¹ For example, signal tracers for electronic products and genetic tracers (that tag information) may count as cases of this.

a process is causal—what they want to know are other features of the process (such as, where it is located, what upstream and downstream entities it connects up, and what intermediates are present along the way). Relatedly, it should be clear that tracer experiments are not Woodwardian interventionist experiments. While interventionist experiments involve manipulating a candidate cause to identify its effects, tracer experiments involve “harmlessly” tagging a candidate cause, so that it (or its effects) can be monitored at a later point in time.¹²

Is it problematic to say that a method of causal study requires knowing something about the causal system before the method is used? In the case of neuroscientific tracers, scientists already know how a type of system generally functions or operates, but they lack other fine-grained information about it. For example, you can know that most neurons transport various materials through axonal transport and synaptic processes, without knowing—for a given neuron or set of neurons—where they are located, how they are connected, and which upstream and downstream elements link up. In these cases, information about the causal system is used to develop a tracer that can answer these questions. We see this in the historical development of tracers—only once axonal transport was discovered were physical tags and tracers used to study neural connectivity. Similarly, only once it was identified that caloric energy flows along ecological pathways, was this knowledge used to develop tracers that illuminated the ordered sequence of species connected in a food chain.

Third, neuroscientific tracers are a unique form of causal investigation because they are used to study causal systems “by proxy,” in which information about one causal system is gained by studying another. For example, although scientists in this area are primarily interested in how *signals* flow along connections of neurons, they tag and study the movement of cellular *materials* along these neurons. Scientists can tag axonal transport and transynaptic processes to learn about signaling pathways because all of these processes flow along the same routes—all are constrained by neuron cell membranes and connections of neurons in sequence. As both of these causal processes are guided by the same physical constraint, studying one process gives you information about the other (Callaway 2008). Why do scientists bother studying a causal system “by proxy”? Why not just “directly” study the particular system they are interested in? Scientists might use this strategy when one causal processes is easier to study than the other. If it is easier to tag and trace the flow of material as opposed to the flow of information, this would explain why neuroscientists study information flow “by proxy.” Of course, this option is only available when both processes are guided by the same constraint, such that gaining information about one reliably provides information about the other. Clarifying exactly how this causal study “by proxy” is justified and the different forms it can take is a rich area for future work (Hooker 2012; Winning and Bechtel 2018; Winning 2018). This work should clarify the justification for this strategy, whether it can be implemented in various ways, and the circumstances under which studying one causal system can reliably provide information about another.

Fourth, tracers appear to be used to study causal systems that involve a fixed sequence of causal steps. These steps are “fixed” or “constrained” in the sense that

¹² For related discussion see Kästner’s distinction between interventions and mere interactions (Kästner 2017, p. 156).

they have a very particular ordering—moving through the causal process requires following this particular ordering without skipping a step or reversing their sequence. At any given step in these processes there are a limited number of “next moves”—a limited number of downstream possibilities that the tagged material could “move to” next. Contrast this with a situation in which a walker is roaming along an open field. If we tag this walker with a tracer, their path is unlikely to provide a similar form of causal information because their walking routes are not very constrained in this area. Once they move to a first location there are a large number of possible downstream locations they could move to next—they could turn around, walk 30° to their right, 45° to their left, and so on. When scientists use tracers to study causal systems they typically study systems with a fixed order of causal steps, in which this order is guided by various constraints. We see this fixed order in metabolic pathways, vascular pathways, developmental pathways, ecological pathways, and many others.¹³ Future work in this area should explore the role of constraints in these causal processes and their relationship to tracer experiments. This work should also examine whether a “constraint” is a particular type of causal factor and, if so, how they should be understood and how they differ from other causes.

3.2 Contributions to existing views: mark transmission and causation

Although this paper has analyzed tracer techniques within an interventionist framework, these techniques are often associated with mark transmission accounts of causation, such as those supported by Reichenbach (1971) and Salmon (1984). Why not use these mark transmission accounts to understand neuroscientific tracers? How does an analysis of these tracer techniques bear on this philosophical literature?

At first glance, neuroscientific tracer experiments appear similar to mark transmission accounts of causation. These accounts have been articulated by Russell (1948), Reichenbach (1971) and Salmon (1984) and they share the view that causation can be defined by the capacity of a process to transmit a mark. In particular, it is suggested that genuine causal processes are capable of transmitting marks, while non-causal processes (or “pseudo processes”) are incapable of this. Although there are different conceptions of what counts as a “mark,” in much of this work various types of properties have counted, such as “constituent material, bonding forces...geometrical shape” (Dowe 2018, p. 202) and “momentum, energy, or electric charge” (Salmon 1997, p. 467). For example, when a minor car crash leaves a dent in a car’s door, this dent is a mark that is transmitted with the car as it continues on its journey through town. On the other hand, the car’s shadow is a mere pseudoprocess because various alterations to the shadow’s shape are not consistently transmitted as it moves with the car. This view is said to accommodate many other ordinary life examples of causal processes, including: scuffs on a fly baseball, snow on the roof of a railcar, carved initials on a flying arrow, and chalk marks on a sequence of colliding billiard balls (Reichenbach

¹³ Note that there are at least two different types of pathways in these cases. There are pathways that involve changes in the physical location of some material over time (neural and vascular pathways) and pathways that involve changes in the constitution of some material over time (metabolic and developmental pathways).

1971; Salmon 1984; Woodward 2016). In all of these cases, some physical mark moves through the causal process in question.

Scientific tracers appear similar to these accounts, because they involve studying causal processes by tracing properties or “marks” that flow through them. In fact, both Reichenbach and Salmon appeal to scientific tracers in motivating their mark transmission views, suggesting that they involve a similar rationale. As Reichenbach states, “radioactive tracers [are] used nowadays to reveal the causal structure of circulation in living organisms” (Reichenbach 1971, p. 200). Relatedly, Salmon claims that “[m]arking methods are sometimes used in practice for the identification of causal processes” and that, more specifically, “[r]adioactive tracers are used in the investigation of physiological processes—for example, to determine the course taken by a particular substance ingested by a subject” (Salmon 1984, p. 154).

How exactly does an analysis of scientific tracers contribute to this philosophical literature? Do these tracer techniques support mark transmission accounts of causation or provide insight into their plausibility? I will outline two ways in which an analysis of neuroscientific tracers contributes to this literature.

First, in contrast to Reichenbach and Salmon’s claims, an analysis of scientific tracers does not support a mark transmission definition of causation. Reichenbach and Salmon are not just arguing that scientists *use* tracers to study aspects or features of causal systems, they are suggesting that causation can be *defined* by tracer transmission and that scientific tracers help illustrate this. They suggest that tracers can be used to discover causal relationships and distinguish them from non-causal ones. This is discussed by Salmon at various points throughout his work:

I believe that the mark method is an extremely useful tool for the *discovery* and study of causal processes. A paradigmatic example is the use of radioactive tracers in studying physiological processes (Salmon 1998, p. 20, emphasis added).

Marking methods are sometimes used in practice for the identification of causal processes (Salmon 1984, p. 154).

One clear problem for this view—and for the claim that tracer methodology supports a mark transmission definition of causation—is that scientists do not use these methods to establish or prove causality.¹⁴ Instead, they use these methods to study features of systems that they *already* recognize as causal. In many ways, this makes sense. Ensuring that you have a suitable tracer (and that it meets the tracer criteria) requires that you already know various things about the causal system of interest—these include what properties of the causal process you should (and should not) tag, on the basis of knowing what reliably moves through the system. Of course, we can attach a visualizable “tag” to anything we want, but this alone is no guarantee that it will track a causal process at all, much less one we are interested in. We know this, in part, because we know that marks can flow through processes that we do not consider causal (Danks 2017, p. 204).

¹⁴ Scientists do use tracers to show that, for example, one neuron (1) is causally connected to another (2) on the basis of the fact that a tracer moves from (1) to (2). However, using tracers to establish this causal connection depends on the prior view of axonal transport processes as causal.

In fact, the experiments that neuroscientists use to establish and prove causality look a lot more like interventionist-type experiments. The causal character of neural transport processes is often justified with various interventionist-type evaluations, such as lesion, neural degeneration, and constriction experiments (Morecraft et al. 2014). These studies reveal dependency relations between cause–effect properties along these processes. When material is prevented from flowing to upstream location A (by physically constricting the neuron) this material will fail to arrive at downstream location B. Relatedly, when material is allowed to flow to A, (in proper conditions) it will eventually flow to location B. A similar rationale guides lesion and degeneration experiments. The fact that degeneration experiments establish a difference-making picture of neural transport is suggested by Nassi et al, who state that “The very fact of Wallerian degeneration made it clear that the distal part of the axon was somehow dependent for its viability on substances supplied by the cell body, and, consequently, that there must exist mechanisms to transport these substances” (Nassi et al. 2015, p. 2). This dependency was “experimentally confirmed” with constriction studies (Weiss and Hiscoe 1948), and soon after, it was “exploited” with radioactively tagged amino acids that were taken up by these processes and identified with autoradiography (Nassi et al. 2015, p. 2). In this sense, the methodology of tracer experiments requires understanding the causal nature of neural transport processes and this causal nature is established by interventionist-type experiments.

Second, this analysis contributes to the literature by defending mark transmission accounts from a commonly accepted criticism. This criticism comes from Hitchcock and it involves the example of a sequence of colliding billiard balls (Hitchcock 1995). In this example, the end of a cue stick is covered in blue chalk and then used to strike a first billiard ball, which collides with a second and so on, until a final ball falls into the corner pocket. According to Hitchcock, these colliding billiard balls are a genuine causal process and this process is “marked” by the blue chalk that moves along them. However, in addition to this blue chalk mark, there is another mark that flows through this causal process, namely, linear momentum. Although both of these marks flow through this causal process, we view the linear momentum—and *not* the blue chalk—as causally (and explanatorily) relevant to the final ball falling into the corner pocket. Hitchcock’s main line of criticism is that mark transmission accounts are unable to distinguish the causally irrelevant chalk mark from the causally relevant linear momentum. As any account of causation should meet this explanatory relevance standard, mark transmission accounts are said to be inadequate. As Hitchcock states “our demand that explanations provide relevant information requires...that we be told which earlier properties the properties specified in the explanandum depend upon” (Hitchcock 1995, p. 311). A similar sentiment is echoed by Woodward who states that “there appears to be nothing in Salmon’s notion of mark transmission or the notion of a causal process that allows one to distinguish between the explanatorily relevant momentum and the explanatorily irrelevant blue chalk mark” (Woodward 2003, p. 352).

The main problem with this criticism is that the chalk mark is not a suitable tracer for this causal process—it would never meet scientific standards for tracers and it fails to meet our intuitive understanding of what a tracer is. Notice that we are very hard pressed to agree that this chalk mark can reliably move through a sequence of

colliding billiard balls. At the very least, this requires that the chalked portion of the first ball perfectly hits the second ball, and that the chalked portion of the second ball perfectly hits the third ball, and so on through this sequence. The movement of chalk through this process seems to rely on sheer luck and this fails to capture the principled rationale that guides scientific tracer experiments.¹⁵ Scientists expect tracers to meet very strict criteria—a tracer should mark some entity that reliably flows through a causal process and it should mark it firmly without falling off. Chalk does not mark any such entity and it can easily fall off to tag causally irrelevant factors. Using chalk as a tracer seems to assume that there is some material that reliably moves through this causal system, but no such material exists. This example requires that the chalk mark will effectively “jump” from one billiard ball to another, perfectly marking only cause–effect relations. What guarantees that the chalk will do this? If two billiard balls collide, and chalk flies into the air, what prevents it from marking entities in close proximity that are causally irrelevant to the system? This case does not teach us that tracers fail to mark causally relevant properties. It teaches us that chalk is not a suitable tracer for this system.

What lessons should we take away from this example? In contrast to Hitchcock’s claims, this example does not suggest that mark transmission accounts fail to provide an account of causation (other things teach us this, as mentioned above). Instead, this example shows us that chalk is not a suitable tracer for this system. A deeper lesson is that, if we want to understand causal reasoning with tracers and tags, we should carefully examine our best scientific work in this area and use this to inform an account of tracer methodology. Relatedly, we should carefully check our philosophical examples to scientific practice to ensure that these examples accurately represent scientific methodology. A basic understanding of tracer criteria in neuroscience shows that there are flaws with this chalk mark–billiard ball example and that there are principled reasons that guide tracer methodology.

This analysis of tracer techniques suggests that they fall short of supporting mark transmission accounts of causation. This is motivated, in part, by the fact that scientists do not use these techniques to establish causality—they use them to study systems that they already view as causal. Despite challenges for defining causation on the basis of tracers (or mark transmission), these techniques do have a principled rationale. In particular, these techniques do not succumb to Hitchcock’s criticisms, as these criticisms apply to a case that fails to meet scientific (and likely ordinary life) standards for tagging and tracer experiments. Appreciating the scientific use of these methods, provides a clearer picture of their implications for philosophical analyses of causation and the study of causal systems.

4 Neuroanatomical tracers: Causal information

What types of insights relevant to debates in philosophy of neuroscience about the nature of causal discovery does a philosophical analysis of tracer technology yield?

¹⁵ As Salmon and Reichenbach viewed their mark transmission accounts as supported by scientific tracer examples, the failure of this example to reflect scientific reasoning about tracers should raise a red flag.

One main insight has to do with the type of causal information that tracers provide. In order to clarify this, consider mechanistic accounts of explanation, which have dominated philosophy of neuroscience for more than a decade. According to these mainstream accounts, causal information in neuroscience is always “mechanistic” (Machamer et al. 2000; Craver 2007; Kaplan and Craver 2011).¹⁶ In this work, the notion of “mechanism” often refers to a set of causal parts, which all mechanically interact to produce some behavior of interest. These mechanisms involve (a) constitutive or part-whole relationships, (b) descriptions that involve significant fine-grained causal detail, (c) causal relationships that are characterized in terms of “mechanical” language, and (d) frequent analogy to the ordinary life notion of a “machine” (Machamer et al. 2000; Bechtel and Richardson 2010; Craver 2007, Ross Forthcoming). First, mechanisms are constitutive in the sense that their causal components are said to stand in a part-whole relationship to the mechanism behavior they explain. In fact, mechanisms are relative to the effects they produce—discovering a mechanism requires first specifying some effect of interest and then “drilling down” to identify its causal parts. This discovery process involves “decomposition” and “localization,” in which a system is divided into identifiable lower-level parts that give rise to the mechanism’s behavior (Wimsatt 1974; Bechtel and Richardson 2010; Bechtel and Levy 2013).

Second, scientists often expect the causal relationships in mechanisms to be described in significant fine-grained detail. To say that you know the “mechanism of action” of a particular drug, means that you know intricate details about how it leads to some downstream effect, as opposed to just knowing that it has this effect. This feature of mechanisms is associated with our interest in understanding how they work and our assumption that this involves identifying significant information about their causal components, organization, and so on. The expectation that mechanisms contain significant causal detail is suggested by scientific claims and expressed by many philosophical accounts of mechanism (Machamer et al. 2000; Darden 2006; Craver 2007; Craver and Darden 2013), although it is not universally accepted by all philosophers.¹⁷ A common view is that mechanism descriptions should contain some kind of “complete” detail, while descriptions that lack full detail are merely mechanism “sketches” or “schemata” (Craver 2007).

Finally, causal relationships in mechanisms are often described in mechanical language, with force, action, or motion terms (Ross Forthcoming). This is related to the fact that scientific mechanisms are often analogized to ordinary life machines, such as car engines and clocks. For example, instead of saying that a machine’s component “causes” a downstream component to do something we are more likely to say that it “bends,” “pushes,” “pulls,” or “compresses” this downstream component. These mechanical descriptions are common in neuroscience—this is evidenced by scientist’s claims that a neurotransmitter “binds” to a receptor, which “opens” an ion channel,

¹⁶ Many of these views rely on a basic interventionist framework (Woodward 2003), but add much more in capturing the notion of “mechanism.”

¹⁷ While many diverse accounts of mechanistic explanation exist, some deny that mechanistic information involves fine-grained or significant causal detail (Bechtel and Levy 2013; Boone and Piccinini 2016; Craver and Kaplan 2020). My analysis relies on a “fine-grained detail” account of mechanism, which is suggested and argued for in other work (Machamer et al. 2000; Craver 2007; Darden 2006, Ross Forthcoming).

and ultimately produces an “influx” of ions into the cell. This mechanical terminology goes hand-in-hand with our assumption that mechanisms should be described in significant detail. Stating that component A “pushes” component B, provides more information than simply saying that A “causes” B.

Although mainstream views in philosophy of neuroscience claim that all causal information is mechanistic, the causal information provided by these neuroscientific tracers appears to be different. Tracers provide information about sequences of causal connections—these are sequences of steps along an individual axon or along a chain of neurons connected in series. These causal connections lack the constitutive feature of mechanisms, they are not described in significant fine-grained detail, and they emphasize mere causal connection over the mechanical detail of these connections. In order to clarify these points consider a further piece of information. Instead of referring to these neural connections as “mechanisms,” neuroscientists often describe them as “pathways,” which they analogize to ordinary life examples of pathways, such as roadways, highways, and city streets (Saleeba et al. 2019; Frick et al. 2013). These analogies are not just colorful, superficial expressions—they capture important features of the causal information identified by neuroscientific tracers.

The causal pathways identified by tracers have at least four main features: they are represented as having a (i) sequence of causal steps, where these steps (ii) track the flow of some entity or signal through a system, (iii) abstract from significant causal detail, and (iv) emphasize the connection aspect of causal relationships.¹⁸ When focused on a single neuron, tracers capture a sequence of causal steps along the length of the neuron. In the same manner that a roadway captures a route along which traffic can flow, these neural pathways outline a route along which material (and signaling information) can flow. These routes capture causal steps that are “abstract” in the sense that they capture the causal relations between spatial locations along the axon (upstream location 1, to intermediate location 2, to downstream location 3) without identifying any other causal information. In particular, tracers do not provide fine-grained mechanical information about how materials (or signals) move along these steps—they capture that these are the available steps, what their order is, their location in the body, and which upstream and downstream tissues they connect up.¹⁹ In other words, causal pathways capture “that” 1 is causally connected to 2, but not “how” they are causally connected, which is typically associated with mechanism information. This is similar to the causal structure of a roadway on a map—this structure tells us “that” a car can move from 1 to 2 to 3, but it does not tell us “how” the engine works to sustain this movement.

The causal pathway information that neuroscientific tracers reveal is not just relevant to studies of single neurons, but also for understanding the complex causal connections among many neurons in the brain and other nervous tissues. This connection information is represented in pathway maps, connectomes figures, and neuron wiring diagrams. In the same way that a roadmap depicts available causal routes that vehi-

¹⁸ For more on these features and other examples of the pathway concept in science, see (Ross Forthcoming) and Ross (2018).

¹⁹ Pathways are also abstract in the sense that they represent complex processes with an economy of causal steps (Ross Forthcoming, p. 13). For example, the process of signal transmission from the spinal cord to the leg can be represented in anywhere from one to three causal connections, which abstract from an innumerable set of molecular steps.

cles can travel along, neural pathway maps capture neural connections along which signals and information can flow. As mentioned above, the causal pathway information in these diagrams abstracts from fine-grain mechanistic detail about how these entities flow along causal channels. Abstracting from this mechanistic information is important because systems with the same causal pathway structure, can have different fine-grained mechanistic details. Suppose a set of roadways or neural pathways depicts connections from upstream location 1 to downstream location 2, which is connected to downstream location 3, but there is no causal connection from any of these to location 4. We know from this pathway structure that a vehicle or neural signal starting at 1 can make its way to 2 and 3, but that it cannot travel to location 4. The fact that we know this is unrelated to the fine-grained mechanistic details about how the vehicles or signals move. This is evidenced by the fact that the same pathway structure can have different mechanistic realizations—cars can have different engines (gas, diesel, electric, etc.) and neural signals can have different fine-grained causal components (ion channel types, density of these types, and so on) (Ross 2018). In these cases, the shared behavior across various systems can be explained by their similar causal pathway architectures, but not by their fine-grained mechanistic details, as these details differ from system to system.

Neuroscientific tracers are used to identify causal connections in some domain—these connections are often referred to as “causal pathways” and this connectionist information is represented in pathway maps, connectome figures, and neuron wiring diagrams. While the causal connections in these representations abstract from fine-grained mechanistic detail they also lack the part-whole and effect-relative nature of mechanisms. They lack these features because the pathway components represented in these maps do not all constitute a unified whole and they do not all “mechanically” interact to produce a single explanatory target. Where a drug’s mechanism of action captures an actual, single process (and single outcome) executed by the drug, pathway maps differ by capturing a variety of possible causal trajectories, as opposed to a single causal process. This is made more clear by recognizing that neural pathway maps, roadmaps, and circuit diagrams represent a different sort of causal structure than a machine. While a mechanism captures a set of causal parts that all interact to produce a particular outcome, pathway maps represent a variety of available causal routes in some domain. These differences are related to the fact that different causal investigative strategies are used to study mechanisms and pathways. Mechanisms are studied through a process of “decomposition and localization”—some explanatory target is fixed and then one “drills down” to identify the lower-level mechanical parts that interact to produce the behavior of interest. Pathways, on the other hand, are identified by tracers which involve a strategy of “expanding out” to identify causal connections, as opposed to “drilling down” to identify causal parts. A tracer is attached to some material or information, after which it flows through causal pathways and available causal routes in the system.

How do neuroscientists use pathway information to understand and explain neural systems? When a comprehensive set of neural connections is represented in a neural pathway map, this “connectome” can be used to better understand how neural circuits work. This is seen in the use of “nanoscale connectomes” for *C. elegans* and *Ciona intestinalis* (DeWeerd 2019). For these organisms, pathway maps are an “important

first step” in understanding animal behavior and they serve as a “starting point for various hypotheses” (Jabr 2012). For example, inferences from structure to function are made by ablating particular neurons, determining where they are located in the connectome, and what resulting outcome is produced. In addition to providing positive suggestions about function, circuit diagrams and tracers also serve an important role in hypothesis exclusion. These methods are viewed as a “powerful winnowing tool” that is capable of excluding various hypotheses on the basis of absent connections (DeWeerd 2019). Tracers can provide information about which areas are not causally connected, which can help in narrowing down the possible functions of a given neural circuit (Kohara et al. 2013).²⁰

Although tracers are able to provide these types of information, there are particular details which they cannot provide. For example, tracers fail to capture more fine-grained information about causal connections, including the dependency-relations they exhibit. Tracers capture directionality and location of flow, but they do not clarify “how” this flow takes place. In fact, neuroscientists admit that there are significant gaps in their understanding of how neural transport processes work, despite the fact that they are tagging and tracing them (Saleeba et al. 2019). As mentioned above, this additional fine-grained information is often irrelevant for particular questions about these systems. If we want to know how traffic flows through a particular city, it does not matter much how the vehicles drive along these roads, such as the different types of engines that might propel their movement. In this case we care about the more abstract causal structure and constraints of the system. It is this information that tracers provide and that pathway maps are intended to convey.

If neural pathway maps represent causal information, but it fails to meet mechanistic standards, how should this information be understood? When scientists discuss the causal information and causal structures that tracers identify, they refer to them in terms of causal “pathways” and not causal “mechanisms.” In this context, “pathway” refers to something quite different than “mechanism.” Pathways are causal structures that involve (i) a sequence of causal steps, in which these steps outline the (ii) flow of material or information, they (iii) abstract from significant fine-grained detail, and they (iv) emphasize the “connection” feature of causal relationships (as opposed to a mechanical one) (Ross Forthcoming). These pathways capture causal connections in some system, which are capable of answering explanatory why-questions that mechanistic information cannot answer. For example, if you want to know whether one neural area is connected to another, you simply need to know whether a causal pathway exists between them—you do not need to know the fine-grained mechanical details of “how” information or material is transmitted from one to the other. Furthermore, where causal systems referred to as “mechanisms” are analogized to machines, causal systems referred to as “pathways” are analogized to roadways, highways, and city streets. These ordinary life pathways often exhibit the same features (i)–(iv) that are characteristic of neural pathways. Scientists appear to be analogizing these neural systems to ordinary life examples of pathways in order to make complicated features of these neural connections more obvious and cognitively accessible.

²⁰ If a tracer fails to move from one neural area to another, this suggests that these areas lack anatomical connection and functional involvement.

Part of what this analysis shows is that we cannot use the single concept of “mechanism” to interpret and understand all causal structures in this domain. There are other causal concepts, such as the notion of “pathway,” that capture distinct causal structures and distinct types of causal information. Recent work suggests that causal structures referred to as “mechanisms” and “pathways” have unique features, involve different causal investigative strategies, and figure in different types of explanations (Ross Forthcoming). Part of what this work reveals is that causal structures and causal investigative strategies in neuroscience are much more diverse than mainstream mechanistic views have suggested.

5 Conclusion

This paper has examined how neuroscientific tracer techniques are used and the particular types of information they provide. An analysis of these techniques bears on mark transmission and mechanistic accounts of causal explanation. This work suggests that, while scientifically-defined tracer methods fail to support mark transmission definitions of causation, they remain legitimate tools of causal study. Furthermore, tracer techniques help capture the importance of “causal pathways” in neuroscience and how they are poorly accommodated by existing mechanistic accounts. These techniques suggest that neuroscience contains a diversity of causal concepts, causal structures, and causal reasoning, which are not all captured with the single notion of “mechanism.” While these conclusions relate to specific questions about the use of tracers in uncovering causal structure, there are many further questions about how this structure bears on our understanding of brain functioning, how it can be extrapolated from animal models to humans, and how it might provide useful principles that generalize across various nervous systems.

An analysis of these tracer techniques also uncovers questions that should be explored in future work. For example, while neuroscientific tracers typically track changes in a material’s location over time, other scientific tracers appear to track changes in a material’s constitution over time. This is seen in tracers for metabolic pathways, which identify changes in a metabolite’s constitution as opposed to changes in its physical location. What are the similarities and differences across these tracers? Additionally, neuroscientific tracers involve an intriguing causal study “by proxy” strategy, in which one causal process is studied in order to learn about another. Under what circumstances can studying one causal system provide information about another? Finally, constraints appear to play a significant role in the causal processes studied with tracer and tagging experiments. How are these constraints to be understood and what exactly is the role that they play? Given the fundamental role of tracer experiments in neuroscience careful analysis of these methods is a welcome move in philosophy and one that is likely to provide many more insights into causal reasoning in science.

Acknowledgements I would like to thank James Woodward, Chris Hitchcock, David Danks, John Bickle, Carl Craver, and audience members at the “NeuroTech” conference, sponsored by the Mellon Institute and Center for Philosophy of Science at the University of Pittsburgh, for helpful comments on this paper.

References

- Bechtel, W., & Levy, A. (2013). Abstraction and the organization of mechanisms. Technical report 2.
- Bechtel, W., & Richardson, R. C. (2010). *Discovering complexity*. Cambridge, MA: The MIT Press.
- Boone, W., & Piccinini, G. (2016). Mechanistic abstraction. *Philosophy of Science*, 83, 686–697.
- Callaway, E. M. (2008). Transneuronal circuit tracing with neurotropic viruses. *Current Opinion in Neurobiology*, 18(6), 617–623.
- Craver, C., & Darden, L. (2013). *In search of mechanisms*. Chicago: The University of Chicago Press.
- Craver, C., & Kaplan, D. M. (2020). Are more details better? On the norms of completeness for mechanistic explanations. *The British Journal for the Philosophy of Science*, 71(1), 287–319.
- Craver, C. F. (2007). *Explaining the brain*. Oxford: Oxford University Press.
- Danks, D. (2017). Singular causation. In M. Waldman (Ed.), *The Oxford handbook of causal reasoning*. Oxford: Oxford University Press.
- Darden, L. (2006). *Reasoning in biological discoveries*. Cambridge: Cambridge University Press.
- DeWeerd, S. (2019). How to map the brain. *Nature*, 571, S6–S8.
- Dowe, P. (2018). Wesley Salmon's process theory of causality and the conserved quantity theory. *Philosophy of Science*, 59, 1–23.
- Ekstrand, M. I., Enquist, L. W., & Pomeranz, L. E. (2008). The alpha-herpesviruses: Molecular pathfinders in nervous system circuits. *Trends in Molecular Medicine*, 14(3), 134–140.
- Frick, A., Ginger, M., Habert, M., Conzelmann, K.-K., & Schwarz, M. K. (2013). Revealing the secrets of neuronal circuits with recombinant rabies virus technology. *Frontiers in Neural Circuits*, 7, 1–15.
- Friston, K. J. (2011). Functional and effective connectivity: A review. *Brain Connectivity*, 1(1), 13–36.
- Heimer, L. (2005). Gross dissection of the human brain. *Neurosurgical Focus*, 18, 1–2.
- Hitchcock, C. (1995). Salmon on explanatory relevance. *Philosophy of Science*, 62, 304–320.
- Hooker, C. (2012). On the import of constraints in complex dynamical systems. *Foundations of Science*, 18(4), 757–780.
- Hu, W., Liu, D., Gu, J., Zhang, Y., Gu, X., & Gu, T. (2013). Neurological function following intra-neural injection of fluorescent neuronal tracers in rats. *Neural Regeneration Research*, 8, 1253–1261.
- Huh, Y., Oh, M. S., Leblanc, P., & Kim, K.-S. (2010). Gene transfer in the nervous system and implications for transsynaptic neuronal tracing. *Expert Opinion on Biological Therapy*, 10(5), 763–772.
- Jabr, F. (2012). The connectome debate: Is mapping the mind of a worm worth it? *Scientific American*.
- Kaplan, D. M., & Craver, C. F. (2011). The explanatory force of dynamical and mathematical models in neuroscience: A mechanistic perspective. *Philosophy of Science*, 78(4), 601–627.
- Kästner, L. (2017). *Philosophy of cognitive neuroscience: Causal explanations, mechanisms, and experimental manipulations*. Berlin: Walter de Gruyter.
- Kohara, K., Pignatelli, M., Rivest, A. J., Jung, H.-Y., Kitamura, T., Suh, J., et al. (2013). Cell type-specific genetic and optogenetic tools reveal hippocampal CA2 circuits. *Nature Neuroscience*, 17(2), 269–279.
- Kristensson, K. (1970). Transport of fluorescent protein tracer in peripheral nerves. *Acta Neuropathologica*, 16(4), 1–8.
- Lanciego, J. L., & Wouterlood, F. G. (2011). A half century of experimental neuroanatomical tracing. *Journal of Chemical Neuroanatomy*, 42(3), 157–183.
- Levy, S. L., White, J. J., & Sillitoe, R. V. (2015). *Wheat germ agglutinin (WGA) tracing: A classic approach for unraveling neural circuitry*. Totowa: Humana Press.
- Lichtman, J. W., & Sanes, J. R. (2008). Ome sweet ome: What can the genome tell us about the connectome? *Current Opinion in Neurobiology*, 18(3), 346–353.
- Machamer, P., Darden, L., & Craver, C. F. (2000). Thinking about mechanisms. *Philosophy of Science*, 67, 1–25.
- Martin, X., & Dolivo, M. (1983). Neuronal and transneuronal tracing in the trigeminal system of the rat using the herpes virus suis. *Brain Research*, 273, 253–276.
- Mi, D., Yuan, Y., Zhang, Y., Niu, J., Wang, Y., Yan, J., et al. (2019). Injection of Fluoro-Gold into the tibial nerve leads to prolonged but reversible functional deficits in rats. *Scientific Reports*, 9, 1–8.
- Morecraft, R. J., Ugolini, G., Lanciego, J. L., Wouterlood, F. G., & Pandya, D. N. (2014). Classic and contemporary neural tract-tracing techniques. In H. Johansen-Berg & T. E. J. Behrens (Eds.), *Diffusion MRI* (pp. 359–399). Amsterdam: Elsevier.
- Nassi, J. J., Cepko, C. L., Born, R. T., & Beier, K. T. (2015). Neuroanatomy goes viral!. *Frontiers in Neuroanatomy*, 9, 80.
- Reichenbach, H. (1971). *The direction of time*. Berkeley: University of California Press.

- Ross, L. N. (2018). Dynamical models and explanation in neuroscience. *Philosophy of Science*, 82(1), 32–54.
- Ross, L. N. (Forthcoming). Causal concepts in biology: How pathways differ from mechanisms and why it matters. *The British Journal for the Philosophy of Science*.
- Russell, B. (1948). *Human knowledge: Its scope and limits*. Routledge: Taylor & Francis.
- Saleeba, C., Dempsey, B., Le, S., Goodchild, A., & McMullan, S. (2019). A student's guide to neural circuit tracing. *Frontiers in Neuroscience*, 13, 1–19.
- Salmon, W. C. (1984). *Scientific explanation and the causal structure of the world* (pp. 135–157). Princeton, NJ: Princeton University Press.
- Salmon, W. C. (1997). Causality and explanation: A reply to two critiques. *Philosophy of Science*, 64, 461–477.
- Salmon, W. C. (1998). *Causality and explanation*. Oxford: Oxford University Press.
- Shi, Y., & Toga, A. W. (2017). Connectome imaging for mapping human brain pathways. *Molecular Psychiatry*, 22(9), 1230–1240.
- Sporns, O. (2007). Brain connectivity. *Scholarpedia*, 2(10), 4695.
- Sporns, O. (2012). *Discovering the human connectome*. Cambridge: The MIT Press.
- Sporns, O., & Bassett, D. S. (2018). Editorial: New trends in connectomics. *Network Neuroscience*, 2(2), 125–127.
- Ugolini, G. (2011). Rabies virus as a transneuronal tracer of neuronal connections. In A. C. Jackson (Ed.), *Advances in virus research* (pp. 165–202). Amsterdam: Elsevier.
- Weiss, P., & Hiscoe, H. B. (1948). Experiments on the mechanism of nerve growth. *The Journal of Experimental Zoology*, 107, 315–395.
- Wimsatt, W. C. (1974). Reductive explanation: A functional account. *Philosophy of Science*, 671–710.
- Winning, J. (2018). Mechanistic causation and constraints: Perspectival parts and powers, non-perspectival modal patterns. *The British Journal for the Philosophy of Science*, 71(4), 1–23.
- Winning, J., & Bechtel, W. (2018). Rethinking causality in biological and neural mechanisms: Constraints and control. *Minds and Machines*, 28(2), 287–310.
- Woodward, J. (2003). *Making things happen*. Oxford: Oxford University Press.
- Woodward, J. F. (2016). Causation and manipulability. *Stanford Encyclopedia of Philosophy*.

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.