

Semantic Perversity*

Francisco Calvo Garzón

RESUMEN

Este artículo consta de dos partes. En la primera sección exploraré una propuesta de Christopher Hookway (1988) para eludir un contraejemplo que Gareth Evans (1975) ofrece contra la tesis quineana de la inescrutabilidad de la referencia. Evans propuso una argumentación que sugiere que uno de los manuales de traducción perversos quineanos es conductualmente incorrecto. Hookway modifica dicho manual perverso para hacerlo conductualmente correcto. Tal y como veremos, la propuesta de Hookway no consigue establecer las condiciones de satisfacción adecuadas. Afortunadamente, una simple modificación permitirá a Hookway hacer frente a la objeción de Evans. En la segunda sección propondré una estrategia quineana alternativa para eludir el contraejemplo de Evans; una estrategia no sujeta a ciertas críticas que podrían hacer peligrar la propuesta de Hookway.

ABSTRACT

This paper consists of two parts. In section I, I explore Christopher Hookway's proposal (1988) to elude a counter-example which Gareth Evans [Evans (1975)] offers against Quine's Thesis of the Inscrutability of Reference. Evans produced a line of argument which suggests that one of Quine's semantically perverse translation manuals is behaviourally incorrect. Hookway modifies the perverse manual to make it behaviourally correct. As we'll see in due course, Hookway's proposal fails to deliver the right satisfaction conditions. Fortunately, just a minor modification will suffice for Hookway to meet Evans' objection. In section II, I shall offer the Quinean a different strategy to bypass Evans' counter, a strategy which is not subject to certain criticisms which may put Hookway's own proposal in jeopardy.

INTRODUCTION

In a nutshell, Quine's Thesis of the Inscrutability of Reference claims that there is no objective fact of the matter as to what the *ontological commitments* of the speakers of a language are. To become acquainted with this polemical thesis, Quine (1960) invites the reader to imagine two linguists whose task is to produce rival translation manuals to account for the expressions of an unknown language. The manuals, when finally completed, should be able to correlate each of the potentially infinite number of sentences uttered by natives with one or more sentences belonging to the linguist's home

language. The linguists are not allowed to correlate native expressions with those of the Home language on the grounds that they pin down the same *idea*. Quine's naturalism forbids them to pair words with language-independent mental acts, acknowledging as a genuine evidential basis only the stimulation of their sensory receptors. Upon this they will try theories in search of true predictions. The linguists, as we shall see shortly, can produce rival translation manuals which are mutually incompatible, and yet fit all possible evidence. The Inscrutability Thesis is the doctrine that there is no fact of the matter as to what the extensions of the terms of a language are. Claims about the ontological commitments that the speakers of a language incur are relative to which translation manual we favour.

Take, for instance, the native sentence "Gavagai." Let us suppose that it has been empirically determined that "Gavagai" relates to portions of space-time in the vicinity of the native speaker, which are *rabbit-related*. We may then translate "Gavagai" with our "Lo, a rabbit." However, it would be rash to impute our ontology to the natives. Quine maintains that the extension of the term 'gavagai' could be taken to be the set of undetached rabbit parts. His conclusion is that it is inscrutable what the expression 'gavagai' refers to: The linguist may assign to the native expression as its extension either the set of rabbits or the set of undetached rabbit parts. Our only hope, in Quine's view, of solving the indeterminacy is by looking at the interaction of such expressions with the *apparatus of individuation* (plurals, identity, etc.). Unfortunately, this hope is thwarted, Quine argues, since the apparatus of individuation is itself inscrutable too¹. Were Quine's radical thesis to earn its keep, objectivism as applied to our ordinary notion of reference, and related semantic notions — truth, meaning, etc. — would be in serious jeopardy.

I. EVANS' COUNTER-EXAMPLE AND THE 'DIVIDE-AND-RULE' STRATEGY

Evans (1975) produced a line of argument which suggests that semantically perverse translation manuals *à la* Quine are *behaviourally* incorrect. In my opinion, the interest of Evans' argument relies in the fact that, unlike some foes of Quine that insist in the need of honouring a mentalistic level of explanation [see Kirk (1986), and the references there], Evans' attack is launched from within a Quinean framework. Evans tries to show that the perverse referential schemes will not be able to cope with all the data that the standard scheme does. And Evans confines himself to a pool of data which Quine would acknowledge as genuine evidential basis: namely, native assent and dissent to the linguist's queries under concurrent observable circumstances. If Evans' attack is sound, it may prove fatal since Quine will not be able to reply by claiming that Evans' criticism relies on non-factual consid-

erations. Evans' anti-Quinean line of argument is a powerful one, and I shall spend some time in this section to review it².

Evans starts by pointing out the divergencies between the task of a translator and the task of a semanticist. The aim of the former is simply to facilitate communication between two linguistic communities. In order to do so, she must devise a manual of translation. Evans does not manifest any concern with the claim that translation suffers from indeterminacy. The reason is simply that a translator is not devoted to revealing any *semantic* truth. The translator's aim is simply to find smooth vehicles of communication, and insofar as this target is achieved, the way the translator dissects native utterances is completely irrelevant to her task. By contrast, the semanticist is involved in the project of constructing a theory of meaning. She is not concerned merely with correlating expressions of native with lumps of home language, but rather with stating what the native expressions actually *mean*³. The sentences of Native are potentially infinite in number. The semanticist, similarly to the translator, will be obliged to dissect native sentences. The target now, however, is to account for the meaning of those previously unencountered native utterances in a recursive way. But in opposition to the case of Radical Translation, Evans claims, not any given set of analytical hypotheses will do. Quine's perverse treatment of certain compound expressions, as we shall see next, is the root of Evans' distrust.

To introduce Evans' counter-example, consider the native expression 'Blanco gavagai'. Natives utter 'Blanco gavagai' *only* when a white rabbit shows up in their visual field. Take two semantic theories of native, one standard and the other perverse⁴. On the one hand the Standard Theory, **ST**, deals with 'Blanco gavagai' in the following way:

ST

Axioms:

- (a) (x)(x satisfies 'gavagai' iff x is a rabbit)
- (a1) (x)(x satisfies 'blanco' $\wedge$ f iff (x is white & x satisfies f))

Theorem:

- (a2) (x)(x satisfies 'blanco' $\wedge$ 'gavagai' iff (x is white & x is a rabbit))⁵.

On the other hand the perverse semanticist offers the following alternative:

PT1

Axioms:

- (b) (x)(x satisfies 'gavagai' iff x is an undetached rabbit part)

(b1) $(x)(x \text{ satisfies 'blanco' } \wedge f \text{ iff } (x \text{ is white } \& x \text{ satisfies } f))$

Theorem:

(b2) $(x)(x \text{ satisfies 'blanco' } \wedge \text{'gavagai' iff } (x \text{ is white } \& x \text{ is an undetached rabbit part.})$

Let us suppose that **ST** is behaviourally adequate. We can, thus, identify the sentence 'Blanco gavagai' with 'There is a white rabbit here'. However, Evans argues, if **ST** is behaviourally adequate, then **PT1** is not behaviourally adequate. There are certain circumstances in which **PT1** fails to reflect correctly the native's linguistic behaviour [Evans (1975), p. 358], — assuming **ST** does correctly reflect the native's linguistic behaviour. The sort of situation Evans is thinking of is for example when native speakers are stimulated by a brown rabbit with a white leg. In this case, **PT1** is not faithful to the evidence since, assuming **PT1**, natives should assent to 'Blanco gavagai?' when stimulated by a white-legged brown rabbit⁶. But, we have assumed that **ST** is behaviourally correct, and hence that natives would assent to the combined construction 'Blanco gavagai?' only in presence of a white rabbit.

There is a further alternative that Evans himself advances. In order to avoid the inconvenient consequences of white-legged brown rabbits, the obvious move is to link the satisfaction conditions of 'blanco' to things which are parts of white rabbits. The perverse theory would then require an axiom of the form:

(b1)* $(x)(x \text{ satisfies 'blanco' iff } x \text{ is an undetached part of a white rabbit}).$

But this move only brings further difficulties: What will the native say about white sheets of paper, snowed landscapes, and so on? It seems that we are obliged to extend the scope of (b1)* in order to talk about white things other than rabbits. Hence, the broader axiom required should run as follows:

(b1)** $(x)(x \text{ satisfies 'blanco' iff } x \text{ is an undetached part of a white thing}).$

But, as Evans notices, the Quinean still faces a similar worry to the one motivated by white-legged brown rabbits. According to (b1)**, 'Blanco gavagai?' should be assented to when a claw of a white-legged brown rabbit is present. For the claw itself is a part of a white thing: namely, a white leg. At this point, Evans doesn't pursue these matters further. It seems there is nothing the Quinean can do.

Hookway, however, proposes a rejoinder to the difficulties which Evans has raised for the Quinean thus far. He contends that the problem arising with (b1)* does not force us to go for (b1)**. If we want to refer to white sheets of paper or snowed landscapes, then the way to do so is by displaying the satis-

faction conditions of ‘blanco’ in a *context-sensitive* way. In order to do so Hookway [Hookway (1988), p. 155] offers a disjunctive axiom. Recasting Hookway’s axiom in our terminology we get:

PT2

Axioms:

- (c) (x)(x satisfies ‘gavagai’ iff x is an undetached rabbit part).
- (c1) (x)(x satisfies ‘blanco’ iff either (a) ‘blanco’ occurs together with ‘gavagai’ and x is an undetached part of a white animal, or (b) ‘blanco’ occurs in some other context and x is white).

Theorem:

- (c2) (x)(x satisfies ‘blanco’ $\wedge$ ‘gavagai’ iff (x is an undetached part of a white animal)).

Hence, if the native utters ‘Blanco gavagai’ we employ the first disjunct of (c1). Otherwise, we use the second.

According to the theorem generated by **PT2**, (c2), native speakers should assent to ‘Blanco gavagai?’ in the vicinity of a white cat or a white cow. The reason is obvious: Any undetached part of a white cat or a white cow is an undetached part of a white *animal*. I ignore what moved Hookway to formulate his proposal in terms of animals. However, it does not cause great inconvenience, for the modification required is minimal. By substituting ‘rabbit’ for ‘animal’ in the first disjunct of (c1), we shall obtain the correct satisfaction theorem. Hence the perverse semantic theory Hookway requires is:

PT3

Axioms:

- (d) (x)(x satisfies ‘gavagai’ iff x is an undetached rabbit part)
- (d1) (x)(x satisfies ‘blanco’ iff either (a) ‘blanco’ occurs together with ‘gavagai’ and x is an undetached part of a white *rabbit* or (b) ‘blanco’ occurs in some other context and x is white).

Theorem:

- (d2) (x)(x satisfies ‘blanco’ $\wedge$ ‘gavagai’ iff (x is an undetached part of a white *rabbit*))

Now, **PT3** is behaviourally correct if **ST** is, as required. natives will *only* assent to ‘Blanco gavagai?’ in presence of a white *rabbit*. When dealing with white cats or white cows the second disjunct, (b), of (d1) will come to the

rescue. Hookway's disjunctive strategy, as reformulated in **PT3**, seems to succeed in eluding Evans' counter-example.

II. SEMANTIC PERVERSITY: OVERCOMING SOME POTENTIAL THREATS

Hookway is careful not to offer his perverse manual as conclusive against Evans, but rather as 'no more than a first approximation to a satisfactory response' [Hookway (1988), p. 155]. Hookway acknowledges the possibility that 'the attempt to develop this proposal consistently would run into technical difficulties' [*ibid.*, p. 155]. There are at least two hurdles. I shall try to make manifest why these two hurdles can jeopardize Hookway's overall enterprise and, afterwards, I shall offer a different perverse route which overcomes both difficulties.

On the one hand, the semantic perversity of **PT3** is rather narrow in scope. **PT3**'s results coincide with the standard ones, as achieved via **ST**, except for rabbit expressions: The satisfaction conditions of 'blanco' are linked to *undetached parts* of white [...] *only* when 'blanco' is coupled with 'gavagai'. In all other cases, **PT3** behaves standardly, taking 'blanco'-related utterances to be associated with whole enduring white cats or white sheets of paper, for example. This hybrid character of **PT3** (i.e., standard-cum-perverse) seems to be alien to Quine's original pursuit. Quine's aim was to produce a *fully perverse* alternative to **ST** in the sense that for *every* standard referent that **ST** picks out, a perverse counterpart is offered.

Now, it seems that when we try to broaden the scope of Hookway's perverse route we are in trouble. If **PT3** is to account for Evans' counter while being fully-perverse, (c1) will have indefinitely many disjuncts. We will require an indefinite number of disjuncts in order to link the satisfaction conditions of 'blanco' to the appropriate wholes of undetached parts of rabbits, cats, cows, paper, etc., etc. And the same will happen with respect to all those axioms required for dealing with any other native colour-word, and indeed, with any other native expression for which a version of Evans' counter can be put forward. Therefore, it *may* be the case that the perverse semanticist will not be able to state a fully-perverse *disjunctive* semantic theory.

However, in fairness to Hookway, we ought to notice that this difficulty is rooted on rather speculative grounds. First, it is unclear why the Quinean should not favour an array of merely hybrid semantic theories, rather than a single fully-perverse one. And second, even if the Quinean wishes to be fully-perverse, it is not obvious that the aforementioned difficulty could not be overcome by some baroque plot which the Quinean has up her sleeve. Nevertheless, I shall not expand on these considerations, for there is still a second hurdle which seems to me crucial against Hookway. Even though we fa-

voured **PT3** as it stands, there is a further worry, raised once more by Evans, which, in my view, Hookway underestimates.

Confronted with **ST** and **PT3**, Evans would argue that there is actually a body of evidence favouring **ST**. Evans' idea [Evans (1981), pp. 124-7] relies on a particular approach to the Theory of Meaning. A Theory of Meaning, by contrast to a Theory of Translation, aims to provide a *psychological explanation* of the speakers' verbal behaviour by singling out certain behavioural dispositions. Given a putative semantic theory, speakers are ascribed a set of dispositions: one corresponding to each axiom of the semantic theory, or to each disjunct if the axiom is itself disjunctive. These linguistic dispositions, Evans argues, constitute tacit knowledge of *one* specific theory of meaning. We can decide which semantic theory is the correct one by looking at the dispositions of the speakers. Some theories provide better explanations of the native's linguistic behaviour than others⁷. In this way, we may find that native speakers do follow, though tacitly, **ST**, and not **PT3**, by observing for instance that mastering the term 'blanco' in contexts which do *not* include 'gavagai', permits natives to *understand* such expression in *all* contexts. If this were to be the case, then this behavioural evidence would favour **ST** over **PT3**, since native speakers would have just *one single disposition* for judging sentences containing 'blanco' as having such-and-such truth-conditions (as opposed to having two different dispositions, as occurs under **PT3**: one to account for (**d1**), the other for (**d2**)).

Hookway [Hookway (1988), pp. 155-62] considers Evans' view that one semantic theory may give a better psychological explanation of a speaker's verbal behaviour than another, but he believes that it poses no serious threat to **PT3**. His reason is that the Quinean would simply reject as *non-factual* any psychologically-based criterion which goes beyond the description of the observable behaviour of speakers. Hookway's Quinean notes that 'unless psychological explanations simply allude to physical mechanisms, they do not enhance our knowledge of (physical) reality' [Hookway (1988) p. 159].

However, Hookway is overlooking a crucial point: Namely, that Evans' argument can be transposed into a form which a physicalist will have to admit as legitimate. The key point is that there must be some relation between speaker's linguistic manifestations and the information content of *physical states* in their brains, such that the canonical route in a theory of meaning leading from its axioms to the theorems produced reflects a neurophysiological causal structure found underlying the competencies of the speakers⁸. This means that there must be a neurophysiological explanation of the way competent speakers understand their language. And this causal explanation will provide us with a picture of the actual route leading from the speaker's dispositions associated with the atomic elements of native to the overall physical states associated with the whole sentences they produce. A semantic theory will thus be *empirically grounded* since the tacit knowledge of the semantic

theory ascribed to a certain speaker of native comes in terms of the causal explanatory states attributed to her (or better said, to her internal information-processing system). Once we know how this internal system operates, it is theoretically plausible that we can determine whether a speaker tacitly follows **ST** or **PT3**. If future neuroscience reveals that there is one single neurophysiological state causally activated when a native utters 'blanco' in all different contexts, then that would count as evidence against **PT3**, since assuming **PT3** we would require two different neurophysiological states: Namely, one state exclusively responsible for 'blanco' when coupled with 'gavagai' and a different one causally responsible for all other 'blanco'-related utterances⁹.

Nevertheless, even if we granted without further ado these hypothetical neurophysiological data against **PT3**, I contend that it does not follow that **ST** is the *only* correct theory. Evans' (1975) attack on Quine's Inscrutability Thesis has been so widely well received by the philosophical community because of an implicit, though misleading, assumption made by foes and sympathizers of Quine alike. Namely, that reference is to be divided over objects in a *monolithic* fashion. Evans [Evans (1975), p. 362] talks in terms of semantic theories that cut the reference of 'gavagai' finer than the standard theory does — e.g., over undetached rabbit parts¹⁰. It is however tacitly assumed that finer cuts, such as the division of the reference of 'gavagai' over undetached rabbit parts, constitute a monolithic block. That is, the axioms that deal with the satisfaction conditions of 'gavagai' and 'gavagai' $\wedge f$ are spelt out such that *any* undetached rabbit part smaller than a whole enduring rabbit satisfies the argument.

However, I contend, we need not cluster *all* undetached rabbit parts under the same semantic theory. Rabbit claws, feet, legs and heads are undetached parts of rabbits. But we can differentiate among them, and articulate semantic theories whose axioms deal with those anatomical parts separately. In this way, 'gavagai', under one particular scheme, might be taken to divide its reference over undetached legs of rabbits, for instance; under another scheme, over undetached tails of rabbits; and so forth. Unfortunately, were the semanticist to specify *which* particular anatomical part of a rabbit her scheme makes use of, it would be fairly easy for the anti-Quinean to rebut the proposal. Simply by pointing; for even though every time you point to a rabbit, you are pointing to an undetached rabbit part, you need not point to, say, its leg in every occasion. Therefore, the semanticist will be able to discard, on inductive grounds *a* particular undetached rabbit part as the target of the native's ostensive behaviour. Nevertheless, there is a better option available to the Quinean.

In what follows, I shall propose a particular way to discriminate among schemes of reference denoting diverse undetached rabbit parts that is not subject to the aforementioned difficulties. We may talk in terms of the percentage of the whole rabbit, including the percentage of its surface, that each

scheme assigns as the extension of ‘gavagai’. In this way, one putative perverse scheme may claim that ‘gavagai’ divides its reference over 5% of the whole rabbit, including 5% of its surface (henceforth abbreviated 5%-urp: i.e., 5% undetached rabbit part). Another scheme over 20%-urp, another over 80%-urp, and so on. Notice that pointing cannot help to solve the referential indeterminacy. Every time you point to a rabbit, you are pointing to a 5%-urp, to a 20%-urp, to an 80%-urp, etc. We may then take natives’ assent to/dissent from any given query as evidence in favour of a ‘x%-urp’ scheme, as opposed to the standard one (although see below).

Let’s see how some semantic theories that cut the reference of ‘gavagai’ over x%-urp can cope with Evans’ white-legged brown rabbit. Take, for instance, a perverse semantic theory that divides the reference of ‘gavagai’ over 5%-urp. Such a theory would include the following axioms:

(a*) (x)(x satisfies ‘gavagai’ iff x is a 5%-urp),

and

(a**) (x)(x satisfies ‘blanco’ $\wedge$ f iff (x is white & x satisfies f)).

Hence, taking the satisfaction conditions for ‘blanco’ in the standard way¹², our putative semantic theory will generate theorem (t):

(t) (x)(x satisfies ‘blanco’[^]‘gavagai’ iff (x is white and x is a 5%- urp)).

However, such a perverse semantic theory would not resist Evans’ attack. A version of Evans’ counter-example would kick in. Think of a brown rabbit which, instead of having a white leg, has 5% of its surface white-coloured. In this case, natives guided by (t) would assent to ‘Blanco gavagai?’ when stimulated by a 5%-white-coloured brown rabbit. Whiteness distributed all over a 5%-urp would not work since it elicits the wrong answer under certain circumstances. Semanticists agreed that natives would not assent to ‘Blanco gavagai?’ unless they are in presence of a white rabbit. And clearly an object which only has 5% of its surface ϕ -coloured does not count as a ϕ -coloured object.

The careful reader may have guessed by now what the next move for the Quinean should be. Evans’ contention about compound expressions such as ‘blanco gavagai’ is that ‘blanco’ has to be distributed in a particular way with respect to the boundaries of the object prompting native’s assent to the query ‘Gavagai?’ The key word is *distribution*. In natural languages, when we say that a rabbit is white, we are assuming that the white feature is distributed more or less uniformly over all the surface of the rabbit. Let’s say that when the percentage of white-coloured surface is equal or bigger than β ,

then we take the rabbit as white¹². Now, my contention is that a perverse scheme that divides the reference of ‘gavagai’ over β %-urp will cope with Evans’ white-legged brown rabbit. Take β for instance as 99%. The perverse theory would then run as follows:

PT4

Axioms:

- (e) $(x)(x \text{ satisfies ‘gavagai’ iff } x \text{ is a } 99\text{-urp})$
 (e1) $(x)(x \text{ satisfies ‘blanco’} \wedge f \text{ iff } (x \text{ is white} \wedge x \text{ satisfies } f))$.

Theorem:

- (e2) $(x)(x \text{ satisfies ‘blanco’} \wedge \text{‘gavagai’ iff } (x \text{ is white} \wedge x \text{ is a } 99\text{-urp}))$.

Now, let’s see how this perverse referential scheme behaves under Evans’ pool of data. The question is: Would the native guided by **PT4** assent to ‘Blanco gavagai?’ when a brown rabbit with a white leg is in his presence? Certainly not, for the native will only assent to the query when the 99% of the surface of the rabbit is white. Hence, Evans’ counter-example is not a counter to **PT4**. Those sympathetic to Evans would have to develop a different version of his counter in which the white portion of the brown rabbit is bigger. But not any bigger portion will do. We require the brown rabbit to have a white part occupying the 99% of its surface. But in this case, we would be confronted with a white rabbit, rather than with a brown one. Therefore, Evans’ example is unable to show that **PT4** misrepresents native usage. A translator guided by this perverse scheme will predict native assent to/dissent from ‘Blanco gavagai?’ in exactly the same sort of situations in which a ‘non-perverse’ translator would. The reason is that rabbits and 99%-urp are *observationally indistinguishable*.

The reader can see that the ‘99%-urp’ scheme differs from **ST** in a non-trivial way. What we need to achieve semantic perversity is a scheme of reference that conforms to all possible evidence, and yet assigns *different* extensions to the native terms from those assigned by **ST**. The following is *a priori*:

$$(x)(y)(x=y \rightarrow (z)(z \text{ is a part of } x \leftrightarrow z \text{ is a part of } y)).$$

This condition establishes the semantic perversity of **PT4**. Since 99 is smaller than 100, there will always be an undetached part of a whole rabbit which does not belong to the 99%-urp: — namely, a 1%-urp. Hence the perversity of **PT4** is *real* in the sense that the set of objects satisfying the property of being white does not coincide with the set of objects contemplated under **ST**.

The indeterminacy, thus, remains unsolved. We haven't got a clue as to whether 'gavagai' divides its reference over rabbits or 99%-urp.

The results obtained by assuming **PT4** are, however, in clear contrast with those obtained via Hookway's strategy, as modified under **PT3**. Both **PT3** and **PT4** can account for Evans' white-legged brown rabbit. Nevertheless, the Quinean has good reasons for preferring the latter theory. If Evans' considerations regarding semantic structure and tacit knowledge were correct, and the neurophysiological evidence were as described above, **PT4** would have a clear advantage over **PT3**. Derivations in **PT4** have exactly the same syntactic structure as derivations in the standard theory, **ST**. Therefore, if the hypothesized neurophysiological data showed that any semantic theory aiming to explain native linguistic behaviour ought to do so by means of non-disjunctive axioms, **PT4** would conform to such a constraint. Whether or not such a constraint is correct is a matter for future research and will clearly depend on what kind of architecture embeds our higher cognitive abilities. However, insofar as the constraint is drawn from a physicalist framework, its bearing is a theoretical possibility that the Quinean cannot ignore. Fortunately, as we've seen, by favouring the '99%-urp' referential scheme, there is nothing the Quinean should fear from Evans' considerations¹³.

*Departamento de Filosofía
Universidad de Murcia
Campus de Espinardo, E-30071, Murcia
E-mail: fjcalvo@um.es*

NOTES

* A version of this paper was presented at the 1997 *University of Glasgow-University of St. Andrews Joint Philosophy Conference* at the Isle of Raasay (Scotland). I am grateful to Bob Hale, Philip Percival and Crispin Wright for helpful comments and suggestions. I am especially indebted to Jim Edwards for many helpful discussions on this and related topics during the four years that I was his student. This research was supported by a grant from *Caja de Ahorros del Mediterráneo* and *The British Council*.

¹ The reasons for this are well-known and I shall not pursue them here. The reader not familiar with them may consult Quine (1960), ch. 2. For a comprehensive review, see Calvo Garzón (1999).

² Many philosophers take Evans' counter-example to have definitely defeated the thesis of referential inscrutability (see, for instance, Kirk (1986), p. 47). A shorter version of my review of Evans' counter-example, and the Quinean strategy I shall be offering to bypass it (see section II below), have appeared in Calvo Garzón (2000a). However, the proposal put forward in this paper is developed for different purposes.

³ In fact, Evans' approach differs from the original project of Radical Translation in more substantial respects. Being concerned with semantics, we need the concepts of truth, denotation, etc. And Evans' approach to such notions must be understood in a full-blooded sense: "[The] semanticist aims to uncover a structure in the language that mirrors the competence speakers of the language have actually acquired." [Evans (1975), pp. 343-4]. We shall see below how Evans tries to exploit this issue to his advantage.

⁴ Although Quine initially employed his parable to illustrate the Indeterminacy of Translation, Referential Inscrutability actually concerns indeterminacy in the Semantic field. By transferring Quine's original formulation into Semantics, we fear no loss: Any theory of Semantics will have to match Native with Home sentences. And in doing so the semanticist relies upon the same body of evidence as the translator does. Namely, native assent to/dissent from queries under concurrent observable circumstances.

⁵ (a2) is obviously a consequence of (a) and (a1). The reader might be expecting that 'theorems' of the standard theory would assign truth to sentences. However, it is simpler to stay with satisfaction for nothing in my ensuing argument hangs on the difference.

⁶ Notice that a brown rabbit's white leg is a white undetached rabbit part.

⁷ Evans' original target was discrediting rival semantic theories which are *extensionally equivalent* (i.e., which deliver the same set of well-formed theorems) to the standard one [Evans (1981)]. We can nevertheless apply Evans's argument to cases where the theories under consideration are not extensionally equivalent, such as ST and PT3.

⁸ This is indeed Evans' original approach. Evans' account of dispositions must be understood in a full-blooded way: 'The decisive way to decide which model is correct is by providing a causal, presumably neurophysiologically based, explanation of comprehension' [Evans (1981), p. 127]. And notice that in Quine's view this is the correct level of analysis: 'To cite a behavioural disposition is to posit an unexplained neural mechanism, and such posits should be made in the hope of their submitting some day to a physical explanation' [Quine (1975), p. 95].

⁹ Cf. Calvo Garzón (2000a, 2000b).

¹⁰ For present purposes I shall ignore Quine's *coarser* cuts. The reader may care to consult Evans (1975). Wright (1997) offers a critical appraisal of all the different Quinean proposals (both finer and coarser) and of Evans' counters to all of them.

¹¹ Note that (a**) coincides with (a1) — i.e., the axiom employed by the standard theory, ST.

¹² I can set up the example in terms of percentage-of-*surface* (rather than volume) since we are restricting our attention to highly observational features such as 'colour' which applies to the external surface of objects. Notice, however, that since the 'x%-urp' scheme was defined in terms of x% of *whole* objects, including x% of their surfaces, we could bypass putative versions of Evans' counter that exploited volume features — like mass.

¹³ Crispin Wright (1997) has recently produced two arguments against Quine's Inscrutability Thesis, based on 'structural' and 'psychological' simplicity, respectively, which may jeopardize the perverse semantic route offered in this paper. For a rejoinder to Wright's arguments see Calvo Garzón (2000a; under review).

REFERENCES

- Calvo Garzón, F. (1999), *A Connectionist Defence of the Inscrutability Thesis and the Elimination of the Mental*, unpublished doctoral dissertation, University of Glasgow.
- (2000a), “A Connectionist Defence of the Inscrutability Thesis”, *Mind and Language*, vol. 15, pp. 465-80.
- (2000b), “State Space Semantics and Conceptual Similarity: Reply to Churchland”, *Philosophical Psychology*, vol. 13, pp. 77-95.
- (under review), “Is Simplicity Alethic for Semantic Theories?”.
- EVANS, G. (1975), “Identity and Predication”, *The Journal of Philosophy*, vol. 72, pp. 343-63.
- (1981), “Semantic Theory and Tacit Knowledge”, in Holtzmann, S. and Leich, C. (eds.) (1981), *Wittgenstein: To Follow a Rule*, London, Routledge & Kegan Paul, pp. 118-37.
- HOOKEY, C. (1988), *Quine*, Oxford, Polity Press.
- KIRK, R. (1986), *Translation Determined*, Oxford, Oxford University Press.
- QUINE, W.V.O. (1960), *Word and Object*, Cambridge, Mass., MIT Press.
- (1974), *The Roots of Reference*, La Salle, Ill., Open Court.
- (1975), “Mind and Verbal Dispositions”, in Guttenplan, S. (ed.) (1975), *Mind and Language*, Oxford, Oxford University Press, pp. 83-95.
- WRIGHT, C. (1997), “The Indeterminacy of Translation”, in Hale, B. and Wright, C. (eds.) (1997), *A Companion to the Philosophy of Language*, Oxford, Blackwell, pp. 397-426.