

Maddy vs. Quine on Innate Concepts

Revisiting A Perennial Debate in Light of Recent Empirical Results

Abstract

In his posthumously published work, Quine abandons his empiricist principle that humans do not have any innate concepts, or knowledge. He does so in light of empirical research that Penelope Maddy capitalizes on to develop her own naturalized epistemology. The empirical research in question is due to the pioneering work of developmental psychologist Elisabeth Spelke. Spelke employs the method of habituation and preferential looking to argue that human infants have innate concepts, and that they have some knowledge about what can and cannot happen to physical objects. Taking into account empirical studies as well as methodological considerations, this article examines whether this research can support these strong philosophical conclusions drawn from it, finding that it likely cannot provide such support.

Keywords: Quine; Maddy; Spelke; nativism; empiricism; naturalism; innate concepts

Version: rev v2

This is a post-peer-review, pre-copyedit version of an article published in *Philosophia*. The final authenticated version is available online at: <https://rdcu.be/bMFAI>

1. Introduction

This article revisits a debate that has, in one form or another, resurfaced across the millennia of western philosophical thinking: How much of our cognitive and linguistic competence is innate? Typically, empiricists reject the idea that substantial parts of our cognitive capacities are innate, that is, present from birth and essentially independent of learning. The dialectical counterparts of empiricists, usually called ‘rationalists’ or ‘nativists’, maintain that humans have to be credited with sophisticated innate cognitive equipment such as knowledge, concepts, or ideas to explain how humans are able to acquire their first language, or to learn to apply concepts.

Throughout his published work, W.V.O. Quine clearly embraces the empiricist position in this debate: Putting an innate sense of similarity¹ aside, he refuses to attribute any innate cognitive capacities to human infants. This conviction shapes many features of Quine’s naturalized epistemology. Naturalized epistemology (see Quine, 1969, esp. pp. 82ff.) tries to explain, in a scientific way, how human infants can acquire a language, which is, according to Quine, tantamount to acquiring a primitive theory of the world. The epistemological purpose of this genetic account, in turn, is to scientifically elucidate the relationship between scientific theories and their evidence.

Quine maintains that an infant makes his first steps into language by acquiring his first observation sentences. These observation sentences are conceived holophrastically as one-word sentences that are keyed to certain stimuli. An example for such an observation sentence is “Red!” To learn these first observation sentences, Quine holds, the child must be inclined to give more weight to the relevant similarities than to the irrelevant ones – where the relevant similarities are those which matter for the meanings of the observation sentences that the infant tries to learn (compare Quine 1995, p. 19, and in particular Quine 1960, p. 83). Given a number of global stimuli, that is, the sum total of sensory input at different times, the child must be biased to generalize along the aspects of the stimuli which matter for his parent, say the red patch in the visual field. Otherwise, the infant’s task would be hopeless: How could it, given the myriad of different aspects in any given global stimulus – olfactory, visual, acoustic, proprioceptive, etc. –, identify the tiny part that matters, namely the red patch in the visual part of the stimulus?

Innate quality spaces are Quine’s theoretical response to this massive underdetermination problem: The infant’s quality spaces situate global stimuli involving red close together, instead of grouping all those global stimuli where, say, a certain distinctive gust of wind is present. This way, the child is able to inductively generalize the stimulus meaning of his first observation sentences.

In contrast to Quine, empirical psychology underwent a strong cognitive turn in the 1980ies. There is a prominent position in contemporary developmental psychology that fundamentally challenges Quine’s empiricism. I call this position nativism (I detail on this conception below, section 2). Briefly, nativists claim to provide convincing empirical evidence for two different claims. First, nativists claim that pre-linguistic infants use a sophisticated set of criteria to

¹ Quine uses innate quality spaces, innate senses of similarity, or innate standards for perceptual similarity for very similar conceptions that fulfill the very same function in his theory.

individuate and refer to instances of a kind of physical object, so-called *Spelke objects*. Second, nativists claim that pre-linguistic infants have knowledge about the physical behavior of Spelke objects that allows them to judge the possibility or impossibility of certain events involving such Spelke objects. For instance, just like adults, infants know that one solid object cannot easily be penetrated by another. Using a common terminology (see, for instance, Glock 2012, 213), nativism claims both that pre-linguistic infants can think of certain things, namely Spelke objects, and think that such-and-such is the case. The latter is tantamount to attributing beliefs to the respective infants. Contemporary nativists mostly support their position with results from a research method that is called habituation and preferential looking, which will be discussed in this article.

Very surprisingly, the research to be considered here in support of nativism, including the nativist conclusions drawn from it, is mentioned very favorably in Quine's posthumously published work. In a note standing under the title *The Innate Foundational Endowments*, dated March and May 1996, that is, roughly one year after the publication of *From Stimulus to Science*, Quine clearly rejects empiricism:²

Descartes thought we had innate knowledge and innate ideas. Locke thought not. I despair of sharpening the issue by defining the term 'idea'. Definition even of 'knowledge' is in trouble since Gettier's challenge of the definition of knowledge as true and warranted belief. However, we need no such sharpening of the issue to see that the evidence favors Descartes over Locke. It is reported that a child in early infancy will show anxiety when glowered at, and will even imitate a frown or smile. The infant will show anxiety also after seeing an object steadily approach a screen, pass behind it, and then fail to emerge in due course at the far side. Further, there is our spontaneous identifying of different perspectives of an angular object as all somehow the same thing. All this is innate, or, as we say, instinctive. Hence innate knowledge, we might say, and no words to express it.

This is a remarkable passage: Quine explicitly accepts that humans have innate knowledge. With this radical change of position, Quine invalidates most of naturalized epistemology, as he pursued it over the course of at least 35 years: Knowledge about the physical world is the end-point of a complex process whose reconstruction is the centerpiece of naturalized epistemology. The evidence that Quine cites very cursorily in this passage stems from two different disciplines (putting aside the reports according to which infants become anxious when glowered at, and imitate facial expressions). The second piece of evidence that Quine cites is the fact that we are able to realize, without reflection, that different perspectives on an angular object (which cause very different proximal stimuli on the retina) are perspectives on *the same* object. Such phenomena are investigated by perceptual psychology.³

In the present article, I delineate (section 3) and critically assess (section 4) the research method that Quine implicitly alludes to when stating that infants show a marked reaction when

²I am quoting from the edition by Dagfinn Føllesdal and Douglas B. Quine: *Confessions of a Confirmed Extensionalist and Other Essays*, see References.

³Quine's example would count as a kind of perceptual constancy, namely shape constancy, in contemporary perceptual psychology. See Burge (2010) for an account of the philosophical significance of contemporary perceptual psychology.

they see an object disappear behind a screen but fail to reappear where it would, had it followed a steady trajectory. The method is called habituation and preferential looking, it was pioneered by Elizabeth Spelke in the mid-eighties, and it is currently the main way in which nativists support their position.

Penelope Maddy relies on results from this research method to support one of three main pillars of her own version of naturalized epistemology, what she calls *Second Philosophy of Logic*: she wants to establish, using purely scientific data, (1) that pre-linguistic infants are able to reason according to a specific logical structure, what she calls *rudimentary logic*, (2) that the world itself has this structure, and (3) that humans are so endowed because the world is so structured (compare Maddy 2007, p. 271). I here only examine the first pillar of her program, the claim that pre-linguistic humans reason according to rudimentary logic.

The nativist claims, if warranted, would undermine the basics of Quine's naturalized epistemology: rather than being the end-point of a complicated development that starts with the acquisition of holophrastic observation sentences, beliefs about the physical world would be with infants almost from the start – and certainly before they acquire their first language. In the following section (section 2) I detail nativism, as defended by contemporary developmental psychologists, I introduce the first pillar of Maddy's *Second Philosophy of Logic*, and I discuss Maddy's claim that nativism vindicates this first pillar.

To conclude this introductory section, a note on my use of the adjective "cognitive". An explanation takes recourse to cognitive structures or abilities if these structures or abilities are inseparable from what the explanation in question would count as thinking. For instance, Quine's quality spaces qualify as cognitive structures, albeit primitive ones, since they are intimately connected with inductive inference, which is clearly an act of thinking. In contrast, explanations of the behavior of infants with recourse to simple preferences for more motion or novelty (see below, section 4.2) do not qualify as cognitive. Finally, Tyler Burge's explanations of animal behavior that invoke structured, intentional perceptual representations do not count as cognitive, since Burge emphasizes that these representations are not achieved by a central cognitive faculty, but rather by the visual system (see Burge 2010, pp. 438ff.). To categorize Burge's perceptual representations as non-cognitive might raise eyebrows, but it serves the purpose of the present study.

2. Spelkean Nativism and Maddy's Pre-Linguistic Rudimentary Logic

The preferential looking paradigm was pioneered by developmental psychologist Elizabeth Spelke. Since her seminal work in the mid-eighties, she has become one of the most influential proponents of nativism, and arguably one of the most influential living developmental psychologists *tout court*.

Spelke is particularly known for the claim that pre-linguistic infants possess a concept of physical object – what came to be simply called the *Spelke object concept*, which refers to *Spelke objects*. The concept of a Spelke object determines identity criteria for the respective objects, including "spatial contiguity, common fate, and continuous motion" (Maddy 2007, 253). According to this specification, the identity of a Spelke object depends on whether it coheres spatially, whether it moves as a unit, and whether it moves without making abrupt changes in its trajectory. Hence, to credit infants with the Spelke object concept means that infants individuate

Spelke objects based on these three criteria. For instance, two oranges that are moving in different directions will count as two Spelke objects, since they are neither spatially contiguous, nor moving into the same direction.

Furthermore, infants that possess the concept of a Spelke object have some knowledge, and hence beliefs, about the behavior of Spelke objects. Compare how Baillargeon, Spelke, and Wasserman (1985) conceive of the significance of the so-called drawbridge experiments (which will be discussed below (section 3.2):

Our experiments provide evidence that by 5 months of age, infants already appreciate two aspects of the behavior of objects. First, they understand that an object continues to exist when occluded, and that it exists not as a disembodied image residing somewhere behind the occluder but as a solid, three-dimensional entity occupying a specific spatial location. Second, they understand that an object can move only through space not occupied by other objects. (Baillargeon, Spelke, and Wasserman 1985, 206)

According to this passage, pre-linguistic infants, indeed, infants as young as five months, understand that occluded objects continue to exist as solid, physical entities. Furthermore, such infants understand that one object cannot simply move through another one. That is, the infants can appreciate the impossibility of this event. To put it differently, once infants have cognized some x and y as Spelke objects, they believe that x persists, and hence expect that it will not just evaporate when occluded, and they also believe that x and y are solid, and hence understand that x cannot simply pass through y . Naturally, they are surprised if it looks as if x would, impossibly, pass through y .

Maddy is convinced that this nativist position decisively supports one of three core claims of her *Second Philosophy of Logic*.⁴ Compare how Maddy states this claim in the following passage:⁵

[...] human beings believe the simple truths of rudimentary logic because their most primitive cognitive mechanisms allow them to detect and represent the KF-structures in the world. (Maddy 2007, 245)

I will sketch Maddy's notion of a KF-structure, or, as she later came to call it 'rudimentary logical structure', in a minute. For now, let me focus on the precise claim that Maddy here makes by stating that the most primitive cognitive mechanisms enable us to "detect and represent" these rudimentary logical structures in the world. Maddy (2007, 264) makes clear that to be "most primitive" implies, at least, *pre-linguistic*.

It is more difficult, however, to get a clear conception of how to read "detect and represent". Such wording suggests a perceptual ability. However, this not only fits ill with her

⁴This is the basic presupposition of her entire case on Maddy (2007, 245–70).

⁵The basic idea behind Maddy's notion of a KF-Structure (read: Kant-Frege-Structure) was to take Kant's notion of logical form and to update it in the way "contemporary students of the subject" would (Maddy 2007, 227). In particular, she suggests replacing Kant's subject-predicate relation with Frege's argument-function relation (Maddy 2007, 227). On the same page, Maddy proposes an analogous amendment of the Kantian tables of forms of judgment and categories to accommodate Frege's notion of quantification. In later work, especially in Maddy (2014), she has dropped this terminology of KF-structures, and I think it's good riddance: Maddy is in no way committed to any Kantian views on logic, epistemology, or metaphysics. Hence the terminology is rather misleading.

specification of these mechanisms as *cognitive*. On other occasions, she writes that the infant *thinks* according to these rudimentary logical structures, and that the infant *understands* them (see, in particular, (Maddy 2014, 100–101)).

Given her overall project, I submit that Maddy intends the latter, cognitive reading of “detect and represent”: She does not merely want to make some claims about the grouping capacities of perceptual systems of infants. Rather, she wants to claim that these infants *think* in a certain way, namely in the form of rudimentary logic. Here is how Maddy specifies rudimentary logical structure (you can read ‘KF-structures’ as ‘structures having rudimentary logical form’, compare footnote 4):

Let’s press a bit further into the KF-structure. A given object a might enjoy the property P and it might fail to enjoy that property, but nothing in this picture precludes the possibility of indeterminate cases. So for any property P , the domain of a KF world will divide into those objects that have P , those that don’t have P , and those for which P is indeterminate; and perhaps even these boundaries between these groupings are somewhat fuzzy. The same goes for objects standing or not standing in the various relations of the KF-world. For compound states: not- (...) obtains if (...) fails; fails if (...) obtains; and is otherwise indeterminate. ((...) and (...)) obtains if (...) and (...) do; fails if either (...) or (...) fails; and is otherwise indeterminate. ((...) or (...)) obtains if (...) or (...) does; fails if both (...) and (...) fail; and is otherwise indeterminate. A universal generalizes conjunction, so (every x) (... x ...) holds if for every a in the domain of our KF-world, (... a ...) obtains; fails if there is an a in the domain such that (... a ...) fails; and is otherwise indeterminate. Dependencies are more troublesome, because their status isn’t always settled by those of their parts. In fact, all that seems fixed is that if (...) and (if (...), then (...)) hold, then (...) must also – under these conditions, (...) cannot fail or be indeterminate. (Maddy 2007, 229)

A central point that Maddy makes in this passage is that rudimentary logic is three-valued, since the world contains many objects for which it is not clear whether a given property P does or does not apply to them. A note on her notation: Maddy uses ‘and’ and ‘or’ to refer to logical conjunctions and disjunctions, respectively. She uses dots ‘...’ and underline ‘_’ to express schemata; hence, ‘(...)’ refers to any expression whatsoever, ‘(...) and (...)’ refers to any two expressions combined by a logical ‘and’.

What I want to emphasize here is less the details of her three-valued logic (which may be interesting in themselves), but rather the general structure of rudimentary logic. This structure includes *objects, properties, relations, and dependencies*. Furthermore, it obviously contains statements, or propositions that can obtain or fail to obtain. Furthermore, these propositions are of some complexity: they do contain, in addition to conjunctions and disjunctions, universal generalizations and, as the last sentence makes clear, *modus ponens*. It seems plausible that if Maddy can show that infants do cognize their environment in such rudimentary logical terms, then this can indeed provide a cognitive foundation for science.

In the following two sections, I focus on Maddy’s case for the claim that human infants think about the world in terms of objects-with-property structures of rudimentary logic (that is, I leave out her analogous cases for relations and dependencies of rudimentary logic). Briefly, this case

proceeds as follows: Maddy (2007, 245–70) reviews habituation and preferential looking experiments, accepts the nativist conclusions that Spelke and her fellow psychologists draw from the results of these experiments, and then assumes that these conclusions vindicate her claims regarding the logical structure of pre-linguistic infants' thoughts.

This last step may need some further elaboration. Even if it is granted that the respective experiments vindicate nativism, it seems unclear whether nativism as such supports Maddy's central claim about the logical structure of pre-linguistic infants' thinking. Even if it is granted that there is evidence that infants think in terms of Spelke objects with properties, indeed, that they think like adults in some respects, this may not guarantee that they think in terms of the objects and properties of rudimentary logic. This is simply because this is not even clear in the case of full-grown adults. Hence, *a fortiori*, the looking preferences of pre-linguistic infants do not exclude alternative conceptions of logical structure, such as the term-based logic proposed by Oderberg (2005). More needs to be said why the evidence does favor Maddy's specific logical structure at the expense of the numerous alternative proposals.

Note, however, that this worry does not question the basic thrust of Maddy's approach: if it can be empirically shown that pre-linguistic infants do draw inferences about their environments, then this establishes the first pillar of her Second Philosophy of Logic – irrespective of the details of the infants' inferential patterns. In the worst case, Maddy might have to make some minor modifications to her concept of rudimentary logic.

In sum, nativism, as it is understood here, claims that pre-linguistic infants can individuate so-called Spelke-objects based on the criteria of cohesion, common motion and continuous motion. Furthermore, pre-linguistic infants can attribute properties such as persistence and solidity to these Spelke objects. Based on this, they are able to understand some aspects of the physical behavior of objects, that is, they think like adults in this regard. Maddy holds that nativism implies the first central claim of her *Second Philosophy of Logic*, according to which pre-linguistic infants think in accordance with rudimentary logic.

3. In Support of Nativism: Preferential Looking Experiments

In this section, I detail Maddy's evidence for the claim that pre-linguistic infants reason according to rudimentary logic. More precisely, I focus on the basic part of rudimentary logical structure, objects with properties.

Maddy's argument for the claim that infants cognize the world in terms of Spelke objects with properties is based on the research method called *habituation and preferential looking* (Maddy 2007, 247). The basic structure of experiments conducted using this method involves two steps. First, *habituate* infants to a certain scene. This involves repeatedly showing a certain arrangement of items, or a certain event, to infants until they lose interest in the scene or event. It is assumed that infants lose interest if their looking time at the scene or event falls below a certain threshold. For instance, the relevant display could be a yellow ball resting on a blue bar. To habituate the infants to this display, they are placed before a little stage, then the display is presented to them, perhaps by removing a screen, or curtain. As soon as they look away, the display is hidden from them. After that, the display is shown again, etc., until their looking time falls below the threshold.

In the second step of the experiment, the infants are presented with several *test displays*. What these test displays show depends entirely on the research interest of the experimenter. To continue the example introduced in the preceding paragraph, if the goal of the experiment is to determine whether infants perceive scenes as novel that only differ in color, the test displays could consist in balls of various colors resting on bars of various colors. If the infants look longer at displays of non-yellow balls resting on non-blue bars than on the habituation display, and given that the form of both balls and bars is the same as in the habituation display, this is taken as evidence that the infants perceive scenes as novel that only differ in colors, since they look longer at such scenes. It is then said that the infants prefer (to look at) the display with colors that differ from the habituation colors. Hence, a basic presupposition of the entire method is that infants find novel displays more interesting than familiar ones, and therefore look longer at novel displays.

The key advantage of this method is that it is applicable to infants at a very young age, since it does not require that infants have mastered any sophisticated behavioral abilities such as grasping or crawling: it suffices that the infant is able to focus her gaze on a specific scene.

In the following, I reconstruct the experimental evidence from this preferential looking method that Maddy enlists to show that human infants think according to the structure of rudimentary logical objects with properties. The reconstruction is not exhaustive, but merely representative: I chose two experiments that have been particularly influential in the research domain and on which Maddy puts special weight. Such a representative sample suffices for the purposes of my critical appreciation of the method below (section 4.2).

3.1. *How Infants Individuate Spelke Objects*

First, consider the following experiment that establishes, according to Maddy (2007, 251), that infants use continuous motion to determine object identity (as a matter of fact, Maddy (ibid.) never explicitly says so, but she approvingly quotes a passage to this effect). The experiment, or better, the two experiments have been conducted by Spelke et al. (1995, 118–22).⁶ In the first experiment, the set-up is as follows: on the stage, there are two screens, one on the left and the other on the right of the stage, behind which rods can move freely. Then, three different groups of four-month-old infants are habituated to three different kinds of events that all take place on this stage (Spelke et al. 1995, 119–20). The first group is shown the continuous bar movement, that is, a bar appears on the very left side of the stage and moves continuously over to the very right side of the stage, passing behind the two screens. The second group is shown the discontinuous bar movement, that is, the bar appears on the left side of the stage, disappears behind the left screen, and then, in due time, another bar appears on the right-hand side of the right screen and moves over to the right edge of the stage. The third group of infants of the same age is habituated to the stage and the setting by showing them events unrelated to the present issue.

Then, all three groups are presented with two different displays of the stage with the screens removed. In the first display, two rods are shown where previously the screens have been located. In the second display, only one rod is shown where previously one of the two screens

has been located.

According to the hypothesis of the experimenters, continuous movement constitutes a core criterion that infants rely on when individuating objects; hence, the group that has been shown the continuous movement should be surprised, and hence look longer, at the display of two rods, since they have represented the situation as involving only one rod. In contrast, the group that has been shown the discontinuous movement should be surprised by, and hence look longer at, the display with only one rod, since they have represented the situation as involving two rods. In both cases, the differences in looking time should be significantly above any differences in looking time at the same displays of infants in the control group – otherwise, the differences in looking time could be due to, say, a general preference of infants for more things.

The results of the experiment are as follows. The infants habituated to the continuous movement looked longer at the display of two rods, while the infants habituated to the discontinuous movement looked longer at the display of just one rod. The infants of the control group, finally, looked longer at the presentation of two rods. Unfortunately, the differences between the looking times of these three groups were too small to be significant (Spelke et al. 1995, 123). Thus, the interpretation of the outcome of the first experiment is summarized by the experimenters as follows:

The experiment therefore provides no clear evidence that infants in the continuous condition perceived a determinate number of objects during familiarization, as predicted by the thesis that infants perceive object identity in accord with the continuity principle. (Spelke et al. 1995, 123)

To remedy this, the experimenters devised a second experiment, identical to the first one, except for the following aspects: the experimenters used a considerably greater number of infants (47 as compared to 24), the single rod paused when behind the screen so as to match the motions of the discontinuous movement, and the experimenters ensured that the noise produced by both experiments was equivalent (Spelke et al. 1995, 124). In this second study, the infants habituated to the continuously moving rod showed a more significant preference for the two-rod display (i.e. they looked significantly longer at it) than the test group; however, the preference of the discontinuously-habituated infants for the one-rod display actually decreased, thus contradicting the predictions of the experimenters.

Thus, Spelke et al. (1995, 135) conclude: “Experiments 1 and 2 provide evidence that 4-month-old infants’ perception of object identity over occlusion is reliably affected by information on the continuity or discontinuity of object motion.” This may be correct, but given that the infants that were habituated to the discontinuous event showed less interest in the display that should be surprising to them, namely the display of just one rod, it seems not entirely clear how exactly this perception is affected by said information.

3.2. *Infants’ Knowledge about the Behavior of Spelke Objects*

The most influential kind of habituation and preferential looking experiment that is taken to show that infants understand to some extent what can and what cannot happen to Spelke objects is the so-called drawbridge experiment introduced by Baillargeon, Spelke, and Wasserman (1985) and refined by Baillargeon (1987). Baillargeon (1987) describes her experiment – which

is mostly a more sophisticated repetition of the experiment conducted by Baillargeon, Spelke, and Wasserman (1985) with even younger infants (3.5 and 4.5 as opposed to 5 months old, see (Baillargeon 1987, 656/659)) – as follows. Infants are habituated to a screen (the ‘drawbridge’) rotating back and forth by 180 degrees on a plain. Then, infants are shown four test scenarios. The first two scenarios just show the drawbridge moving 180 degrees and 112 degrees respectively. Scenario three and four show the drawbridge moving towards the box, and then either stopping at the box or apparently going through the box. Scenarios 1-3 are called possible, while scenario 4 is called impossible. Again, the idea is that if the infants look longer at the impossible movement, it is because they understand that it is impossible, as it would otherwise be more familiar to them than the restricted movement. The outcome of these studies is that, overall, infants did indeed look significantly longer at the impossible scenario.

Here is how Maddy comments on a passage from Baillargeon, Spelke, and Wasserman (1985, 195–6):

Our reasoning was as follows. If infants understood that (1) the box continued to exist, in its same location, after it was occluded by the screen, and (2) the screen could not move through the space occupied by the box, then they should perceive the impossible event to be novel, surprising, or both. (Baillargeon, Spelke, and Wasserman 1985, 195–6)

[...]

In fact, the infants looked longer at the impossible event, evidence for their understanding of something like (1) and (2). (Maddy 2007, 248)

In this passage, Maddy clearly, even though with some hedging (“something like”), approves of the experimenter’s interpretation of the infants’ looking preferences: these preferences are explained with recourse to the infants’ understanding that it is not possible for the screen to move through the box.⁷ Baillargeon (1987, 661–63) addresses a further observation that is of considerable interest (not only because of one of the characters involved): those infants that had to be presented with the habituation display (rotation of the drawbridge of 180 degrees) many times to fulfill the habituation threshold did *not* show a preference for the impossible event – even if the box was replaced by a small figure of Mr. Potato Head. This hints at one of the studies that I will detail below (section 4.2).

To sum up, in this section, I have considered a representative – but not exhaustive – sample of the preferential looking experiments that Maddy enlists to support her claim that pre-linguistic infants reason according to rudimentary logical structure. The first experiment is taken to show that pre-linguistic infants individuate objects based on the criterion of common motion: objects that move continuously, even when occluded, are represented over time and space as one

⁷ Surprisingly, what Maddy quotes is not the summary of the results of the more sophisticated experiments conducted by Baillargeon (1987), but rather the research hypothesis of Baillargeon, Spelke, and Wasserman (1985). However, the summary given by (Baillargeon 1987, 663) of her results is essentially equivalent to the research hypothesis of Baillargeon, Spelke, and Wasserman (1985) quoted by Maddy: Baillargeon maintains that her results indicate that infants attribute permanence and solidity, or impenetrability, to Spelke objects, which allows them to suppose that objects cannot move through other objects. What is more interesting is that Baillargeon (1987, 661–63) addresses a phenomenon not mentioned by Maddy

persisting object. The second experiment is taken to show that pre-linguistic infants understand that objects persist when occluded and that they are not easily penetrable.

4. Preferential Looking Experiments Do Not Vindicate Nativism

In this section, I critically assess whether preferential looking experiments do indeed justify crediting pre-linguistic infants with the sophisticated cognitive abilities that nativists want to attribute to them. I first (section 4.1) highlight the high degree of sophistication of these cognitive capacities. After pointing out an empirical problem regarding attributing such sophisticated cognitive abilities to infants, I suggest that doing so needs to be supported by very strong and specific empirical evidence. This latter suggestion is based on a methodological principle introduced by Glock (2013).

Then (section 4.2), I show that the preferential looking experiments pioneered by Spelke are subject to strong experimental-methodical critique, critique that suggests that non-cognitive explanations rather than cognitive ones can explain the looking behavior of infants in such experiments. Finally (section 5), I conclude by suggesting that the evidence that can be provided by habituation and preferential looking experiments might be too unspecific in principle to warrant attribution of such high cognitive abilities to pre-linguistic infants.

4.1. *Nativism Needs Very Specific Evidence*

A common understanding of ‘understanding’, ‘believing’, ‘expecting’, etc. requires to be able to distinguish between what one understands, believes, and expects, and what actually is the case: a being that believes that *p* must understand that there is a reality that may contradict her belief. Similarly, to expect, and to be surprised, the infant needs to be able to think something like the following: *I* have expected such-and-such to occur, but what *actually* occurs is very different. Hence, given the fully intentional interpretations of the infants’ behavior in the preferential looking experiments, nativists like Spelke are committed to attribute to the infant the ability to make a subjective-objective distinction: the infant needs to be aware that her expectations, her concepts are hers, and that they may differ from what actually is the case. This ability is usually called epistemic self-consciousness.

A problem with attributing such an epistemic self-consciousness to infants is that it contradicts empirical research. Empirical work shows that infants do not develop the ability to distinguish between what they think is the case and what actually is the case, or simply the capacity to think that they may be wrong, until well after their acquisition of their first language. According to the current developmental accounts of self-consciousness, infants do not develop even a psychological – let alone an epistemic – self-consciousness until about 10 months of age (compare the excellent overview on (Bermúdez 1998, 271–75), with empirical details on the development of psychological self-consciousness in infants on (Bermúdez 1998, 229–65)). Bermúdez (1998, 69–71) argues against the attribution of concepts to infants along lines that are broadly analogous to the one suggested here.

Can nativists respond to this worry by delineating a different understanding of understanding, expecting, being surprised, etc. that does not require epistemic self-consciousness, and that is much less cognitively demanding in general? Such notions are proposed by Glock (2013, 214–

16) and Burge (2010, 267–76). Perhaps they could, but this is exactly not the choice of nativists. Spelke explicitly embraces the sophisticated cognitive structures that have to be in place for an infant to possess the Spelke object concept, and to understand that some things can and others cannot happen to such objects:

The object concept, on this view, arises from an intuitive theory of the physical world and its behavior. It is present, albeit in impoverished form, near the beginning of life. (Spelke, 2004, 333–34)

Spelke is not very explicit about what exactly she means by an “intuitive theory of the physical world and its behavior”, but it is evident that she is willing to bite the bullet and attribute not only a single object concept, but much more substantial cognitive structures to the infant. It is plausible that the notion of a theory of the physical world can be specified in such a way as to support the claim that the infant understands much of the physical behavior of objects and is indeed surprised if objects appear to violate this physical theory.

Note that it is not clear whether being able to individuate Spelke objects based on the three criteria mentioned can be separated from being able to understand some of the physical events that can and that cannot happen to these objects. This means that, for the present case, it may be that thinking of a certain instance of a Spelke object implies being able to think that certain things can or cannot happen to them. For instance, consider the nativist interpretation of the infants’ looking behavior in the rod experiment. Infants habituated to the continuous movement looked longer at the display of two rods because they had represented their habituation event as involving one rod only. In contrast, infants habituated to the discontinuous movement should have looked longer at the one-rod display (which they actually did not), because they had represented the habituation event as involving two rods. Can this explanation be separated from the following: infants understand that it is possible that one rod continuously moves across the entire stage, while they understand that it is not possible that one rod discontinuously moves across the entire stage?

Finally, note that the nativists’ reasoning takes the general form of an inference to the best explanation: there is a phenomenon to be explained, namely the looking preferences of pre-linguistic infants in preferential looking experiments. Nativists claim that attributing the ability to refer to Spelke objects, and a partial understanding of the things that can and cannot happen to these objects is the best, or the only explanation of this phenomenon.

As such, nativism would almost certainly be a target for the methodological principle called Morgan’s canon, according to which, for any behavioral pattern of a given animal, we should always adopt the explanation that attributes the least complex cognitive abilities to the respective animal (compare (Glock 2013, 147)).

Against Morgan’s canon, Glock (2013, 147–49) rightly maintains that there are other relevant explanatory virtues than the simplicity of cognitive abilities attributed. For instance, to explain a given behavioral pattern, attributing relatively simple cognitive capacities can be inferior to an explanation that attributes more complex cognitive capacities if the simpler capacities can only explain a very limited range of behavioral patterns – perhaps only the single behavioral pattern at issue. More generally speaking, explanatory unification and generality should be accepted as explanatory virtues along with the simplicity of the cognitive capacities attributed.

Hence, if it would turn out that the nativist explanation of the looking preferences of infants in the respective experiments can only be avoided by positing simpler cognitive capacities that are highly *ad hoc*, this would not threaten the status of the nativist explanation as the best explanation for the infants' behavior. I take this to be the essence of Glock's more modest alternative to Morgan's canon, called *Glock's handgun* (Glock 2013, 148).

However, even when aimed at merely with Glock's handgun, and not with Morgan's canon, it takes very substantial evidence for the nativist explanation to hold its ground: claiming that pre-linguistic infants are able to think like adults about parts of the physical world involves attributing to these infants about as sophisticated cognitive abilities as possible. Furthermore, according to Glock's handgun, the simplicity of the cognitive capacities used is not irrelevant – it is just not the only explanatory virtue that counts. Hence, while such an explanation is very likely to be pulverized by Morgan's canon, it may survive Glock's handgun if it offers a high degree of unification and generality, and if it is extremely well-supported by the empirical evidence. In the remainder of this article, I argue that habituation and preferential looking experiments cannot provide such evidence.

4.2. The Experiments Do Not Provide Sufficiently Specific Evidence

In this section, I critically assess the evidence that preferential looking experiments can provide for nativism. In the recent empirical literature, it was the interpretation of the infants' looking time in the drawbridge scenario that has received the greatest amount of critical consideration.

Rivera, Wakeley, and Langer (1999) provide an elaborate and empirically supported alternative to the nativist explanation for the looking behavior of infants in the drawbridge experiments. Remember that nativists infer identification of Spelke objects and understanding of physical possibility and impossibility from the fact that infants look longer at the so-called impossible event. This is the scenario where the drawbridge moves a full 180 degrees and thereby apparently, and impossibly, passes right through a box in its way.

Rivera, Wakeley, and Langer (1999, 433) hypothesized that the longer looking time may not be due to the infants' understanding that this event is impossible, but simply to their preference for more motion. Of course, Baillargeon (1987) tried to forestall this explanation by testing the looking time of a so-called control group of infants of the same age that did receive the same habituation, namely the repeated display of the "drawbridge" performing a 180 degrees movement, see Baillargeon (1987, 657–59). Then, the control group was shown merely two different movements of the "drawbridge", one of them 180 degrees, the other 112 degrees, without any boxes. It turned out that the infants in the control group did not look longer at the 180 degrees movement, which Baillargeon (1987, 659) took as evidence that it is not simply the infants' preference for more motion that accounts for their looking longer at the impossible event.

Rivera, Wakeley, and Langer (1999) were aware of this control group; however, they suspected that what made the infants in the control group look for roughly the same amount of time at both movements is that the *novelty* of the shorter movement – which was not shown to the infants during habituation – cancels out the effect that *more motion* would otherwise have on the infants'

looking time. Hence, these researchers suspected that it was an artifact of the habituation phase, causing infants to find the 180 degrees movement familiar and therefore less interesting, that infants did not display their preference for more motion in the control group.

Furthermore, Rivera, Wakeley, and Langer (1999) hypothesized, if it was indeed the case that the infants understand that objects are impenetrable, then there would be no habituation needed for these infants to prefer the impossible event. Therefore, the research group conducted a first experiment, which is analogous to the one by Baillargeon (1987), just without a control group and, importantly, without any habituation phase. It turned out that even without habituating, the infants looked longer at the so-called impossible scenario, where the drawbridge appears to go through the box, see Rivera, Wakeley, and Langer (1999, 430–31).

To confirm the suspicion that it was a combination of novelty and more motion that explains the infants' looking time rather than their finding one event impossible, the researchers conducted a second experiment, which is analogous to what Baillargeon (1987) conducted as the control condition experiment, again with the important difference that there is no habituation phase. In this experiment, infants were simply shown two displays, one of them showing the drawbridge moving 112 degrees, the other showing the drawbridge moving 180 degrees. Put briefly, Rivera, Wakeley, and Langer (1999, 432–33) report that the infants looked reliably longer at the 180 degrees motion than at the 112 degrees motion – without there being any box in the way. Hence, infants not previously habituated to any movements found both the possible and the impossible 180 degrees movement more interesting than the 112 degrees movement.

Based on this finding, Rivera, Wakeley, and Langer (1999) conclude that what explains the infants' looking time even in the case of the so-called impossible display is a simple preference for more motion. Furthermore, they suggest the results of the habituation test conducted by Baillargeon (1987) actually reinforces their point: as the 112 degrees movement was novel to the control group, it is remarkable that they did not show a pronounced preference for it; their finding more motion more interesting can explain this absence of a novelty-effect.

In addition to this critique of the drawbridge experiments, Parke and Gauvain (2002, 281–4), an introduction to developmental psychology, cite research by Bogartz, Shinskey, and Schilling (2000) that also casts doubts on the implications that nativists like Spelke draw from the drawbridge experiments. In this study, the researchers wanted to corroborate their hypothesis that familiarity and unfamiliarity can account for the prolonged looking times in habituation and preferential looking experiments in general, and in the drawbridge experiments from Baillargeon, Spelke, and Wasserman (1985) and Baillargeon (1987) in particular.

To verify their hypothesis, they conducted two experiments. The main methodological difference to Baillargeon, Spelke, and Wasserman (1985) and Baillargeon (1987) is that they split up the infants into three groups, each of which was then habituated to a different event (each event being shown three times during habituation phase, see (Bogartz, Shinskey, and Schilling 2000, 410)): one to the drawbridge moving 180 degrees without obstacle, one to the drawbridge moving 180 degrees and apparently going through the obstacle, and one to the drawbridge moving 120 degrees and stopping at the obstacle. Then, each group was tested by being shown each of the three displays four times.

Before considering the results, a remark on effects of habituation: the looking time, during habituation, of those infants habituated to the impossible event did not show any significant

difference to the looking time of the infants habituated to the two possible events (Bogartz, Shinskey, and Schilling 2000, 411). This should give pause to those who think that the main variable influencing infants' looking times is what they conceive, based on the properties they attribute to objects, as possible or impossible events. If infants really thought that objects are impenetrable, and hence, that it is impossible for one object to move through another, the infants that were shown the impossible event should have looked longer at it already during habituation. Note, furthermore, that the same study also brings bad news for non-cognitivist explanations of the drawbridge-experiment. Bogartz, Shinskey, and Schilling (2000, 411) note that they could not confirm (Rivera, Wakeley, and Langer 1999), who claim that infants simply find more motion more interesting. Based on the latter study, one would expect the infants in the former study to look reliably longer at those displays with more motion.

The results of the test trials, found on Bogartz, Shinskey, and Schilling (2000, 412), are as follows: In short, infants simply looked longer at those events that were new to them, given their habituation; those habituated to the 180 degrees movement – with or without box – found the possible event much more interesting than the impossible event. The only group that found the impossible event more interesting is the one that had been habituated to the 120 degrees movement with an obstacle.

In a second experiment, Bogartz, Shinskey, and Schilling (2000) further examined the effects of familiarity for the outcome of the drawbridge experiment. This time, they hypothesized that infants are not, as it were, content with being shown a display 3 times – they want to see it more often. The experiment proceeded as follows: They divided the infants up in four groups, which were each habituated to either the 120 or the 180 degrees movement and habituated 3 or 7 times respectively. So, during habituation, group 1 saw the 120 degrees movement 3 times, group 2 saw the 120 degrees movement 7 times, group 3 saw the 180 degrees movement 3 times, group 4 saw the 180 degrees movement 7 times. The four test displays then consisted of 180 or 120 degrees movements, with or without obstacle. In short, the results suggest that the 3 habituation events increased, rather than decreased, the respective infants' interest in the same displays, while seven habituation displays clearly decreased their interest in the respective displays. Again, whether the displays were possible or impossible from an adult's point of view was entirely insignificant (Bogartz, Shinskey, and Schilling 2000, 421–2).⁸

It is true that this is just one of several experiments that nativists enlist to support their case. However, the drawbridge experiments are by far the most famous instances of habituation and preferential looking experiments. Furthermore, the critiques of Rivera, Wakeley, and Langer (1999) and Bogartz, Shinskey, and Schilling (2000) focus on the impact of habituation on the results of the drawbridge experiments. As, by definition, all the habituation and preferential looking experiments include habituation, it is likely that versions of their critiques apply to other preferential looking experiments as well.

For instance, consider the rod experiments discussed above (section 3.1, page 8), where the behavior of the infants habituated to the discontinuous movement did not fit the nativist explanation: they showed less interest in the one-rod display than they should, given that they

⁸Baillargeon, Spelke, and Wasserman (1985) and Baillargeon (1987) report to have shown the infants the habituation display 6-9 times, depending on the individual decrease in looking time.

individuate objects based on continuous motion and hence should have represented the situation as involving two, not one rod. Perhaps, given the habituation, which involves two screens on the stage, the two-rod display presented infants with just the right mixture of novelty and familiarity to capture their attention (this is at least suggested by (Bogartz, Shinskey, and Schilling 2000), who found that, to a certain point, infants prefer displays that they have seen already)? Or perhaps infants have a basic preference for ‘more things’ on the stage, which caused them to look longer at the display of two rods than at the one of just one rod (which is supported by the behavior of the infants in the control group); however, they also prefer novel displays, which implies a looking preference for the display showing just one rod, and which, combined with the preference for more things, yields the looking behavior observed in the experiment.

Finally, note that, according to Tyler Burge’s philosophy of perceptual psychology, the behavior of the infants in these experiments should be explained by attributing to the visual system of the respective infants the capacity to produce certain perceptual constancies (Burge 2010, 438–50). In contrast, nativists explicitly favor a cognitive explanation of the looking times of infants.⁹

These three alternative explanations of the infants’ looking behavior in the rod experiment are all *prima facie* solid explanations with at least some empirical support. They are not *ad hoc*, and they allow for at least the same level of unification and generality as the nativist explanation. Hence, none of them is threatened by Glock’s handgun. None of them, however, dovetails with the nativist’s claims – and it seems very difficult to conceive of experimental set-ups using the preferential looking method that can eliminate these, and all the other conceivable alternative explanations. This is because the evidence that can possibly be provided by preferential looking experiments – the looking times of infants – seems too unspecific to warrant the strong conclusions that nativists want to draw from them.

This critical assessment of the preferential looking method is reinforced by the so-called *paradox of early permanence*: as Meltzoff and Moore (1998, 202–3) detail, the paradox consists in inconsistent empirical findings. On the one hand, preferential looking experiments are taken to show that 3-month old infants “understand the continued existence and movements of occluded objects” (Meltzoff and Moore 1998, 202). On the other hand, infants do not recover hidden objects, that is, they do not search for them, until about 8 months of age, suggesting that they do not understand that hidden objects persist until that age. It certainly supports critics of nativist interpretations of preferential looking experiments that the results obtained through such experiments are in tension with the results obtained by other experimental methods. Even more, these latter results are based on much more specific and sophisticated behavior than mere looking time.

5. Conclusion

In this article, I have examined a contemporary instance of the perennial dispute between empiricists and nativists, a dispute that centers around the question whether human infants should be credited with innate knowledge, concepts, or ideas.

⁹Another proposal to explain the behavior of infants in the rod experiment is due to Hatfield (2009, 241–48).

Hatfield only attributes reference to a bounded trackable volume, and not to instances of the kind Spelke object to infants.

I have first introduced the main characters of this contemporary debate. I have sketched Quine's empiricism, as he develops and defends it in his published writings. Then, I have discussed a passage in his posthumous work in which he explicitly abandons his empiricism in favor of nativism. Finally, I have introduced Maddy's reliance on the nativist arguments proposed by developmental psychologist Elisabeth Spelke.

Second, I have reconstructed a representative sample of the experiments that Maddy invokes in support of her nativist position. The first experiment aims to show that infants use the criterion of continuous motion to individuate Spelke objects. It aims to show this based on the looking times of infants at displays of one or two rods, after having been habituated to continuous and discontinuous rod movements respectively. The second kind of experiment, the so-called drawbridge experiment, is taken to show that infants understand that occluded objects persist and are not penetrable. This is inferred from the experimental datum that they are looking longer at the so-called impossible event, where the drawbridge appears to penetrate a box, or Mr. Potato head, than at various so-called possible events.

Third, I have critically assessed these experiments on two levels. First, I have outlined the high degree of sophistication of the cognitive abilities that nativists want to attribute to infants to explain their behavior in preferential looking experiments. I have suggested that being able to understand, to expect, and to be surprised in the nativist understanding of these terms requires an epistemic self-consciousness. However, empirical research suggests that pre-linguistic infants do not have such a consciousness. More generally speaking, I have pointed out that the nativists' inference is abductive in form: it infers nativism as the best explanation of the infant's looking preferences in the experiments. Based on the methodological principle called Glock's handgun, I have urged that it takes much in terms of empirical evidence to support this abductive inference.

On the second level, I have discussed empirical studies that imply that the habituation and preferential looking method probably cannot provide such empirical evidence. I have focused on the drawbridge experiments, as conducted by Baillargeon, et al. (1985) and Baillargeon (1987). I have argued that there are strong reasons, arising simply from paying close attention to the experimental set-up, against taking the results of these experiments to support nativism. First, Rivera, Wakeley, and Langer (1999) provide evidence that a combination of simple preferences of the infants, namely a preference for novelty and more motion, can explain the experimental results, once the side-effects of habituation (infants find the 180 degrees movement familiar and hence less interesting) are considered. Bogartz, Shinskey, and Schilling (2000) also focus on unintended effects of the habituation phase of the drawbridge experiments. They show that the most reliable way to predict the looking time of infants is simply to determine which event is new to them, given the habituation that they have been subjected to. In analogy to Rivera, Wakeley, and Langer (1999), they show that no cognitive explanation is needed to explain the infants' looking times in these drawbridge experiments (let alone a cognitive explanation of the nativist's caliber). Rather, simple preferences for novelty and more motion suffice, once the side-effects of habituation are taken into account.

I have then urged that these observations question the nativist's claim that her explanation is the best one for the looking preferences of infants in the respective experiments. There are good, non-*ad-hoc* explanations of the infants' looking preferences for both the rod and the drawbridge experiment that do not attribute cognitive capacities of the nativist's caliber to the infants. Hence,

it seems questionable why we should accept the nativist explanation. This holds even if we are not aiming Morgan's canon, but merely Glock's handgun at this nativist explanation. Furthermore, I have urged that it lies in the nature of habituation and preferential looking experiments that such alternative explanations are always possible.

This problem is reinforced by the fact that these nativist explanations themselves stand in tension with other empirical insights, such as the development of epistemic self-consciousness and the recalcitrant data from experiments based on search behavior of infants (leading to the so-called paradox of early permanence, see above, section 4.2).

These results do not support the conclusion that the dispute between nativists like Spelke and her critics, such as Rivera, has been decided in favor of the opponents of nativism. Rather, the results imply that the dispute is undecided, and probably undecidable on the basis of habituation and preferential looking experiments.

This verdict is fallible: it could be that researchers become so astute at devising new set-ups that the heavily cognitive explanation favored by nativists for the looking preferences of infants becomes the only possible one. For instance, it could be that experiments along the line of the rod experiment discussed above (section 3.1) establish beyond doubt that infants represent one single rod in the continuous display, but two rods in the discontinuous display. This requires, for instance, that the experimental data can exclude the explanation that (Burge 2010, 438–50) urges, according to which the respective infants' visual systems achieve a specific kind of perceptual constancy. Similarly, it is not entirely impossible that researchers develop experiments in the drawbridge-tradition that leave as the only acceptable explanation for the infants' looking times that they have some understanding of what is physically possible, and what is not physically possible.

However, as it were, I would not hold my breath until such evidence is available. It is just terribly difficult to develop experimental set-ups such that the looking time of infants in these experiments leaves only one explanation, namely that infants think like adults about parts of the physical world. The looking time of infants seems too unspecific a behavior to achieve that. Alternative, much simpler explanations such as preferences for novelty or more motion appear impossible to exclude with this experimental method.

This is of course not to say that the habituation and preferential looking method as such is useless. With regard to Quine's conception of perceptual similarity standards, the method could be put to use to chart these standards (compare Quine 1974, 17–18). Furthermore, it is clear that, given the conceptual and ontological shallowness of these similarity standards, they are not threatened by the worries concerning nativism that have been developed in this article. The doubts expressed in this article only concern the attempt to support nativist claims with these habituation and preferential looking experiments.

In short, the examinations of this article suggest that Quine can answer Maddy's nativist challenge. Quine does not have to be troubled by anti-empiricist, nativist claims as long as they are not supported by more solid evidence than the one currently, and probably possibly, provided by habituation and preferential looking experiments with infants. Maddy's *Second Philosophy of Logic*, in contrast, is questioned by these findings.

References

- Baillargeon, Renée. 1987. "Object Permanence in 3 1/2- and 4 1/2-Month-Old Infants." *Developmental Psychology* 23 (5): 655–64.
- Baillargeon, Renée, Elizabeth Spelke, and Stanley Wasserman. 1985. "Object Permanence in Five-Month-Old Infants." *Cognition* 20: 191–208.
- Bermúdez, José. 1998. *The Paradox of Self-Consciousness*. The MIT Press.
- Bogartz, Richard S., Jeanne L. Shinsky, and Thomas H. Schilling. 2000. "Object Permanence in Five-and-a-Half-Month-Old Infants?" *Infancy* 1 (4): 403–28.
- Burge, Tyler. 2010. *Origins of Objectivity*. Oxford University Press.
- Glock, Hans-Johann. 2012. "Thought, Judgment and Perception." *Grazer Philosophische Studien* 86. Rodopi: 207–21.
- . 2013. "Animal Minds: Philosophical and Scientific Aspects." In *A Wittgensteinian Perspective on the Use of Conceptual Analysis in Psychology*, 130–52. Palgrave MacMillan.
- Hatfield, Gary. 2009. "Getting Objects for Free (or Not): The Philosophy and Psychology of Object Perception." In *Perception and Cognition: Essays in the Philosophy of Psychology*, 212–55. Oxford University Press.
- Huntley-Fenner, Gavin, Susan Carey, and Andrea Solimando. 2002. "Objects Are Individuals but Stuff Doesn't Count: Perceived Rigidity and Cohesiveness Influence Infants' Representations of Small Groups of Discrete Entities." *Cognition* 85: 203–21.
- Maddy, Penelope. 2007. *Second Philosophy. A Naturalistic Method*. Oxford University Press.
- . 2014. "A Second Philosophy of Logic." In *The Metaphysics of Logic*, edited by Penelope Rush, 93–109. Cambridge University Press.
- Meltzoff, Andrew N., and Keith M. Moore. 1998. "Object Representation, Identity, and the Paradox of Early Permanence: Steps Toward a New Framework." *Infant Behavior and Development* 21 (2): 201–35.
- Oderberg, David S. 2005. "Predicate Logic and Bare Particulars." In *The Old New Logic*, edited by David S. Oderberg, 183–210. The MIT Press.
- Parke, Ross D., and Mary Gauvain. 2002. *Child Psychology. A Contemporary Viewpoint*. McGraw-Hill.
- Quine, Willard Van Orman. (1969a). "Epistemology Naturalized". In: *Ontological Relativity and Other Essays*. Columbia University Press, pp. 69–90.
- . 1974. *The Roots of Reference*. Open Court Publishing Co.
- . 2008. *Confessions of a Confirmed Extensionalist and Other Essays*. Ed. by Dagfinn Føllesdal and Douglas B. Quine. Harvard University Press.
- Rivera, Susan M., Ann Wakeley, and Jonas Langer. 1999. "The Drawbridge Phenomenon: Representational Reasoning or Perceptual Preference?" *Developmental Psychology* 35 (2): 427–35.
- Spelke, Elizabeth. 2004 "Where Perceiving Ends and Thinking Begins: The Apprehension of

Objects in Infancy.” In *Perception*, 321–36. Blackwell.

Spelke, Elizabeth, Grant Gutheil, and Gretchen Van de Walle. 1995. “The Development of Object Perception.” In *Visual Cognition. An Invitation to Cognitive Science*, 2nd edition, 2:297–330. The MIT Press.

Spelke, Elizabeth, Roberta Kestenbaum, Daniel J. Simons, and Debra Wein. 1995. “Spaciotemporal Continuity, Smoothness of Motion and Object Identity in Infancy.” *The British Journal of Developmental Psychology* 13: 113–42.