

Mysterious Quantifiers*

Lavinia Picollo

This is the penultimate version. Published version forthcoming in
the *Journal of Philosophy*.

How do the logical terms of our language acquire their meaning? In answering the analogous question for empirical terms, we can appeal to facts involving our causal interaction with the external world, which plausibly play a meaning-determining role. But for logical terms, involving a realm with which no direct causal interaction is possible, no such resources are available – we certainly do not stumble upon negation on our way to work.¹

The problem is compounded if, as I will assume, we attempt to explain how logical terms acquire their meanings against a broadly naturalistic background. Naturalism, as I understand it, rules out any appeal to mechanisms for which no scientific explanation can be envisioned, such as a ‘naturalness’ of certain properties that allows them to serve as ‘reference magnets’, cognitive powers that surpass those of Turing machines, or a faculty of ‘intuition’ that somehow connects us to a realm of non-causal facts.

A promising solution is to appeal to the rules of *use* of the logical terms in our language. The most widespread use-theoretic metasemantic view of logical terms is *logical inferentialism*, according to which the rules of use of logical terms take the form of inferences between statements (or axioms, or meta-inferences), which fully determine the meaning of such terms.² Logical inferentialism is a metasemantic view about how meaning is determined, and as such it is compatible with different views in the metaphysics of meaning. One, radical, brand

*I am deeply indebted to Daniel Waxman for his extensive and invaluable help, without which this work would not have been possible. I am also grateful to Gil Sagi and Volker Halbach for their inspiration, which motivated me to write this piece, and to Ethan Jerzak, Beau Madison Mount, Julien Murzi, Carlo Nicolai, Pedro del Valle-Inclán, José Zalabardo, the NUS Philosophy reading group, and an anonymous referee for their insightful and constructive comments on earlier drafts, which greatly improved this work.

¹The arguments of this paper apply equally well to mental as to linguistic content. However, for brevity, I will speak in linguistic terms only.

²One motivation for adopting logical inferentialism is the broadly naturalist metasemantic picture sketched above. But even if one does not share this motivation, and is willing to consider further metasemantic resources, it is worth investigating how much inferentialism can achieve.

Note also that the formulation of inferentialism offered here is somewhat narrower than some found in the literature, since it makes no room for the meaning of logical terms to be determined by language-entry and language-exit rules as in Sellars [1] and Brandom [2]. I will return to discuss this broader kind of inferentialism in the last section.

of logical inferentialism *identifies* the meanings of logical terms with the inference rules that govern them.³ By contrast, I am concerned with inferentialism combined with the dominant, truth-conditional view of meaning, roughly that specifying the meaning of a logical term involves specifying its contribution to the truth-conditions of sentences in which it figures.

The fundamental question for inferentialism, so understood, is: how can the inference rules governing an expression fix its meaning? The natural strategy appeals to a principle of charity that allows us to ‘read off’ truth-conditions from the rules. This involves the assumption that these rules are valid, from which the expression’s truth-conditions can be derived.

The tricky part is getting the *intended* meanings out of this strategy, partly because the background assumption of naturalism imposes significant constraints. On plausible assumptions about creatures like us, any rules we are in a position to follow involve only finitely many premises.⁴ Moreover, the class of (schematic) rules by which we are disposed to infer is computable – there is an algorithm to decide whether an inference belongs to the class. To aid perspicuity, rules are typically articulated in some formal language within a logical calculus – for instance, natural deduction calculi for first-order languages. These deductive systems are intended as – perhaps idealized – representations of our ordinary inferential practice. The constraints on our inferential practice mentioned above naturally carry over to the formal calculi intended to model it: in particular, they must be computable, in the sense that proofs in the calculus must be finite and the property of being a proof must be algorithmically checkable.

This leads to serious difficulties. Carnap [8] showed that ordinary calculi for first-order classical logic fail to exclude non-standard interpretations of negation, disjunction, the material conditional, and the universal and existential quantifiers. This issue is now known as “Carnap’s categoricity problem”.⁵ It not only affects classical logic but most non-classical logics too, including intuitionistic logic.

In order to solve the issue, several philosophers have proposed different calculi that arguably do succeed in fixing the meaning of all *propositional* connectives. One proposal, from Carnap himself, is to use a multiple-conclusion sequent calculus; another, advocated by Smiley [9] and Rumfitt [10], is to use bilateral systems; and a third proposal, pursued by Murzi [11], is to use a sequent calculus allowing single and empty conclusions. However, none of them improves

³Some exponents of this view are Putnam [3], Tennant [4] and, more recently, Button and Walsh [5] and Button [6].

⁴This is not entirely uncontroversial, however. Warren [7], for example, has argued to the contrary.

⁵An analogous – and perhaps more familiar – issue arises in mathematics (cf. Putnam [3]). Due to the compactness of first-order logic, no consistent first-order theory of arithmetic is categorical, in the sense that its axioms pick out a unique structure up to isomorphism. And, despite the fact that second-order theories can be categorical, they cannot be appealed to either, as the second-order consequence relation cannot be captured by a computable calculus. So we face a real challenge in explaining how arithmetical vocabulary gets its intended meaning. The problem is not confined to arithmetic; it is just as pressing for the language of set theory and other rich areas of mathematics.

the situation with the quantifiers: the usual truth-clauses for the quantifiers simply do not hold in every interpretation of the sequent calculus or bilateral systems. Carnap himself notoriously proposed to resolve this residual problem by incorporating infinitary rules into his sequent calculus; but this is hardly suitable from a naturalistic perspective, as we cannot plausibly follow infinitary rules, and moreover fails to deliver the desired truth-conditions, as we will see later. Still, the apparent availability of ways to explain the meaning of the propositional connectives might suggest that it is only a matter of time before the quantifiers succumb to similar strategies.

My aim in this paper is to argue that no such strategy will succeed. Attempts to use naturalistic inference rules to supply the traditional truth-conditions for quantifiers face a much more serious problem than has previously been supposed. In particular, I will argue that, against this background, there is no way to vindicate the thesis that our quantifiers satisfy the standard, Tarskian picture – what are usually called the “objectual” truth-conditions. The best we can achieve is an unfamiliar and radically different set of *sentential* truth-conditions, to adopt Garson’s [12] terminology. Roughly, the link between a generalization and its instances in a model is severely undermined. Using fixed-point techniques, I explore sentential quantification in detail, as well as the corresponding class of interpretations. If I am right, our quantifiers exhibit logical behaviour very different from objectual quantifiers. I go on to show that this has important implications in different areas of philosophy, including logic and metaphysics. Quantification is far more mysterious than many suppose.

1 Carnap’s Categoricity Problem

I begin by introducing Carnap’s categoricity problem for classical logic.⁶ First, I show how it affects the propositional connectives in a natural deduction calculus, and how it can be resolved by a multiple-conclusion sequent calculus. Second, I show how the problem unfolds for the quantifiers even in the sequent calculus setting.

Axioms and rules of inference confer meaning on an expression by constraining the possible interpretations of the language. Axioms must be *true* in every admissible interpretation and inference rules must be *valid* – that is, every admissible interpretation preserves truth from premises to conclusion.⁷

But what are the possible interpretations of the language? For first-order languages they are usually understood as models. A model is an ordered pair $\mathcal{M} = \langle \mathcal{D}_{\mathcal{M}}, \mathcal{I}_{\mathcal{M}} \rangle$, where $\mathcal{D}_{\mathcal{M}}$ is a non-empty set, the domain, and $\mathcal{I}_{\mathcal{M}}$ the interpretation function, that assigns semantic values built from $\mathcal{D}_{\mathcal{M}}$ to the non-logical expressions. The truth-value of each sentence is fully determined by the

⁶For simplicity and familiarity I focus exclusively on classical logic throughout. It is an interesting question how alternative logical traditions would be affected by my arguments, but I will not pursue the issue here.

⁷I take axioms to be sentences, and rules to be inferences between sentences, rather than formulae, except where explicitly discussed (cf. Section 2.2).

information provided by the domain and the interpretation function, together with the standard truth-conditions for logical terms, which are held constant across models.

However, for our purposes, such models are unsuitable, because they assume precisely what is in question – the interpretation of logical terms. In investigating how inference rules can fix the meaning of logical terms, we must also consider interpretations in which the truth-values of complex sentences do not obey standard truth-conditions – or any particular truth-conditions for that matter. So let us instead conceive of models as triples $\langle \mathcal{D}_{\mathcal{M}}, \mathcal{I}_{\mathcal{M}}, \mathcal{V}_{\mathcal{M}} \rangle$, where the valuation function $\mathcal{V}_{\mathcal{M}}$ assigns truth-values t or f to every complex sentence of the language *arbitrarily* – atomic sentences are evaluated as usual, as predication is not in question in this context. So, for instance, there are models in which every logically complex sentence is false, others in which all and only conjunctions are; and so on.

The question is then whether there is a computable calculus – a faithful (and perhaps idealized) representation of (perhaps a fragment of) our inferential practice – whose axioms and rules of inference pick out exactly the ‘standard’ models of the language, in which the logical constants have their intended meanings, i.e. obey standard truth-conditions.

1.1 The propositional connectives

Carnap [8] showed that, despite being sound and complete with respect to standard semantics, both natural deduction and Hilbert-style calculi for propositional logic fail to restrict the class of interpretations appropriately, as they are compatible with non-standard interpretations in which most connectives, including negation, disjunction, and the conditional – but, interestingly, excluding conjunction – have unintended meanings.

Consider the following introduction and elimination rules for negation and

conjunction in a natural deduction calculus (written in horizontal notation):⁸

$$\begin{array}{ll}
(\wedge I) & \phi, \psi \vdash \phi \wedge \psi \\
(\neg I) & \text{If } \Gamma, \phi \vdash \psi \text{ and } \Gamma, \phi \vdash \neg\psi, \text{ then } \Gamma \vdash \neg\phi \\
(\wedge E1) & \phi \wedge \psi \vdash \phi \\
(\neg E) & \neg\neg\phi \vdash \phi \\
(\wedge E2) & \phi \wedge \psi \vdash \psi
\end{array}$$

On the assumption that the inference rules for conjunction are valid, they yield the following conditions on the class of admissible interpretations:

$$\begin{array}{ll}
(\mathcal{V}\wedge I) & \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\phi) = \mathbf{t} \text{ and } \mathcal{V}_{\mathcal{M}}(\psi) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\phi \wedge \psi) = \mathbf{t}) \\
(\mathcal{V}\wedge E1) & \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\phi \wedge \psi) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\phi) = \mathbf{t}) \\
(\mathcal{V}\wedge E2) & \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\phi \wedge \psi) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\psi) = \mathbf{t})
\end{array}$$

Together, these conditions entail – and are entailed by – the standard truth-conditions for conjunction:

$$(\wedge) \quad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\phi \wedge \psi) = \mathbf{t} \Leftrightarrow \mathcal{V}_{\mathcal{M}}(\phi) = \mathbf{t} \text{ and } \mathcal{V}_{\mathcal{M}}(\psi) = \mathbf{t})$$

So, the natural deduction rules for conjunction successfully fix the intended meaning of conjunction, as they rule out all (and only) interpretations in which $(\wedge)$ does not hold.

By contrast, consider the standard truth-conditions for negation:

$$(\neg) \quad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\neg\phi) = \mathbf{t} \Leftrightarrow \mathcal{V}_{\mathcal{M}}(\phi) = \mathbf{f})$$

$(\neg I)$ is unlike the other rules so far considered in that it licenses an inference conditional on other inferences holding. $(\neg I)$ is not an ordinary inference but

⁸This choice of rules for negation is not uncommon but is somewhat non-standard from a proof-theoretic perspective. Following Gentzen, proof-theorists typically formulate the natural deduction rules for negation involving a *falsum* constant $\perp$ governed by:

$$(\perp E) \quad \perp \vdash \phi$$

The rules for negation then become:

$$\begin{array}{ll}
(\neg I^*) & \text{If } \Gamma, \phi \vdash \perp, \text{ then } \Gamma \vdash \neg\phi \\
(\neg E^*) & \text{If } \Gamma, \neg\phi \vdash \perp, \text{ then } \Gamma \vdash \phi
\end{array}$$

These rules arguably represent our inferential practice more faithfully. They also facilitate the distinction between classical logic, including the three rules, intuitionistic logic, in which $(\neg E^*)$ is dropped, and minimal logic, which also drops $(\perp E)$. In this context, however, it is more convenient to work with $(\neg I)$ and $(\neg E)$ instead, so to avoid the complications that the introduction of an extra logical constant to the language would bring about.

Regarding Carnap's categoricity problem, the traditional rules present similar difficulties to mine. Assuming that $\perp$ gets its intended interpretation, the rules manage to fix the standard interpretation of negation, but they face the directly analogous problem of how $\perp$ gets its intended interpretation.

An alternative approach in the literature, which originated in work by Tennant, is to understand *falsum* as “a kind of structural punctuation mark” indicating a ‘logical end’ rather than a propositional constant whose interpretation should be fixed by rules of inference (Tennant [13, p. 204]); formally, *falsum* is represented by an empty set of conclusions instead of $\perp$ (see also Rumfitt [10], Steinberger [14]). This strategy also has the advantage of solving Carnap's categoricity problem for negation (cf. Murzi [11]).

a *meta*-inference. How do meta-inferences constrain the class of admissible interpretations? Of course, they rule out interpretations in which they are not valid, but what does it mean for a meta-inference to be valid? We can answer this question by extrapolating from axioms and ordinary rules. Axioms – i.e. sentences – must be *true*, while ordinary inferences – i.e. inferences between sentences – must be *truth*-preserving, that is, *valid*. The natural next step is to say that meta-inferences – i.e. inferences between ordinary inferences – must be *validity*-preserving – a notion oftentimes referred to as “global validity”: if the ‘premises’ are valid, then so is the ‘conclusion’.⁹

If this is correct, we may read off the following conditions on admissible interpretations from $(\neg\text{I})$:

$$\begin{aligned} (\mathcal{V}\neg\text{I}) \quad & \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \text{ and } \mathcal{V}_{\mathcal{M}}(\phi) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\psi) = \mathbf{t}) \text{ and} \\ & \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \text{ and } \mathcal{V}_{\mathcal{M}}(\phi) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\neg\psi) = \mathbf{t}) \Rightarrow \\ & \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\neg\phi) = \mathbf{t}) \end{aligned}$$

where “ $\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t}$ ” is short for “for every $\phi \in \Gamma, \mathcal{V}_{\mathcal{M}}(\phi) = \mathbf{t}$ ”. $(\neg\text{E})$, in turn, restricts the class of interpretations by the following principle:

$$(\mathcal{V}\neg\text{E}) \quad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\neg\neg\phi) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\phi) = \mathbf{t})$$

By the soundness of the rules of the calculus, $(\mathcal{V}\neg\text{I})$ and $(\mathcal{V}\neg\text{E})$ follow from $(\neg)$. But the converse implication does not hold. Consider, for example, a model whose valuation function assigns $\mathbf{t}$ to every sentence. This model trivially satisfies both $(\mathcal{V}\neg\text{I})$ and $(\mathcal{V}\neg\text{E})$, as it makes their consequents (as well as their antecedents) true. Similarly, a model that maps every tautology to $\mathbf{t}$ and every other sentence to $\mathbf{f}$ is one that also satisfies both $(\mathcal{V}\neg\text{I})$ or $(\mathcal{V}\neg\text{E})$. But these models are clearly incompatible with $(\neg)$: the former assigns the same truth-value to every sentence and its negation, and the latter to every contingent sentence and its negation. This is, in a nutshell, Carnap’s categoricity problem for negation.

Contrary to what we might have expected, the inference rules fail to pin down the intended meaning of negation. Analogous results can be proven for disjunction and the conditional, and for other natural deduction proof systems, single-conclusion sequent calculi, and Hilbert-style calculi.¹⁰ However, we might still be able to establish the intended meanings of these logical expressions if our practice could be modelled by a stronger (but still computable) calculus in which all non-standard interpretations are excluded. And, indeed, this seems to be the case. Bilateral systems and sequent calculi which allow for multiple or zero conclusions, for example, appear to be successful in this respect.¹¹

⁹Alternative ways of construing meta-inference validity have been considered in the literature. See, for instance, Garson [12, Ch. 3 and 4]. In Section 2.2 I discuss – and dismiss – the most salient ones.

¹⁰Cf. Garson [12, §2.2].

¹¹Bilateral systems have been proposed as adequate formalizations of our actual inferential practice, e.g. by Smiley [9] and Rumfitt [10], and multiple-conclusion sequent calculi by Restall [15]. For a defence of sequent calculi allowing for only one or zero conclusions, see Tennant [13], Steinberger [14], and Murzi [11].

The right to appeal to some of these alternative proof systems in this context has been called into question.¹² However, I do not intend to pursue this issue here. I will assume that we can legitimately appeal to some kind of sequent calculus or to bilateral system to explain how the propositional connectives latch on to their intended meanings. I will focus on quantification, for which different, more serious issues arise. As will become clear in the next section, neither sequent calculi nor bilateral systems – or indeed *any* calculus – suffices to pin down the standard meaning of the universal and existential quantifiers.

Consider a multiple-conclusion sequent calculus for first-order logic.¹³ of bilateral systems and sequent calculi that allow only singletons or the empty set in the conclusions. Sequents are expressions of the form “ $\Gamma \vdash \Delta$ ”, where Γ and Δ are sets of sentences. Their intended reading is that *all* members of Γ entail *some* member of Δ . The rules of the sequent calculus allow us to infer a sequent from other sequents. Consider, for example, the sequent rules for negation:

$$\frac{\Gamma \vdash \phi, \Delta}{\Gamma, \neg\phi \vdash \Delta} (\neg\text{L}) \qquad \frac{\Gamma, \phi \vdash \Delta}{\Gamma \vdash \neg\phi, \Delta} (\neg\text{R})$$

Take $(\neg\text{L})$. According to this rule, if all members of Γ entail ϕ or some member of Δ , then all members of Γ together with $\neg\phi$ must entail some member of Δ . An argument with premises in Γ and conclusion ϕ is proof-theoretically valid just in case $\Gamma \vdash \phi$ is derivable in it.

Sequent rules are, like $(\neg\text{I})$ above, *meta*-rules – they licence the derivation of valid arguments from valid arguments. Thus, $(\neg\text{L})$ and $(\neg\text{R})$ impose the following conditions on admissible interpretations:

$$(\mathcal{V}\neg\text{L}) \quad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \text{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\phi) = \text{t} \text{ or } \mathcal{V}_{\mathcal{M}}(\Delta) \neq \text{f}) \Rightarrow \\ \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \text{t} \text{ and } \mathcal{V}_{\mathcal{M}}(\neg\phi) = \text{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\Delta) \neq \text{f})$$

$$(\mathcal{V}\neg\text{R}) \quad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \text{t} \text{ and } \mathcal{V}_{\mathcal{M}}(\phi) = \text{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\Delta) \neq \text{f}) \Rightarrow \\ \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \text{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\neg\phi) = \text{t} \text{ or } \mathcal{V}_{\mathcal{M}}(\Delta) \neq \text{f})$$

where “ $\mathcal{V}_{\mathcal{M}}(\Delta) \neq \text{f}$ ” is short for “it is not the case that, for every $\phi \in \Delta$, $\mathcal{V}_{\mathcal{M}}(\phi) = \text{f}$ ” – or, equivalently, “for some $\phi \in \Delta$, $\mathcal{V}_{\mathcal{M}}(\phi) = \text{t}$ ”.

The sequent rules for negation, unlike their natural deduction counterparts, do pin down the intended meaning of negation. Taken together, $(\mathcal{V}\neg\text{L})$ and $(\mathcal{V}\neg\text{R})$ not only follow from $(\neg)$ but also entail it.¹⁴ (Note that a model in

¹²See, for instance, Rumfitt [16], Steinberger [17], and Murzi & Hjortland [18] for objections to the adequacy of multiple-conclusion sequent calculi as formalizations of our actual inferential practice, and Button & Walsh [5, §13.6] for a critical assessment of bilateral systems.

¹³I work here with a multiple-conclusion sequent calculus, but everything of significance holds equally

¹⁴To see it, let $\mathcal{M}$ be a model and let Γ be the set of truths in $\mathcal{M}$ and Δ the set of falsities – $\Gamma = \{\psi : \mathcal{V}_{\mathcal{M}}(\psi) = \text{t}\}$ and $\Delta = \{\psi : \mathcal{V}_{\mathcal{M}}(\psi) = \text{f}\}$. Trivially, $\mathcal{V}_{\mathcal{M}}(\Gamma) = \text{t}$ and $\mathcal{V}_{\mathcal{M}}(\Delta) = \text{f}$. We also obtain the following:

$$(\text{A}) \quad \forall \mathcal{M}' (\mathcal{V}_{\mathcal{M}'}(\Gamma) = \text{t} \text{ and } \mathcal{V}_{\mathcal{M}'}(\Delta) = \text{f} \Rightarrow \mathcal{V}_{\mathcal{M}} = \mathcal{V}_{\mathcal{M}'})$$

which every sentence is true does not satisfy the consequents of $(\mathcal{V}\neg\text{L})$ and $(\mathcal{V}\neg\text{R})$, as Δ might be empty.) Analogous results can be shown to hold for the other propositional connectives.

1.2 The quantifiers

Neither sequent calculi nor bilateral systems fix the intended meaning of the first-order quantifiers. I illustrate this for the universal quantifier; parallel considerations extend to the existential quantifier in the obvious way.¹⁵

The standard sequent rules for the universal quantifier are the following:

$$\frac{\Gamma, \phi[t/v] \vdash \Delta}{\Gamma, \forall v \phi \vdash \Delta} (\forall\text{L}) \qquad \frac{\Gamma \vdash \phi[c/v], \Delta}{\Gamma \vdash \forall v \phi, \Delta} (\forall\text{R})$$

where t is any closed term of the language and, in $(\forall\text{R})$, c is any individual constant not occurring in Γ , Δ , or ϕ . $(\forall\text{L})$ allows us to conclude that Γ together with $\forall v \phi$ entails some member of Δ whenever Γ together with any instance of $\forall v \phi$ does so. $(\forall\text{R})$, in turn, allows us to infer that Γ entails either $\forall v \phi$ or some member of Δ from the premise that Γ entails either $\phi[c/v]$ or some member of Δ , for some arbitrary constant c .

Carnap noticed that, unlike the sequent rules for the propositional connectives, the rules for the quantifiers do not suffice to pin down their standard ‘objectual’ meanings.

Let a c -variant of $\mathcal{M}$ be a model whose domain, interpretation function, and valuation function are identical to $\mathcal{M}$ ’s with the possible exception of the semantic value of the individual constant c and the truth-values of sentences in which c occurs. The standard truth-conditions for the universal quantifier can be spelled out as follows:

$$(\text{O}\forall) \qquad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \text{t} \Leftrightarrow \forall \mathcal{M}' \sim_c \mathcal{M} \mathcal{V}_{\mathcal{M}'}(\phi[c/v]) = \text{t})$$

where c is any constant that does not occur in ϕ and $\sim_c$ says of two models that they are c -variants of each other. In other words, $\forall v \phi$ is true in a model just in case ϕ is true of every *object* in the domain of the model. The quantifier characterized by these truth-conditions is usually referred to as the “objectual” universal quantifier.

$(\forall\text{L})$ and $(\forall\text{R})$, on the other hand, impose the following respective conditions on the class of admissible interpretations of the language:

$$(\mathcal{V}\forall\text{L}) \qquad \begin{aligned} &\forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \text{t} \text{ and } \mathcal{V}_{\mathcal{M}}(\phi[t/v]) = \text{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\Delta) \neq \text{f}) \Rightarrow \\ &\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \text{t} \text{ and } \mathcal{V}_{\mathcal{M}}(\Delta) \neq \text{f} \end{aligned}$$

If $\mathcal{V}_{\mathcal{M}}(\neg\phi) = \text{t}$, then $\mathcal{M}$ is a counterexample to the consequent of $(\mathcal{V}\neg\text{L})$, so $\exists \mathcal{M}' (\mathcal{V}_{\mathcal{M}'}(\Gamma) = \text{t} \text{ and } \mathcal{V}_{\mathcal{M}'}(\phi) = \text{f} \text{ and } \mathcal{V}_{\mathcal{M}'}(\Delta) = \text{f})$. By (A), $\mathcal{V}_{\mathcal{M}}(\phi) = \text{f}$ too. Thus, we have established the left-to-right direction of $(\neg)$. The proof of the converse is symmetric.

¹⁵Interestingly, a similar result holds of the identity predicate, although I will not get into this issue in this paper.

$$\begin{aligned}
(\mathcal{V}\forall R) \quad & \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\phi[c/v]) = \mathbf{t} \text{ or } \mathcal{V}_{\mathcal{M}}(\Delta) \neq \mathbf{f}) \Rightarrow \\
& \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\forall v\phi) = \mathbf{t} \text{ or } \mathcal{V}_{\mathcal{M}}(\Delta) \neq \mathbf{f})
\end{aligned}$$

where Γ, Δ , and ϕ in $(\mathcal{V}\forall R)$ are c -free. Although, by the soundness of the calculus, $(\mathcal{V}\forall L)$ and $(\mathcal{V}\forall R)$ follow from $(O\forall)$,¹⁶ they do not entail it. $(\mathcal{V}\forall L)$ suffices to establish the left-to-right direction of $(O\forall)$,¹⁷ but the converse simply does not follow. This is because the sequent rules for the universal quantifier are compatible with interpretations in which a universal claim $\forall v\phi$ is false despite ϕ being true of every object in the domain (e.g. the rules allow an interpretation whose domain consists only of one object, the one-place predicate P is assigned the whole domain as its extension, but in which $\forall xPx$ is false), as will hopefully become clear in Section 3.2.¹⁸ Very roughly, the underlying reason why such non-standard interpretations are admitted by the rules is that the rules only spell out truth-conditions of sentences in terms of the truth-values of other sentences in any given model; they are insensitive to what lies in the domain of the model. This is admittedly a somewhat cryptic diagnosis, but it will be expanded and discussed at length in Section 4.1.

Similar results can be obtained for all usual systems, including bilateral, natural deduction, Hilbert-style, and other sequent calculi.¹⁹ Carnap's categoricity problem for the quantifiers appears to be insensitive to the choice of calculus. One possible reaction to this recalcitrant problem is to simply bite the bullet – we could resign ourselves to a language without objectual quantification. This is the option I ultimately favour; I explore it thoroughly in Sections 3 and 4.

¹⁶To derive $(\mathcal{V}\forall L)$ from $(O\forall)$, assume $(\mathcal{V}\forall L)$'s antecedent, $\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t}$, and $\mathcal{V}_{\mathcal{M}}(\forall v\phi) = \mathbf{f}$. Let c be an individual constant not occurring in ϕ or t , and let $\mathcal{M}' \sim_c \mathcal{M}$ be such that, for every c -free sentence ψ , $\mathcal{V}_{\mathcal{M}'}(\psi[c/t]) = \mathbf{t}$ just in case $\mathcal{V}_{\mathcal{M}'}(\psi) = \mathbf{t}$ – c -free predicates are true of c in $\mathcal{M}'$ if and only if they are also true of t . (We know that $(O\forall)$ does not rule out $\mathcal{M}'$ because, in that case, $\mathcal{M}$ would be automatically ruled out too.) By $(O\forall)$, $\mathcal{V}_{\mathcal{M}'}(\phi[c/v]) = \mathbf{t}$, so $\mathcal{V}_{\mathcal{M}'}(\phi[t/v]) = \mathbf{t}$. Thus, we have that $\mathcal{V}_{\mathcal{M}}(\phi[t/v]) = \mathbf{t}$ too. So, by $(\mathcal{V}\forall L)$'s antecedent and $\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t}$, it follows that $\mathcal{V}_{\mathcal{M}}(\Delta) \neq \mathbf{f}$, as desired.

To show that $(O\forall)$ entails $(\mathcal{V}\forall R)$, assume $(\mathcal{V}\forall R)$'s antecedent, $\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t}$, and $\mathcal{V}_{\mathcal{M}}(\Delta) = \mathbf{f}$. It follows that $\mathcal{V}_{\mathcal{M}}(\phi[c/v]) = \mathbf{t}$. What is more, since c does not occur in Γ or Δ , if $\mathcal{M}' \sim_c \mathcal{M}$, we also have that $\mathcal{V}_{\mathcal{M}'}(\Gamma) = \mathbf{t}$ and $\mathcal{V}_{\mathcal{M}'}(\Delta) = \mathbf{f}$. So, by $(\mathcal{V}\forall R)$'s antecedent, $\mathcal{V}_{\mathcal{M}'}(\phi[c/v]) = \mathbf{t}$. By $(O\forall)$, we have that $\mathcal{V}_{\mathcal{M}}(\forall v\phi) = \mathbf{t}$, as we wanted.

¹⁷To see it, assume that $\mathcal{V}_{\mathcal{M}}(\forall v\phi) = \mathbf{t}$ and let $\mathcal{M}' \sim_c \mathcal{M}$, where c does not occur in ϕ . Thus, $\mathcal{V}_{\mathcal{M}'}(\forall v\phi) = \mathbf{t}$ too. Letting $\mathcal{M}'$ play the role of $\mathcal{M}$, the proof-strategy of fn 14 delivers the desired result.

¹⁸The proof of this result is not trivial and requires some setup, which Section 3.2 provides. It might be instructive, though, to see why the proof-strategy used for negation in footnote 14 cannot work in this case. Since c may not occur in Γ or Δ in $(\mathcal{V}\forall R)$, these sets can no longer be identified with the set of *all* truths and the set of *all* falsities in $\mathcal{M}$. The best we can do is let Γ and Δ contain all c -free truths and all c -free falsities. But, then, instead of (A) we obtain a weaker principle:

$$(B) \quad \forall \mathcal{M}' (\mathcal{V}_{\mathcal{M}'}(\Gamma) = \mathbf{t} \text{ and } \mathcal{V}_{\mathcal{M}'}(\Delta) = \mathbf{f} \Rightarrow \mathcal{M} \equiv_c \mathcal{M}')$$

where $\equiv_c$ says of two models that they assign the same truth-values to all c -free sentences. (B), unlike (A), does not suffice to establish the right-to-left direction of $(O\forall)$, as $\mathcal{M} \equiv_c \mathcal{M}'$ is entailed by, but does not entail, $\mathcal{M} \sim_c \mathcal{M}'$ – two models could assign the same truth-values to all c -free sentences but, e.g. have different domains.

¹⁹See Carnap [8, §28] and Garson [12, Ch. 14].

Alternatively, one could dispute the philosophical significance of the results presented in this section. I consider three possible responses along these lines in the next section and find them wanting.

2 Three Resistance Attempts

2.1 Open-endedness

McGee [19] proposed a route to bypass Carnap's negative results presented in the previous section that appeals to the *open-endedness* of our inference rules. In this section I present McGee's argument and explain why it fails.

McGee's argument relies on the following assumptions:

- (**Language-existence**) For any class of models, there is a possible well-behaved expansion of our language in which there is a sentence true in just those models.
- (**Open-endedness**) The (schematic) rules that govern the logical constants are valid not only in our current language but also in any possible well-behaved expansion thereof.

We expand a language via the addition of new vocabulary. An expansion is *well-behaved* if the meaning of the old vocabulary remains the same, or at least the truth-values of the sentences of the old language are preserved.

The argument proceeds roughly as follows: by Language-existence, for each formula ϕ of our language there exists a sentence ψ in some possible well-behaved expansion $\mathcal{L}$ which is true in a model $\mathcal{M}$ just in case $\phi[c/v]$ is true in all c -variants of $\mathcal{M}$, where c is any constant that does not occur in ϕ . Thus, ψ expresses that everything, *in the objectual sense*, satisfies ϕ – in other words, objectual quantification is expressible in $\mathcal{L}$. By Open-endedness, $(\forall\forall\text{L})$ and $(\forall\forall\text{R})$ must hold of this language as well. The instance of $(\text{O}\forall)$ for ϕ then follows easily from an instance of $(\forall\forall\text{L})$ and $(\forall\forall\text{R})$ where Γ and Δ are sets of sentences of $\mathcal{L}$.²⁰ Since $\mathcal{L}$ is a well-behaved expansion of our language and the instance of $(\text{O}\forall)$ for ϕ holds in $\mathcal{L}$, it must also hold in our current language. It follows that *our* quantifiers obey objectual truth-conditions.

From a technical point of view, McGee's result is unimpeachable. However, its philosophical value is questionable. The problem is an ambiguity in the

²⁰Let ψ be true in $\mathcal{M}$ just in case $\mathcal{V}_{\mathcal{M}'}(\phi[c/v]) = \mathbf{t}$ for each $\mathcal{M}' \sim_c \mathcal{M}$. We can safely assume that c does not occur in ψ . Thus,

$$(C) \quad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\psi) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\phi[c/v]) = \mathbf{t})$$

By Open-endedness, we may let Γ be $\{\psi\}$ in $(\forall\forall\text{R})$. Letting Δ be empty, by (C), $(\forall\forall\text{R})$ entails

$$(D) \quad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\psi) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t})$$

Therefore, if $\mathcal{V}_{\mathcal{M}'}(\phi[c/v]) = \mathbf{t}$ for each $\mathcal{M}' \sim_c \mathcal{M}$, it follows that $\mathcal{V}_{\mathcal{M}}(\psi) = \mathbf{t}$ which, by (D), entails that $\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t}$.

crucial term “possible language”.²¹ McGee’s official view is that he is dealing with *mathematically possible* languages: “perhaps an ordered pair consisting of a collection of expression types, appropriately arrayed in grammatical categories, and a function assigning to each expression a semantic value. This is opposed to the conception of a language as a concrete system of human conventions and practices” (McGee [19, p. 62]). On such a conception, Language-existence is indeed plausible: mathematically possible languages are simply abstract objects of a certain kind.²² Open-endedness, however, is hard to defend. Presumably its justification is something like this: the meaning of an expression is determined by its rules of use, which reflect not only how the expression is *actually* used but the way in which the members of the linguistic community are *disposed* to use it.²³ Now we may in principle expand our language to include any new vocabulary (keeping all of the old language in place) and, in any such expansion, our dispositions to use logical terms would still hold; so we can safely assume that the schematic rules that reflect those dispositions extend to the new vocabulary. So there is nothing wrong with construing inference rules as open-ended in this context – many inferentialists have been sympathetic to such an approach. The difficulty is that this justification makes sense only when applied to languages that we are in principle able to speak or otherwise use; it is only in such languages that we can meaningfully be said to be disposed to draw inferences at all. What would need to be shown then is that the abstract languages that express objectual quantification are ones which we are able to speak; but McGee simply has not shown this.

If instead we conceive of a possible language as a language which could in principle be used, the argument fails on different grounds. On such a conception, the reasoning of the previous paragraph can be used to support Open-endedness. However, in these terms, Language-existence then amounts to the claim that *we can in principle speak* a well-behaved expansion of our language in which there is a sentence true in just any class of models. But this claim is question-begging in the present context. We have been asking whether there is some naturalistically acceptable metasemantic explanation of how creatures like us can speak a language whose quantifiers have objectual truth-conditions. Assuming Language-existence would then be tantamount to simply assuming that we can speak a language in which objectual quantification is expressible; but the entire question is how, if at all, this is possible.²⁴

²¹See Field [20] for a similar rebuttal of an analogous argument regarding mathematical vocabulary, also by McGee.

²²It is not totally obvious, however, that a sentence true in just any given class of models can always be added to our language *without altering the truth-values of any old sentence*, but I will not pursue this issue here.

²³See, for instance, Warren [21, §10.II] for an argument roughly along these lines.

²⁴Warren [21, Ch. 10] and Murzi and Topey [22] attempt to overcome the problems just mentioned by arguing that we *are* in fact able to speak the language expansions at issue in, e.g. “metaphysically possible worlds where we become part of a community that includes angels” who are capable of speaking the relevant languages and communicating them to us (Murzi and Topey [22, p. 3415]). This is an intriguing rejoinder; addressing it properly, however, would require getting into deep and tricky issues involving dispositions and their role in metasemantics. Let me just register here my scepticism about the relevance of nomologically

2.2 Local and hybrid validity

In Section 1.2 I pointed out that the sequent rules fail to fix the intended meanings of the quantifiers, as they are validated by a wide range of interpretations, including some which violate the objectual truth-conditions. However, this result relies on a particular understanding of meta-inference validity, namely, *global validity*. This account can be disputed, along with the negative results from Section 1.2, and has been so by Murzi and Topey [22]. Murzi and Topey put forward an alternative understanding of validity which they then deploy in their solution of Carnap’s categoricity problem.²⁵ In this section I consider Murzi and Topey’s account of meta-inference validity and explain why it ultimately fails.

According to global validity, for which I argued briefly in Section 1.1, a meta-inference is valid just in case (ground-level) validity is transmitted from premises to conclusion. More formally:

(Global validity) The inference from $\Gamma \vdash \Delta$ to $\Sigma \vdash \Xi$ is valid just in case:²⁶

$$\forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\Delta) \neq \mathbf{f}) \Rightarrow \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\Sigma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\Xi) \neq \mathbf{f})$$

There are, however, alternative ways of understanding meta-inference validity which give rise to alternative conditions on the class of admissible interpretations of the language. For instance, according to the *local* account, a

impossible scenarios in constraining meaning.

One concern is whether Open-endedness really covers scenarios where we commune with supernatural creatures. The basic rationale is that dispositions, not just actual use, determine meaning. But it is far from clear that our dispositions manifest in nomologically impossible worlds. (Consider the analogy of a mirror with the disposition to shatter on being struck: in a situation in which the laws of nature are radically different so that mirrors spontaneously shatter, it is unclear that the mirror’s shattering is a manifestation of the disposition it has in our world.) What is more, dispositional properties are often closely tied to causation and natural laws (cf. Armstrong [23], Lewis [24], Bird [25]); as Warren argues, the relevant modality is “physical or nomological, rather than metaphysical” (Warren [21, p. 34]). Even if our inferential dispositions *do* manifest in nomologically impossible scenarios, it simply does not follow that our behaviour there is *relevant* to determining meaning. (Just as the shattering of a mirror on being struck in physically impossible scenarios isn’t obviously relevant to determine whether it is fragile.) The details of how to fashion a precise account of meaning out of dispositions are notoriously difficult, but my suspicion is that a broadly naturalistic solution (cf. Warren [26]) will appeal primarily to physically possible (and perhaps even feasible) scenarios.

²⁵Murzi and Topey’s proposal draws upon influential work by Bonnay and Westerståhl. In [27] Bonnay and Westerståhl propose a solution to Carnap’s categoricity problem for first-order logic that relies on three semantic assumptions: non-triviality (there is at least one false statement in every model), compositionality (the semantic value of a compound expression is determined by the semantic values of its immediate constituents and the mode of composition), and topic-neutrality (the logical constants are permutation invariant). But, as Murzi and Topey rightly point out, these assumptions do not seem to be available in this context, as the only principles we are allowed to use are those which can be justified exclusively by reference to our inferential dispositions. Murzi and Topey’s proposal, while inspired by Bonnay and Westerståhl’s, avoids their semantic assumptions.

²⁶I formulate global validity and other definitions of validity just for meta-inferences with one premise, but it should be clear how they generalize to multiple premises.

meta-inference is valid just in case, in any given interpretation, if the premises are truth-preserving, so is the conclusion. More formally:

(Local validity) The inference from $\Gamma \vdash \Delta$ to $\Sigma \vdash \Xi$ is valid just in case:

$$\forall \mathcal{M} ((\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\Delta) \neq \mathbf{f}) \Rightarrow (\mathcal{V}_{\mathcal{M}}(\Sigma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\Xi) \neq \mathbf{f}))$$

The key difference between global and local validity lies in the scope of the quantifiers ranging over models. Unlike global validity, local validity does not require that meta-inferences preserve ground-level validity but just truth-preservation.

According to Murzi and Topey, there are broadly inferentialist grounds to prefer the local over the global account. In practice, they argue, we not only infer statements from statements, but also arguments from arguments. If this is right, then local validity looks better suited to capturing the import of a rule, as all inferences from an invalid argument are globally but – not locally – valid.

In the context of Carnap’s categoricity problem, there is another tempting reason to understand validity as local rather than global: on a local reading, the sequent (as well as the standard natural deduction) rules for the universal quantifier do entail its objectual truth-conditions, as expressed by (O $\forall$).²⁷

However, something is fundamentally wrong with the local account of validity – Garson [12, §3.5] describes it as an “incompleteness” problem. Very roughly, the local readings of certain meta-inferences entail the semantic validity of arguments that are not derivable by the rules themselves. In other words, local readings in some cases ascribe greater meaning-conferring powers to meta-inferences than they actually have.

A first example concerns the material conditional. Some valid principles involving only this connective – e.g. Peirce’s Law – are not derivable using just the introduction and elimination rules for it – i.e. conditional proof and *modus ponens* – in a standard natural deduction calculus.²⁸ Their proof requires also the rule of *reductio ad absurdum*. Yet, the local readings of the conditional rules entail *by themselves* the validity of Peirce’s Law.²⁹ But this is unacceptable from an inferentialist point of view; there is in principle nothing incoherent

²⁷To see it, note that the local readings of ($\forall$ L) and ($\forall$ R) are:

$$\begin{aligned} (\mathcal{V}_l \forall L) \quad & \forall \mathcal{M} ((\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \text{ and } \mathcal{V}_{\mathcal{M}}(\phi[t/v]) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\Delta) \neq \mathbf{f}) \Rightarrow \\ & (\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \text{ and } \mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\Delta) \neq \mathbf{f})) \end{aligned}$$

$$\begin{aligned} (\mathcal{V}_l \forall R) \quad & \forall \mathcal{M} ((\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\phi[c/v]) = \mathbf{t} \text{ or } \mathcal{V}_{\mathcal{M}}(\Delta) \neq \mathbf{f}) \Rightarrow \\ & (\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t} \text{ or } \mathcal{V}_{\mathcal{M}}(\Delta) \neq \mathbf{f})) \end{aligned}$$

where Γ, Δ , and ϕ in ($\mathcal{V}_l \forall R$) are *c*-free. Since ($\mathcal{V} \forall L$) entails the left-to-right direction of (O $\forall$) (cf. fn 17) and follows logically from ($\mathcal{V}_l \forall L$), ($\mathcal{V}_l \forall L$) also entails (O $\forall$)’s left-to-right direction. For the right-to-left direction, assume that $\forall \mathcal{M}' \sim_c \mathcal{M} \mathcal{V}_{\mathcal{M}'}(\phi[c/v]) = \mathbf{t}$, where *c* does not occur in ϕ . It follows that $\mathcal{V}_{\mathcal{M}}(\phi[c/v]) = \mathbf{t}$. Now let Γ be $\{\psi : \mathcal{V}_{\mathcal{M}}(\psi) = \mathbf{t} \text{ and } \psi \text{ is } c\text{-free}\}$ and Δ be $\{\psi : \mathcal{V}_{\mathcal{M}}(\psi) = \mathbf{f} \text{ and } \psi \text{ is } c\text{-free}\}$ in ($\mathcal{V}_l \forall R$). It follows that, if $\mathcal{V}_{\mathcal{M}}(\phi[c/v]) = \mathbf{t}$, then $\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t}$. Therefore, $\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t}$.

²⁸Peirce’s Law is the schematic principle $((\phi \rightarrow \psi) \rightarrow \phi) \rightarrow \phi$.

²⁹See Garson [12, §1.8].

about a linguistic community that infers according to the standard rules for the conditional but refuses to endorse Peirce's Law or the *reductio* rule.

Another, even more devastating counterexample to local validity is that local readings of $(\forall L)$ and $(\forall R)$ not only entail $(O\forall)$ but also:³⁰

$$(*) \quad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\phi[c/v]) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t})$$

This principle says that a universal claim is true whenever *any* single instance of it is true. The leap from an instance to its universal generalization is not only an inference that the quantifier rules by themselves do not license; it is patently unsound. Local validity assigns to the universal quantifier a grossly unintended meaning.

Nevertheless, Murzi and Topey argue that the local account can be salvaged. They show that, relative to a particular sequent calculus developed in Murzi [11], there's a version of local meta-inference validity for which incompleteness does not arise.

One feature of Murzi's calculus is that, following Schroeder-Heister [28], it contains higher-order rules – allowing the assumption and discharge of rules – to model inferences between arguments. Another salient feature is that sequents must have *at most* one conclusion, and possibly none, where empty conclusion sets are conceived as a formalization of *falsum* (cf. fn 8). These two features allow the formulation of *reductio ad absurdum* as a purely structural rule (i.e. in which no logical terms occur). Murzi and Topey go on to show that Peirce's Law is derivable in their calculus using only the rules for the conditional and structural rules. What is more, as they point out, the result generalizes to all connectives, thus eliminating all sources of incompleteness for the propositional case: any propositionally valid principle can be derived using only the operational rules for the connectives involved in it plus structural rules.

To get rid of the incompleteness problem affecting the quantifiers, Murzi and Topey allow open formulae to occur in sequents, and replace $(\forall R)$ with:

$$\frac{\Gamma \vdash \phi}{\Gamma \vdash \forall v \phi} (\forall R')$$

which features a variable v instead of a constant. Naturally, v cannot occur free in Γ .

Since inference rules may now involve open formulae, variable assignments must be brought into the picture on the semantic side. An interpretation will now be given by a pair $\mathcal{M}, g$ consisting of a model $\mathcal{M}$ and an variable assignment g over $\mathcal{M}$; valuations assign truth-values relative to the variable assignment. Truth-conditions must then be relative to variable assignments too.

Finally, variable assignments must be incorporated into the definition of local meta-inference validity. There are at least two possibilities: to go 'fully local', and say that a meta-inference is valid just in case truth-preservation in every model and assignment is preserved from premises to conclusion, or to adopt a

³⁰To see it, let both Γ and Δ in $(\forall_i \forall R)$ be empty, so $\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t}$ and $\mathcal{V}_{\mathcal{M}}(\Delta) = \mathbf{f}$.

hybrid approach, according to which a meta-inference is valid just in case, in every model, the conclusion is truth-preserving under every assignment whenever the premises are. To avoid the incompleteness problem for the quantifiers, Murzi and Topey take the hybrid approach. Formally:

(Hybrid validity) The inference from $\Gamma \vdash \Delta$ to $\Sigma \vdash \Xi$ is valid just in case:

$$\forall \mathcal{M} (\forall g (\mathcal{V}_{\mathcal{M},g}(\Gamma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M},g}(\Delta) \neq \mathbf{f}) \Rightarrow \forall g (\mathcal{V}_{\mathcal{M},g}(\Sigma) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M},g}(\Xi) \neq \mathbf{f}))$$

While the quantifiers over models take a wide scope, as in local validity, the quantifiers over assignments take a narrow scope, as in global validity.

Hybrid validity might seem promising. As Murzi and Topey point out, while the local reading of $(\forall R')$ and the hybrid reading of $(\forall R)$ both entail the unsound principle $(*)$, the hybrid reading of $(\forall R')$ does not – incompleteness does not arise for the quantifiers either in Murzi and Topey’s setup. What is more, the hybrid readings of the quantifier rules play an important role in Murzi and Topey’s argument that our quantifiers obey objectual truth-conditions (cf. Murzi and Topey [22, §2.4]).

Note, however, that whether the hybrid account delivers the desired results is highly sensitive to the details of the calculus under consideration. For example, Murzi and Topey’s proposed solution to the incompleteness problem depends crucially on the fact that they are using a calculus where *falsum* is formalized by the empty set and where *reductio* is correspondingly formulated as a higher-order structural rule. Perhaps these features of the calculus can be justified. Murzi and Topey argue that higher-order rules faithfully formalize our dispositions to infer arguments from arguments, and cite interesting work by proof theorists who argue that sequents with empty conclusions are a better formalization of our practices of reasoning with contradictions.³¹ But the opposite perspective is also arguable; for example, a prominent view (to which I am sympathetic) is that inference is fundamentally best understood as a transition between contentful mental states, and so it makes little sense to speak of an inference from an argument to another argument.³² (Of course one might infer from *the judgement that* one argument is valid to *the judgement that* another is, but that provides no reason to countenance higher-order rules).

But even if it is conceded that a calculus along Murzi and Topey’s lines is the only reasonable way to formalize the propositional fragment of our inferential practice, the situation with the quantifiers seems much worse. Their preferred hybrid account of validity gives entirely unacceptable results when applied to $(\forall R)$, as hybrid validity collapses into local validity if only sentences are involved. Their response is to replace $(\forall R)$ with $(\forall R')$, a rule for the universal quantifier that allows open sentences to figure in sequents and lets variables play the role of arbitrary witnesses. But this move seems at best *ad hoc* and at worst unfaithful to our inferential practice.

³¹Cf. Tennant [13], Steinberger [14]. See also fn 8.

³²See, for instance, Boghossian [29], Wright [30], and Warren [31].

First, except perhaps to ease presentation, there seem to be no good antecedent reasons to *prefer* the use of variables over constants as arbitrary witnesses in proofs; to the extent $(\forall R')$ is plausible, it is because it seems like a mere notational variant of the more usual $(\forall R)$. What is more, allowing open sentences into sequents gives rise to an additional choice point – whether to formulate validity using what I have called “fully local” or “hybrid” readings – with little principled reason for taking the route that Murzi and Topey adopt. And, finally, there are strong inferentialist reasons for thinking that calculi which deal with open formulae do not faithfully model our inferential behaviour. As many in the inferentialist tradition have emphasized, “whole declarative sentences [are] prior in the order of explanation” because they are the vehicles of contentful mental states like assertion and belief (Brandom [32, p. 12]). It is difficult to see how open sentences can play a similar role.³³

So there is a substantial contrast between the hybrid account of validity and the global account, which I favour. The hybrid account not only appears to be *ad hoc* – its sole motivation here has been the role it plays in securing the intended truth-conditions for the quantifiers – but gives rise to devastating problems unless supplemented with what also appear to be *ad hoc* and unprincipled conditions on the calculi to which it is applied. By contrast, the global account enjoys initial plausibility and can apply to any calculus whatsoever, with no such restrictions. The wrinkle, of course, is that it does not help establish the objectual clauses for the quantifiers.

Perhaps it is time to concede that no inference rules of an ordinary calculus suffice to fix objectual truth-conditions for the quantifiers. Might an extraordinary calculus do the work? Carnap himself made an interesting proposal along these lines, to which I turn next.

2.3 Infinitary rules

Faced with the inadequacy of the standard sequent rules, Carnap [8, §31] proposed to strengthen the calculus by the addition of a new, ‘transfinite’ inference rule, which he argues is necessary to confer upon the universal quantifier its objectual meaning:

$$\frac{\{\Gamma \vdash \phi[t/v], \Delta : t \text{ is a closed term}\}}{\Gamma \vdash \forall v \phi, \Delta} (\forall R^\infty)$$

$(\forall R^\infty)$ has infinitely many sequent premises, one for each closed term of the language. Informally, it says that, if each instance of $\forall v \phi$ (or a member of Δ) follows from Γ , then so does $\forall v \phi$ – or that all the instances, put together, entail the universal claim. In arithmetical contexts, $(\forall R^\infty)$ is also known as the ω -rule.

³³Although an argument could be made that free variables are needed in inferences to model the use of anaphoric pronouns in natural language. Thanks to an anonymous referee for pointing this out.

$(\forall R^\infty)$ restricts the class of admissible interpretations of the language by the following principle:

$$(\mathcal{V}\forall R^\infty) \quad \forall t \forall \mathcal{M} (\mathcal{V}_\mathcal{M}(\Gamma) = \mathbf{t} \Rightarrow \mathcal{V}_\mathcal{M}(\phi[t/v]) = \mathbf{t} \text{ or } \mathcal{V}_\mathcal{M}(\Delta) \neq \mathbf{f}) \Rightarrow \\ \forall \mathcal{M} (\mathcal{V}_\mathcal{M}(\Gamma) = \mathbf{t} \Rightarrow \mathcal{V}_\mathcal{M}(\forall v \phi) = \mathbf{t} \text{ or } \mathcal{V}_\mathcal{M}(\Delta) \neq \mathbf{f})$$

where t ranges over closed terms. Taken together, $(\mathcal{V}\forall L)$ and $(\mathcal{V}\forall R^\infty)$ are provably equivalent to the following truth-conditions for universal statements:

$$(\text{Sub}\forall) \quad \forall \mathcal{M} (\mathcal{V}_\mathcal{M}(\forall v \phi) = \mathbf{t} \Leftrightarrow \forall t \mathcal{V}_\mathcal{M}(\phi[t/v]) = \mathbf{t})$$

A quantifier whose meaning is given by $(\text{Sub}\forall)$ is usually referred to as a “substitutional” universal quantifier.

$(\text{Sub}\forall)$ states that $\forall v \phi$ is true in a model just in case ϕ is true of every *named* object in the domain of the model. By contrast, objectual truth-conditions for the universal quantifier require that ϕ is true of *all* objects in the domain, not just the named ones. It follows that Carnap’s infinitary rule is unsound with respect to the intended semantics. To bypass this issue, Carnap focuses exclusively on interpretations of the language in which every object in the domain is denoted by a term. Indeed, if we let $\mathcal{M}$ range exclusively over such models, $(O\forall)$ and $(\text{Sub}\forall)$ become equivalent.

Carnap’s strategy is not particularly attractive for the naturalistically minded. One issue is how an exclusive focus on interpretations in which every object is named may be justified by appealing only to rules of use. Another is that, even if such a justification were possible, focusing exclusively on those interpretations seems highly undesirable, as it would deem our language useless for theorizing about uncountably many objects, such as the real numbers or sets. This is because, from a naturalistic standpoint, the languages finite creatures like us are able to speak must be computable, in the sense that there’s an algorithm for deciding what counts as a well-formed expression of the language. But no such algorithm is possible for a language with uncountably many expressions. But even putting these issues aside, a naturalistic metasemantics simply cannot appeal to an infinitary rule to explain how meanings in our language are fixed, because any calculus enabling infinitary inferences will not be computable.

We can extract two important lessons from Carnap’s attempted solution. First, note that $(\forall L)$ and $(\forall R^\infty)$ articulate an inference pattern typically associated with *infinite conjunctions*, a pattern that limited beings like us could not follow. Just as infinite conjunctions entail and are entailed by their conjuncts, $(\forall L)$ and $(\forall R^\infty)$ let universal claims entail and be entailed by their instances. Accordingly, $(\text{Sub}\forall)$ spells out truth-conditions for universal claims that match those of an infinite conjunction – a universal claim is true exactly when its instances are, just as an infinite conjunction is true exactly when its conjuncts are true. Substitutional universal quantification is just a particular kind of infinite conjunction – one whose conjuncts differ only by a closed term. Our language has no room for either.

Moreover, what distinguishes the objectual from the substitutional quantifier is not a fundamental difference in their respective truth-conditions but rather a

somewhat accidental feature of the language in which they are embedded, that is, whether every object in the domain has a name.³⁴ The only reason why $(\forall R^\infty)$ fails to confer the universal quantifier its objectual meaning is that some objects might lack names, so some ‘instances’ of a universal claim might not be expressible. In certain unfortunate cases, some of these ‘instances’ might have no linguistic correlate. Nonetheless, the objectual quantifier reduces the truth of a universal claim to the truth of its ‘instances’ just as much as its substitutional cousin. Under the right circumstances, when every object has been named, the objectual quantifier will collapse into the substitutional quantifier, exhibiting inferential patterns characteristic of an infinite conjunction too. We should be as suspicious of objectual quantification as we are of substitutional quantification.

Second, and more importantly, note that, although Carnap’s hunch that finitary rules cannot fix the objectual meaning of the quantifiers seems to be on the right track, infinitary rules do not seem to do a satisfactory job either, unless we controversially assume that every object in the domain of a model has a name. The issue regarding the inability of our inference rules to pin down objectual meanings for the quantifiers runs deeper, as it will become clear in the following section. As long as the meaning-determining rules of our language only take the form of inferences between sentences (or formulae, for that matter), expressing just *intralinguistic* relations, they will only link the truth-conditions of statements to the truth-conditions of other statements. As such, we cannot expect them to explain in a fine-grained manner how our quantifiers would interact with *extralinguistic* objects in the domain of an interpretation. As Garson put it, “it would be a lot to ask of rules governing wffs that they require the presence of an objectual domain in their models related to the quantifier truth condition in the right way” (Garson [12, p. 216]). Carnap’s infinitary proposal focuses only on inferences between statements and is, thus, as bound to fail as the finitary strategies we considered before.

3 Sentential Quantification

I hope to have convinced you that neither the substitutional nor the objectual quantifier can be a part of our language. Does this mean that our language simply lacks quantification? I do not think so. The quantifier-words that we use every day are obviously meaningful; ordinary quantification *is* whatever these words mean. Their truth-conditions, however, are not objectual or substitutional, as I have shown, but *sentential*. (The motivation behind this nomenclature will become clear soon.)

3.1 Sentential truth-conditions

Sentential quantifiers are precisely those whose meanings are determined exclusively by the usual inference rules. The truth-conditions for the universal

³⁴Occasionally some objects will lack a name due to mathematical impossibility (e.g. if there are uncountably many), but not always.

sentential quantifier, for instance, consist just of the (global) readings of $(\forall L)$ and $(\forall R)$, that is, the conjunction of $(\mathcal{V}\forall L)$ and $(\mathcal{V}\forall R)$, which can be expressed succinctly by the following clause:³⁵

$$(S\forall) \quad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t} \Leftrightarrow \forall \mathcal{M}' \equiv_c \mathcal{M} \forall t \mathcal{V}_{\mathcal{M}'}(\phi[t/v]) = \mathbf{t})$$

where c is any individual constant not occurring in ϕ and $\equiv_c$ is a relation (different from $\sim_c$; cf. fn 18) that holds between two models whenever they assign the same truth-values to all c -free sentences. Dual truth-conditions characterize the sentential existential quantifier.

According to $(S\forall)$, a universal claim is true in a model $\mathcal{M}$ just in case its instances are true in every ‘nearby’ interpretation, that is, in every interpretation that assigns the same truth-values as $\mathcal{M}$ to all c -free sentences, for any appropriate constant c . Thus, $(S\forall)$ spells out truth-conditions for universal claims exclusively in terms of the truth-values of other *sentences*, without any mention of extralinguistic objects.³⁶ It is in this sense that these truth-conditions are ‘sentential’, as opposed to ‘objectual’.

Consider the relationship between a sentential universal generalization and its instances. It follows from $(S\forall)$ that, if some instance $\phi[t/v]$ is false in $\mathcal{M}$, $\forall v \phi$ is false too. In fact, $\forall v \phi$ is false in $\mathcal{M}$ whenever there is an object in $\mathcal{D}_{\mathcal{M}}$ of which ϕ is not true, regardless of whether it has a name in the model. $(S\forall)$

³⁵To see it, first assume $(\mathcal{V}\forall L)$ and $(\mathcal{V}\forall R)$. Following the same proof-strategy as in fn 14, $(\mathcal{V}\forall L)$ can be easily shown to entail:

$$(E) \quad \forall \mathcal{M} (\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t} \Rightarrow \forall t \mathcal{V}_{\mathcal{M}}(\phi[t/v]) = \mathbf{t})$$

Assume that $\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t}$ and $\mathcal{M}' \equiv_c \mathcal{M}$, where c does not occur in ϕ . It follows that $\mathcal{V}_{\mathcal{M}'}(\forall v \phi) = \mathbf{t}$ too. By (E) , $\forall t \mathcal{V}_{\mathcal{M}'}(\phi[t/v]) = \mathbf{t}$. This establishes the left-to-right direction of $(S\forall)$. For the converse, assume $\forall \mathcal{M}' \equiv_c \mathcal{M} \forall t \mathcal{V}_{\mathcal{M}'}(\phi[t/v]) = \mathbf{t}$ and $\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{f}$, and let Γ and Δ be as in fn 18 – so (B) holds here too. By $(\mathcal{V}\forall R)$, $\exists \mathcal{M}'' (\mathcal{V}_{\mathcal{M}''}(\Gamma) = \mathbf{t}$ and $\mathcal{V}_{\mathcal{M}''}(\phi[c/v]) = \mathbf{f}$ and $\mathcal{V}_{\mathcal{M}''}(\Delta) = \mathbf{f})$. By (E) , $\mathcal{V}_{\mathcal{M}''}(\forall v \phi) = \mathbf{f}$. But, at the same time, by (B) , $\mathcal{M}'' \equiv_c \mathcal{M}$, so $\mathcal{V}_{\mathcal{M}''}(\forall v \phi) = \mathbf{t}$. This completes the proof of $(S\forall)$.

Now, to derive $(\mathcal{V}\forall L)$ from $(S\forall)$, assume the antecedent and that $\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t}$ and $\mathcal{V}_{\mathcal{M}}(\Delta) = \mathbf{f}$. It follows that $\mathcal{V}_{\mathcal{M}}(\phi[t/v]) = \mathbf{f}$. By $(S\forall)$, $\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{f}$ too. The proof of $(\mathcal{V}\forall R)$ is slightly more convoluted. First we need to show that:

$$(F) \quad \forall \mathcal{M} (\forall \mathcal{M}' \equiv_c \mathcal{M} \mathcal{V}_{\mathcal{M}'}(\phi[c/v]) = \mathbf{t} \Rightarrow \forall \mathcal{M}' \equiv_c \mathcal{M} \forall t \mathcal{V}_{\mathcal{M}'}(\phi[t/v]) = \mathbf{t})$$

(F) is a structural principle stating that, if $\phi[c/v]$ is true in every model that assigns the same truth-values as a given model to all c -free sentences, then so is $\phi[t/v]$, for any closed term t . To see why it is true, assume that $\forall \mathcal{M}' \equiv_c \mathcal{M} \mathcal{V}_{\mathcal{M}'}(\phi[c/v]) = \mathbf{t}$ but $\mathcal{M}' \equiv_c \mathcal{M}$ and $\mathcal{V}_{\mathcal{M}'}(\phi[t/v]) = \mathbf{f}$. Let $\mathcal{M}'' \equiv_c \mathcal{M}'$ be such that, for every c -free sentence ψ , $\mathcal{V}_{\mathcal{M}''}(\psi[c/t]) = \mathbf{t}$ just in case $\mathcal{V}_{\mathcal{M}''}(\psi) = \mathbf{t}$. (We know that $(S\forall)$ does not rule out $\mathcal{M}''$ because, in that case, $\mathcal{M}'$ – and $\mathcal{M}$ – would be automatically ruled out too.) Since c does not occur in ϕ , we must have that $\mathcal{V}_{\mathcal{M}''}(\phi[t/v]) = \mathbf{f}$ and, thus, that $\mathcal{V}_{\mathcal{M}''}(\phi[c/v]) = \mathbf{f}$. But this contradicts our initial assumption, as $\mathcal{M}'' \equiv_c \mathcal{M}'$ and, so, $\mathcal{M}'' \equiv_c \mathcal{M}$. Back to the proof of $(\mathcal{V}\forall R)$, assume its antecedent and that $\mathcal{V}_{\mathcal{M}}(\Gamma) = \mathbf{t}$ and $\mathcal{V}_{\mathcal{M}}(\Delta) = \mathbf{f}$. Since c does not occur in Γ or Δ , if $\mathcal{M}' \equiv_c \mathcal{M}$, then $\mathcal{V}_{\mathcal{M}'}(\Gamma) = \mathbf{t}$ and $\mathcal{V}_{\mathcal{M}'}(\Delta) = \mathbf{f}$ too, so $\mathcal{V}_{\mathcal{M}'}(\phi[c/v]) = \mathbf{t}$. By (F) , $\forall \mathcal{M}' \equiv_c \mathcal{M} \forall t \mathcal{V}_{\mathcal{M}'}(\phi[t/v]) = \mathbf{t}$. Finally, by $(S\forall)$, it follows that $\mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t}$, completing the proof.

³⁶In this sense, substitutional quantification may also be called “sentential”. However, since it already enjoys a well-established name, I think it is safe to continue to use terminology the way I have been doing so.

entails the left-to-right direction of both the substitutional and the objectual truth-conditions for the universal quantifier.³⁷

By contrast, the converse ‘substitutional’ step from the truth of the instances $\phi[t/v]$ to the truth of $\forall v\phi$ in $\mathcal{M}$ is not clearly supported by (S $\forall$), as there could in principle be an $\mathcal{M}' \equiv_c \mathcal{M}$ in which some of those instances turn out to be false. (Later we will provide countermodels to confirm this suspicion). Similarly, it is unclear whether the ‘objectual’ step from ϕ being true of every object in $\mathcal{D}_{\mathcal{M}}$ to the truth of $\forall v\phi$ is supported. At first it might seem impossible for $\forall v\phi$ to ever be true (unless it is a logical truth), since other models assigning the same truth-values to c -free sentences might have extra elements in their domain that do not satisfy ϕ . But this is too quick: according to (S $\forall$), if $\forall v\phi$ is true then there are no such models, as $\forall v\phi$, being c -free, must be true in those models too. What we would like (but do not obviously have) is some independent grip on whether such models exist.

This leads to one of the main sources of conceptual difficulty with (S $\forall$): it is not a recursive clause like, for instance, $(\neg)$ or $(\wedge)$ are. According to (S $\forall$), the truth-value of a universal claim in a given interpretation $\mathcal{M}$ does not exclusively depend on the truth-values of simpler sentences in $\mathcal{M}$. Instead, to decide whether $\forall v\phi$ is true in $\mathcal{M}$ one needs to decide whether its instances are true, *not only in* $\mathcal{M}$, but also in all other interpretations that assign the same truth-values to c -free sentences as $\mathcal{M}$. The problem is that $\forall v\phi$ itself is precisely such a sentence. But its truth-value is just what we set out to determine in the first place! As an algorithm to decide whether a universal statement is true in an interpretation, (S $\forall$) fails miserably. But that is not its job.

One might attempt to parse (S $\forall$) by focusing on the conditions it imposes on the admissible interpretations of the language. We may ask, for example: do models in which $\forall v\phi$ is false even though ϕ holds of everything in the domain violate (S $\forall$)? But this is also hard to envision. The main reason is that the clause is, in a sense, *impredicative*, as its right-hand side contains a quantifier ranging over the class of admissible interpretations. According to (S $\forall$), an admissible interpretation $\mathcal{M}$ is one in which a universal claim is true just in case its instances are true in *every admissible interpretation* that assigns the same truth-values to c -free sentences as $\mathcal{M}$. So, to decide whether a given interpretation is admissible, one must already know the extension of the class of admissible interpretations.

If the goal is to better understand what the class of admissible models looks like according to (S $\forall$), checking whether models satisfy (S $\forall$) *individually* is, therefore, not a good strategy. A better strategy might be to focus on whole classes of models and check whether they *fit* (S $\forall$), that is, if whenever $\forall \mathcal{M}'$ in (S $\forall$) is allowed to range over that class of models, (S $\forall$) holds of every member of the class.

Of course, there are too many classes of models to try them one by one, even if we rule out those in which the propositional connectives have unintended

³⁷We can be sure of this as (V $\forall$ L), which follows from (S $\forall$) (cf. fn 35), entails the left-to-right directions of both (O $\forall$) and (Sub $\forall$) by itself (cf. fn 17).

meanings. What is more, there could be not just one but several classes of models that fit $(S\forall)$. And although we are only interested in the largest, as nothing in our inferential practice would justify focusing on a narrower class, we do not yet have a guarantee that any exists. If it did not, the inferentialist would be in a particularly bad position, as there would be no unique semantics our inference rules pin down. This would entail, for instance, that there is no unique class of models in which truth should be preserved for an argument to be valid in (classical) first-order logic, leading perhaps to different (although extensionally equivalent) notions of logical entailment.

3.2 Sentential semantics

A fixed-point construction will be of some help here. It will not only guide us towards the classes of models that fit $(S\forall)$, but also allow us to establish the existence of a largest class.³⁸

We start by hypothesizing a class $\mathfrak{M}_0$ of admissible interpretations, that is, a range for $\forall\mathcal{M}'$ in $(S\forall)$. We then revise our initial hypothesis according to $(S\forall)$, by taking the class $\mathfrak{M}_1$ of all members of $\mathfrak{M}_0$ that satisfy $(S\forall)$ whenever $\forall\mathcal{M}'$ is allowed to range over $\mathfrak{M}_0$. We then proceed to revise $\mathfrak{M}_1$, by taking the class $\mathfrak{M}_2$ of all members of $\mathfrak{M}_1$ that satisfy $(S\forall)$ whenever $\forall\mathcal{M}'$ is allowed to range over $\mathfrak{M}_1$, and so on, until no further improvements are possible, that is, until we reach a fixed point – i.e. a class of interpretations that fits $(S\forall)$.

More formally, for any ordinal α , let $\mathfrak{M}_{\alpha+1}$ be the class of models $\mathcal{M}$ in $\mathfrak{M}_\alpha$ such that

$$(J) \quad \mathcal{V}_{\mathcal{M}}(\forall v\phi) = \mathbf{t} \Leftrightarrow \forall\mathcal{M}' \in \mathfrak{M}_\alpha (\mathcal{M}' \equiv_c \mathcal{M} \Rightarrow \forall t \mathcal{V}_{\mathcal{M}'}(\phi[t/v]) = \mathbf{t})$$

$\mathfrak{M}_{\alpha+1}$ is the class of models taken from $\mathfrak{M}_\alpha$ that satisfy $(S\forall)$ when $\forall\mathcal{M}'$ ranges over all members of $\mathfrak{M}_\alpha$ itself. At limits, we take the intersection of all previous stages, that is, $\mathfrak{M}_\lambda = \bigcap_{\alpha < \lambda} \mathfrak{M}_\alpha$, for each limit ordinal λ .

The revision process is obviously monotonic: models may be discarded at each step, but are never added. This property will help us establish the existence of a fixed point to which the revision process converges.³⁹ But note that there might be more than just one fixed point – different starting hypotheses could lead us to different fixed-point classes. However, we are not interested in any such fixed points but only in the largest one, if there is one. Fortunately, that class exists; it is the fixed point that results from starting the revision process by hypothesizing the widest possible class of interpretations (as far as the universal quantifier is concerned), as we will see next.

³⁸Note that the sole purpose of this construction here is for us to better understand what the class of models admitted by the quantifier rules looks like. I am *not* claiming that the steps of the construction reflect any aspect of our understanding of the quantifiers or of our inferential practice itself – as, at least on some interpretations, the steps of Kripke's famous fixed-point construction are supposed to correspond to steps in our understanding of the word “true” (cf. Kripke [33, p. 701]).

³⁹In fact, since $(\forall L)$ and $(\forall R)$ are sound with respect to standard model-theoretic semantics, the class of standard model-theoretic interpretations is such a fixed point – it fits $(S\forall)$ (cf. fn 16).

So let $\mathfrak{M}_0$ be that class, namely, the class of all interpretations in which the propositional connectives have their standard meanings – $(\neg)$ and $(\wedge)$ hold of every member of the class. In $\mathfrak{M}_0$, quantifier-free sentences are evaluated as usual but generalizations are assigned truth-values *arbitrarily*. So, for instance, in some models in $\mathfrak{M}_0$ there will be true universal claims with false instances. And, while every sentence of the form $\neg(\phi(c) \wedge \neg\phi(c))$ will be true in every model in $\mathfrak{M}_0$, some will also make $\forall v\neg(\phi(v) \wedge \neg\phi(v))$ false. What is more, $\forall v\phi$ and $\forall v\neg\phi$ will both be true in some members of $\mathfrak{M}_0$ (but not $\forall v\phi$ and $\neg\forall v\phi$).

$\mathfrak{M}_1$ improves on $\mathfrak{M}_0$ in two ways. First, by the left-to-right direction of (J), $\mathfrak{M}_1$ leaves behind all models in which a universal claim is true but some instances are false. As a consequence, for every formula, ϕ , either $\forall v\phi$ or $\forall v\neg\phi$ must be false in every model in $\mathfrak{M}_1$, as all models satisfy $(\neg)$.

Second, since all quantifier-free logical truths hold in every model in $\mathfrak{M}_0$, by (J)'s right-to-left direction, all logical truths with exactly *one* quantifier will be true in every model in $\mathfrak{M}_1$ (and all logical falsities with exactly one quantifier will be false). On the other hand, logical truths with two nested quantifiers – e.g. of the form $\forall u\forall v\phi$ – will still be false in some models in $\mathfrak{M}_1$, as their instances – $\forall v\phi[t/u]$ – will be false in the relevant models in $\mathfrak{M}_0$. Logical truths with two quantifiers are guaranteed to be true in every model in $\mathfrak{M}_2$; in general, logical truths with n quantifiers will be true in every member of $\mathfrak{M}_n$ (and logical falsities with n quantifiers will be false in every member of $\mathfrak{M}_n$).

Importantly, there are *no further ways* in which our revision rule may improve a given hypothesis. The only models that are ruled out by the process are those that either verify a universal statement but not its instances – which happens already at the first revision step – or falsify a logical truth. Any other model $\mathcal{M}$ that could in principle be ruled out at a revision step, $\alpha + 1$, would have to falsify $\forall v\phi$ while verifying every instance $\phi[t/v]$, none of which is a logical truth. In addition, each of these instances, despite not being logically true, would have to be true as well in *every* model in $\mathfrak{M}_\alpha$ that assigns the same truth values as $\mathcal{M}$ to all *c-free* sentences, where c does not occur in ϕ . This means that, in $\mathfrak{M}_\alpha$, no model would verify each *c-free* truth of $\mathcal{M}$ and, at the same time, falsify $\phi[c/v]$. But this is not possible, as such a model is guaranteed to exist in $\mathfrak{M}_\alpha$ given our starting point.⁴⁰ To see why, allow me to focus on the

⁴⁰A formal proof of this result would be unnecessarily tedious; I just give the proof-strategy here. Note, first, that, at each stage α of the process, models in $\mathfrak{M}_\alpha$ may assign truth-values to complex sentences freely, provided that they do not violate the restrictions imposed by $(\neg)$, $(\wedge)$, and α -many iterations of (J). Since these truth-conditions can be read off corresponding inference rules of the sequent calculus (as shown in Section 1.1 and Section 3.1), it follows that, if there is no model of a set Γ in $\mathfrak{M}_\alpha$, then Γ must be provably inconsistent in the calculus – i.e. $\Gamma \vdash \emptyset$ must be derivable – appealing *exclusively* to the (operational) sequent rules for negation and conjunction, $(\forall L)$ (unless $\alpha = 0$), and at most α -many applications of $(\forall R)$. Moreover, exploiting the nice meta-theoretic features of the sequent calculus, it can be shown, by transfinite induction on α , that if c does not occur in Γ and $\Gamma \vdash \phi[c/v]$ is derivable by at most α -many applications of $(\forall R)$ but $\Gamma \vdash \forall v\phi$ is not, then $\vdash \phi[c/v]$ is provable in the calculus.

Putting these results together, let Γ be the set of *c-free* truths of $\mathcal{M}$, and assume there is no model of $\Gamma, \neg\phi[c/v]$ in $\mathfrak{M}_\alpha$. It follows that $\Gamma, \neg\phi[c/v] \vdash \emptyset$ is derivable in the calculus by at most α -many applications of $(\forall R)$, and so is $\Gamma \vdash \phi[c/v]$. Trivially, $\Gamma \vdash \forall v\phi$ is not derivable by

first revision step as an example. Further steps are slightly more convoluted but can be dealt with essentially in the same way. Recall that models in $\mathfrak{M}_0$ must abide by $(\neg)$ and $(\wedge)$ but are otherwise at absolute liberty to assign truth-values to complex sentences. So, if there is no model in which every c -free truth of $\mathcal{M}$ and $\neg\phi[c/v]$ are simultaneously true, it must be because the set of c -free truths of $\mathcal{M}$ entails $\phi[c/v]$ in classical propositional logic. But, since c does not occur in this set, this is only possible if $\phi[c/v]$ is a logical truth.

It follows from the previous considerations that, once we are through the finite stages of the revision process, no more models will be ruled out. For the appropriate truth-values of each logical truth and each logical falsity will be settled after *finitely* many steps, as each of them contains only a finite number of quantifiers. The process then reaches a fixed point at ω , $\mathfrak{M}_\omega$, which is the intersection of all finite stages.

Every model in $\mathfrak{M}_\omega$ can be easily seen to satisfy $(S\forall)$ if $\forall\mathcal{M}'$ is taken to range over $\mathfrak{M}_\omega$ itself – $\mathfrak{M}_\omega$ fits $(S\forall)$. Of course, $(\neg)$ and $(\wedge)$ hold too. What is more, no model $\mathcal{M}$ left out by $\mathfrak{M}_\omega$ could satisfy $(S\forall)$ for *any* possible range $\mathfrak{M}$ of $\forall\mathcal{M}'$ (to which $\mathcal{M}$ belongs), as it would either violate the left-to-right direction of $(S\forall)$, or make some logical truth $\forall v\phi$ with n nested quantifiers false. But in that case, there must also be a model in $\mathfrak{M}$ in which a logical truth of the form $\phi[t/v]$, with $n-1$ nested quantifiers, is false, which in turn requires the existence of a further model in $\mathfrak{M}$ in which a logical truth with $n-2$ nested quantifiers is false, and so on. Ultimately, this entails that there must be a model in $\mathfrak{M}$ in which a logical truth with no quantifiers – i.e. a tautology – is false. But those models have been excluded already by the rules for negation and conjunction, as they must violate either $(\neg)$ or $(\wedge)$. So $\mathfrak{M}_\omega$ is the largest class of models with respect to which the inference rules for the universal quantifier, as well as the other logical terms, are valid – it is the class of interpretations in which the logical constants, including the quantifiers, have exactly the meanings that the rules fix for them. From an inferentialist point of view, $\mathfrak{M}_\omega$, the semantics determined by the calculus, is the *true* semantics of our logical vocabulary – to which I refer as “sentential semantics”.

Recall that, unless $\forall v\phi$ is a logical truth, models in which this sentence is false yet every one of its instances is true will not be discarded at *any* stage of the revision process. In other words, substitutional truth-conditions for the quantifiers fail to hold in sentential semantics, unsurprisingly. But note as well that in some of those models that are immune to the revision process ϕ will be true of every member of the domain regardless of whether it has a name in the model, yet $\forall v\phi$ will be false. So objectual truth-conditions also fail to hold in this semantics. In other words, $(O\forall)$ does not follow from $(S\forall)$ or, equivalently, from the conjunction of $(\forall\forall L)$ and $(\forall\forall R)$, as anticipated in Section 1.2, and neither does $(Sub\forall)$. This is, precisely, Carnap’s categoricity problem for the quantifiers: the sequent rules fail to pin down either the objectual or the substitutional interpretation of the quantifiers.

at most α -many applications of this rule, or we would have that $\forall v\phi$ is true in $\mathcal{M}$. Therefore, $\vdash \phi[c/v]$ must be derivable in the calculus, so $\phi[c/v]$ is a logical truth.

It follows that, according to sentential semantics, the truth of a generalization in a model does not supervene on the truth of its instances in that model, whether expressible in the language or not. So, unlike substitutional or objectual quantification, sentential quantification is not *truth-functional*. The meaning of the sentential quantifiers is not fully compositional.

Trivially, the sequent calculus is sound and complete with respect to sentential semantics.⁴¹ But, as is well known, it is also complete with respect to standard model-theoretic semantics. So, although standard semantics is, strictly speaking, incorrect from an inferentialist point of view, for the purpose of defining logical consequence and other logical notions, it will not make any difference which semantics we choose – standard or sentential. There are other contexts, however, in which the choice has a substantive impact, as will be seen in the next section.

4 Implications, Objections, and Replies

I have argued that our quantifiers are not objectual but sentential. This is a radical departure from orthodoxy, which has the potential to illuminate existing problems, but which will also prompt natural concerns. This section applies the account to an old issue in the metaphysics of logic but, before, it considers a major source of discomfort, stemming from the existence of legitimate interpretations of the language in which a universal generalization $\forall v\phi$ is false despite ϕ being true of each object in the domain. (This discomfort may likely intensify upon the realization that, given the duality of the universal and the existential quantifier, in such interpretations there are formulae ψ – e.g. $\neg\phi$ – satisfied by no object in the domain, yet $\exists v\psi$ turns out to be true.)

4.1 A semantic concern

An implication of the sentential truth clause is that there exist models that falsify a universal claim $\forall v\phi$ even though ϕ is true of each object in the domain. This might be used as the basis for a semantic objection. A natural way to think of a model is as a representation of a scenario. In particular, $\mathcal{M}$ represents a scenario in which what *exists* is exactly what is in $\mathcal{D}_{\mathcal{M}}$. But, then, the worry goes, “ $\forall$ ” cannot really mean *everything*, for in such cases *everything that exists in the scenario* satisfies the condition ϕ and, yet, $\forall v\phi$ comes out false. If sentential quantification is all we have, one might be forgiven for thinking that we do not have quantification at all.

This semantic complaint is ungrounded, however, because it relies on a premise – that models represent scenarios in which the objects that exist lie

⁴¹Letting $\mathcal{M}$ range over $\mathfrak{M}_{\omega}$, $(\neg)$, $(\wedge)$, and $(S\forall)$ are equivalent to the validity in $\mathfrak{M}_{\omega}$ of the sequent rules for negation, conjunction, and the universal quantifier, respectively. Thus, since $(\neg)$, $(\wedge)$, and $(S\forall)$ all hold in $\mathfrak{M}_{\omega}$, the sequent rules are sound with respect to this class. Conversely, since $\mathfrak{M}_{\omega}$ is the *largest* class of models satisfying $(\neg)$, $(\wedge)$, and $(S\forall)$, the calculus must also be complete – if $\Gamma \vdash \phi$ is not derivable, there must be some models in which Γ is true and ϕ false that do not violate $(\neg)$, $(\wedge)$, and $(S\forall)$ and, thus, belong to $\mathfrak{M}_{\omega}$.

within the domain – which has little or no support. Note that sentential truth-conditions make no reference at all to a domain – the truth-values of quantified claims are determined exclusively in terms of the truth-values of other expressions. The only function served by the domain of a model is to provide a range of semantic values for the non-logical terms.⁴² So, while models are equipped with a distinguished set which we might call, in fidelity to tradition, a “domain”, the name is misleading, because the connection between this set and the quantifiers has been severed. There is therefore no reason why we should expect a model to represent a scenario in which everything that exists belongs to its domain. Consequently we should not expect the fact that ϕ is true of every object in $\mathcal{D}_{\mathcal{M}}$ to entail that $\forall v\phi$ is true in $\mathcal{M}$; we have not been told that the members of $\mathcal{D}_{\mathcal{M}}$ are *all* there is in this context. To put it differently, the ‘instances’ themselves do not suffice for a universal claim; a premise to the effect that those instances cover *all* possible cases is needed. Just to be clear, I’m not denying that, e.g. “everything is material” is true whenever the predicate “is material” holds of everything. In general, sentential semantics is perfectly compatible with a universal claim $\forall v\phi$ begin true just in case *everything* satisfies ϕ – a straightforward consequence of Tarski’s well-known satisfaction schema. What is being denied is that models successfully represent scenarios in which everything lies within the domain.

If the domain of a model does not convey what exists in the scenario it purports to represent, what can be added to the description of a model in order to achieve this effect? Purely metalinguistic statements such as “everything belongs to $\mathcal{D}_{\mathcal{M}}$ ” or “everything in $\mathcal{M}$ belongs to $\mathcal{D}_{\mathcal{M}}$ ” are hopeless, as there will be either plainly false – $\mathcal{D}_{\mathcal{M}}$, for instance, does not belong to $\mathcal{D}_{\mathcal{M}}$ – or simply tautological – it is not very informative to assert that everything in $\mathcal{D}_{\mathcal{M}}$ is a member of $\mathcal{D}_{\mathcal{M}}$. The only way for a model to explicitly represent a scenario with a certain ontology is to ensure that certain object-language claims expressing (in)existence facts are true in the model. To go back to the preliminary diagnosis offered at the very end of Section 2.3, in a semantics determined by inferences between statements, extralinguistic facts about a model will have hardly any influence on the truth-values of sentences in that model.

For instance, a model can represent that *only* the referents of c_1, c_2 , and c_3 exist if it makes the following sentence true:

$$(1) \quad \forall x (x = c_1 \vee x = c_2 \vee x = c_3)$$

Naturally, this strategy works only if we have finitely many objects in the domain, all of which have names in the model. What about infinite cases? Con-

⁴²This function is completely inessential. In fact, as far as the meanings of logical terms are concerned, domains could drop out of the picture entirely. We could even let go of interpretation functions – their role of fixing the truth-values of atomic statements could be taken up directly by the valuation function, which would remain the single component of a model. Zalabardo [34] and Ben-Yami [35] make similar points regarding (their respective versions of) substitutional semantics. Note also that, by disposing of domains and interpretation functions, we get rid of all ‘representational’ components of Tarski’s model-theoretic semantics, to use Etchemendy’s [36] terminology, thereby obtaining what appears to be a purely ‘interpretational’ semantics for first-order logic.

sider the following schema:

$$(2) \quad \forall \mathcal{M}' \sim_c \mathcal{M} \mathcal{V}_{\mathcal{M}'}(\phi[c/v]) = \mathbf{t} \Rightarrow \mathcal{V}_{\mathcal{M}}(\forall v \phi) = \mathbf{t}$$

A model satisfies this schema if, for any formula ϕ , whenever ϕ is true of every object in $\mathcal{D}_{\mathcal{M}}$, $\forall v \phi$ is true. This schema approximately expresses that what exists is contained within $\mathcal{D}_{\mathcal{M}}$. (But only approximately: for all it says, there could be objects in the scenario that do not belong to $\mathcal{D}_{\mathcal{M}}$ but simply happen to satisfy ϕ whenever every member of $\mathcal{D}_{\mathcal{M}}$ does. To adapt a quote from Quine [37, p. 210], it could just happen that the objects that do not belong to $\mathcal{D}_{\mathcal{M}}$ are *inseparable* from the ones that do. Against the backdrop of scientific naturalism, if determining the truth-values of all sentences does not suffice to fix the range of the quantifiers, then nothing else can.)

The reader might be puzzled at this point. Is not (2) the direction of the objectual clause for the universal quantifier that sentential semantics invalidates? It is. Could we then impose it on all models and thereby enforce objectual truth-conditions for our quantifiers? Not really. (2) is simply a condition that some but not all models happen to satisfy. We could only ‘impose it’ on all interpretations if we could argue that the quantifier rules rule out all models in which it does not hold. But if I have argued correctly so far, this cannot be done. (By analogy: we can write down a condition that is satisfied by some model of a given first-order language if every object in its domain has a name in the model. Some models satisfy this condition; some do not. But it does not follow that the condition can in some way stem from our inferential practice.)

Notice also that the use of (2) to demarcate the ontology of $\mathcal{M}$ *presupposes* the fact that “ $\forall$ ” is a universal quantifier; it does not purport to explain this fact. By contrast, in the objectual semantics, (2) essentially characterizes the meaning of “ $\forall$ ”. It cannot, on pain of explanatory circularity, also demarcate the ontology of the relevant scenario. *Nothing* in an objectual model indicates that the domain exhausts the ontology of the model. On the objectual semantics, the truth of some ‘instances’ suffices for the truth of the universal claim without any additional information that those instances exhaust all the possibilities in this context. As I argued before, objectual quantifiers are closer in meaning to an infinite conjunction than to the quantifiers we use in our language.

To summarize, a model in which $\forall v \phi$ is false represents a scenario in which ϕ is *not* true of every object, regardless of whether it holds of everything in the model’s domain. The mistake is to suppose a tighter link than holds, in general, between the domain of a model and the objects which the associated scenario represents as existing.

4.2 A unified picture of quantification

Let me close this section on a positive note. In [37] Quine launched an objection against substitutional quantification. His main charge was that, unlike objectual quantification, “the key idiom of ontology”, its “substitutional counterfeit [...] turns its back on reference”. What is more,

Its nonreferential orientation is seen in the fact that it makes no essential use of namehood. That is, additional quantifications could be explained whose variables are place-holders for words of any syntactical category. (Quine [37, p. 209])

Quine is likely to be correct that substitutional quantification is counterfeit, but for the wrong reasons. For note that, although sentential quantification is not oblivious to ontology, as I argued in this subsection, like its substitutional cousin it also makes no essential use of namehood. We could replace c and t in $(S\forall)$ with, for example, appropriate meta-variables ranging over monadic predicates of the language, thereby obtaining adequate truth-conditions for universal quantification into monadic predicate position. In general, we could replace c and t in $(S\forall)$ with appropriate meta-variables ranging over expressions of *any* type, and each of the instances of the resulting principle would automatically serve as adequate truth-conditions for quantification of the corresponding order.

But where Quine sees a flaw, I see a virtue. By putting forward isomorphic truth-conditions for quantification into different types of expression, sentential semantics provides a unified picture of quantification without incurring any controversial claims about purported domains of objects quantifiers other than first-order supposedly range over. From an inferentialist perspective, a unified picture is a desirable outcome, as the rules of inference that typically govern first- and higher-order quantifiers share the same structure – first-order quantification is not too different from higher-order quantification on this view.

5 Concluding Remarks

I have argued that our quantifiers do not mean what we typically think they mean. They do not satisfy substitutional or even objectual truth-conditions; only sentential truth-conditions. This is a bold claim, but I have shown that it follows from two plausible assumptions: naturalism – which does not seem to leave us with much besides ordinary rules of use to explain how the meaning of logical terms is determined – and inferentialism – the view that these meaning-determining rules take the form of inferences between statements. Thus, to resist the conclusion of this paper, one would have to reject at least one of these views.

One route is to give up naturalism and explore some supernatural mechanism by which our quantifiers would acquire an objectual interpretation. But if one prefers to explain meaning in a scientifically respectable way, the only way to resist the main conclusion of this paper is to give up inferentialism. There is one initially appealing ‘loophole’: as I have understood inferentialism, it involves inferences between linguistic expressions. But those working in the inferentialist tradition have also explored so-called language-entry and language-exit rules, which generalize the notion of inference to include transitions between sentences and non-doxastic mental states – e.g. perceptual or observational states on the

entry side, states involved in intentional action on the exit side.⁴³ It is reasonable to ask whether such a generalized inferentialism might fare better at fixing the objectual meanings of the quantifiers.

While the issue deserves further exploration, let me briefly record a pessimistic verdict here. Any such language-entry rules for the universal quantifier would have to guarantee what the ordinary inference rules cannot, i.e. that we are able to conclude a universal claim whenever we have its ‘instances’ (regardless of whether they have a linguistic correlate). Perhaps entry rules help in *some* cases, perhaps where the generalization is concerned with material objects available to perception. For instance, we could simply observe that the relevant objects in a certain context are so and so, and nothing else. But it is hard to see how entry rules help more generally, for example, where the relevant objects are abstract or infinitely many. For now, then, we should remain sceptical that we have objectual quantification in our language.

References

- [1] W. Sellars, Inference and Meaning, *Journal of Symbolic Logic* 21 (2) (1956) 203–204. doi:10.2307/2268776.
- [2] R. Brandom, *Making it Explicit: Reasoning, Representing, and Discursive Commitment*, Harvard University Press, Cambridge, MA, 1994. doi:10.2307/2956391.
- [3] H. Putnam, Models and reality, *Journal of Symbolic Logic* 45 (3) (1980) 464–482. doi:10.2307/2273415.
- [4] N. Tennant, *The Taming of the True*, Oxford University Press, Oxford, 1997. doi:10.1093/acprof:oso/9780199251605.001.0001.
- [5] T. Button, S. Walsh, *Philosophy and Model Theory*, Oxford University Press, Oxford, 2018. doi:10.1093/oso/9780198790396.001.0001.
- [6] T. Button, Mathematical Internal Realism, in: J. Contant, S. Chakraborty (Eds.), *Engaging Putnam*, De Gruyter, Berlin, 2022, pp. 157–182. doi:10.1515/9783110769210-007.
- [7] J. Warren, Infinite reasoning, *Philosophy and Phenomenological Research* 103 (2021) 385–407. doi:10.1111/phpr.12694.
- [8] R. Carnap, *Formalization of Logic*, Harvard University Press, Cambridge, MA, 1943. doi:10.2307/2181359.
- [9] T. Smiley, Rejection, *Analysis* 56 (1) (1996) 1–9. doi:10.1111/j.0003-2638.1996.00001.x.

⁴³See, for instance, Sellars [1] and Brandom [2].

- [10] I. Rumfitt, “Yes” and “No”, *Mind* 109 (436) (2000) 781–823. doi:10.1093/mind/109.436.781.
- [11] J. Murzi, Classical harmony and separability, *Erkenntnis* 85 (2) (2020) 391–415. doi:10.1007/s10670-018-0032-6.
- [12] J. W. Garson, *What Logics Mean. From Proof Theory to Model-Theoretic Semantics*, Cambridge University Press, Cambridge, 2013. doi:10.1017/CBO9781139856461.
- [13] N. Tennant, Negation, absurdity and contrariety, in: D. M. Gabbay, H. Wansing (Eds.), *What is negation?*, Kluwer Academic Publishers, Dordrecht, 1999, p. 199–222. doi:10.1007/978-94-015-9309-0_10.
- [14] F. Steinberger, *Harmony and logical inferentialism*, Phd thesis, Cambridge University, Cambridge (2009).
- [15] G. Restall, Minimalists about truth can (and should) be epistemicists, and it helps if they are revision theorists too, in: J. C. Beall, B. Armour-Garb (Eds.), *Deflationism and Paradox*, Oxford University Press, Oxford, 2005, pp. 97–106.
- [16] I. Rumfitt, Knowledge by Deduction, *Grazer Philosophische Studien* 77 (1) (2008) 61–84. doi:10.1163/18756735-90000844.
- [17] F. Steinberger, Why Conclusions Should Remain Single, *Journal of Philosophical Logic* 40 (3) (2011) 333–355. doi:10.1007/s10992-010-9153-3.
- [18] J. Murzi, O. T. Hjortland, Inferentialism and the categoricity problem: Reply to raatikainen, *Analysis* 69 (3) (2009) 480–488. doi:10.1093/analysis/anp071.
- [19] V. McGee, ‘Everything’, in: G. Sher, R. Tieszen (Eds.), *Between Logic and Intuition. Essays in Honor of Charles Parsons*, Cambridge University Press, Cambridge, 2000, pp. 54–78. doi:10.1017/CBO9780511570681.005.
- [20] H. Field, *Truth and the Absence of Fact*, Oxford University Press, Oxford, 2001.
- [21] J. Warren, *Shadows of Syntax: Revitalizing Logical and Mathematical Conventionalism*, Oxford University Press, New York, 2020. doi:10.1093/oso/9780190086152.001.0001.
- [22] J. Murzi, B. Topey, Categoricity by Convention, *Philosophical Studies* 178 (10) (2021) 3391–3420. doi:10.1007/s11098-021-01606-3.
- [23] D. M. Armstrong, Dispositions as categorical states, in: T. Crane, D. M. Armstrong, C. B. Martin (Eds.), *Dispositions: A Debate*, Routledge, New York, 1996, pp. 15–18.
- [24] D. K. Lewis, Finkish dispositions, *Philosophical Quarterly* 47 (187) (1997) 143–158. doi:10.1111/1467-9213.00052.

- [25] A. Bird, Causation and the manifestation of powers, in: A. Marmodoro (Ed.), *The Metaphysics of Powers: Their Grounding and Their Manifestations*, Routledge, New York, 2010, pp. 160–168.
- [26] J. Warren, Killing Kripkenstein’s Monster, *Noûs* 54 (2) (2020) 257–289. doi:10.1111/nous.12242.
- [27] D. Bonnay, D. Westerståhl, Compositionality Solves Carnap’s Problem, *Erkenntnis* 81 (4) (2016) 721–739. doi:10.1007/s10670-015-9764-8.
- [28] P. Schroeder-Heister, A natural extension of natural deduction, *Journal of Symbolic Logic* 49 (4) (1984) 1284–1300. doi:10.2307/2274279.
- [29] P. Boghossian, What is inference?, *Philosophical Studies* 169 (1) (2014) 1–18. doi:10.1007/s11098-012-9903-x.
- [30] C. Wright, Comment on Paul Boghossian, “What is Inference”, *Philosophical Studies* 169 (1) (2014) 27–37. doi:10.1007/s11098-012-9892-9.
- [31] J. Warren, Functionalism about inference, *Inquiry: An Interdisciplinary Journal of Philosophy* (forthcoming). doi:10.1080/0020174x.2022.2075452.
- [32] R. Brandom, *Articulating Reasons*, Harvard University Press, Cambridge, MA, 2000. doi:10.4159/9780674028739.
- [33] S. Kripke, Outline of a theory of truth, *Journal of Philosophy* 72 (19) (1975) 690–716. doi:10.2307/2024634.
- [34] J. L. Zalabardo, Logic Without Metaphysics, *Synthese* 198 (S22) (2021) 5505–5532. doi:10.1007/s11229-019-02124-w.
- [35] H. Ben-Yami, Truth and Proof Without Models: A Development and Justification of the Truth-Valuational Approach (Manuscript, 2024).
- [36] J. Etchemendy, *The Concept of Logical Consequence*, Harvard University Press, Cambridge, MA, 1990. doi:10.2307/2220083.
- [37] W. V. O. Quine, Ontological Relativity, *Journal of Philosophy* 65 (7) (1968) 185–212. doi:10.2307/2024305.