

Sep 11, 2025

Possible Worlds: A Logical Model of Coherence between Identity and Veracity

DOI

<https://dx.doi.org/10.17504/protocols.io.n92ld6wzxcg5b/v1>

Omar Pavel Ancka Quispe¹, omarancka²

¹Facultad de Ingeniería Eléctrica y Electrónica, Escuela de Ingeniería Electrónica, Universidad Nacional del Callao, Callao, Perú;

²Universidad Nacional del Callao

Omar Ancka Quispe



omarancka

Create & collaborate more with a free account

Edit and publish protocols, collaborate in communities, share insights through comments, and track progress with run records.

Create free account

OPEN  ACCESS



DOI: <https://dx.doi.org/10.17504/protocols.io.n92ld6wzxcg5b/v1>

External link: <https://philarchive.org/rec/ANCMPU>

Protocol Citation: Omar Pavel Ancka Quispe, omarancka 2025. Possible Worlds: A Logical Model of Coherence between Identity and Veracity. protocols.io <https://dx.doi.org/10.17504/protocols.io.n92ld6wzxcg5b/v1>

License: This is an open access protocol distributed under the terms of the [Creative Commons Attribution License](#), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited

Protocol status: In development

We are still developing and optimizing this protocol

Created: September 11, 2025

Last Modified: September 11, 2025

Protocol Integer ID: 227059

Keywords: Logical semantics, Model theory, Semantic evaluation, Ontological identity, Veracity and coherence, Agent logic, Truth-tellers and liars, Possible worlds semantics, logical model of coherence, agent logic, applications in agent logic, ontological identity, finite tabular semantic, ontological identity to each individual, semantic operator, epistemological analysis, statements in finite world, tabular semantic, semantic model, coherence, logical model, coherence between identity, truth value, global consistency, evaluation of the global consistency, classic problems of truth, construction of all possible world, agent, formal structure, statement, finite world, consistency rule, veracity, possible world, explicit alternative to classical inferential method, context

Abstract

This article presents a general logical-semantic model that allows the evaluation of the global consistency between ontological identities and statements in finite worlds of individuals. Although the model is inspired by the classic problems of truth-teller and liar — such as those found in the so-called “islands of knights and knaves”-, its formal structure allows for a broader application in contexts where it is necessary to analyze the coherence between what an agent is and what it says. The model is based on assigning an ontological identity to each individual (e.g., honest or liar, reliable or faulty sensor) together with a statement that expresses a proposition. Two semantic operators are established: one representing the agent’s identity and another evaluating the truth value of its statement in a possible world. From these assignments, a finite tabular semantics is defined that allows the construction of all possible worlds and the application of a consistency rule that filters those worlds which present coherence between identity and statement. This approach offers a formal and explicit alternative to classical inferential methods, allowing applications in agent logic, diagnosis, model theory and epistemological analysis.

Attachments



Possible Worlds - A ...

403KB

Guidelines

1. Introduction

The problems of truth-teller and liar have been widely used both in the teaching of logic and in the exploration of epistemological principles about truth, reference, and consistency. In these scenarios, one considers a set of individuals who may be truthful (who always tell the truth) or liars (who always lie), and whose statements involve references to themselves or to others.

Although these problems can be formalized within the framework of propositional or predicate logic, they generally require the assumption of ad hoc premises or axioms to be solved. In contrast, this paper proposes a new semantic model that does not require prior assumptions about the veracity of statements or the use of deductive inference rules, but rather exhaustively explores the space of possible configurations.

The model introduces two operators: one that determines the ontological identity of an individual, and another that interprets the truth content of his statement. By a systematic tabulation over all possible combinations of identities in the Universe $\mathcal{M}$, and an evaluation rule based on the consistency between what an agent is and what it says, worlds with inconsistent configurations are eliminated and those satisfying where the semantic matching rule are preserved. This methodology, which we call Tabular Semantic Model of Coherence between Identity and Veracity, provides a rigorous, intuitive and extensible tool for the analysis of such logical problems.

2. Formal Definitions of the Model

2.1 Universe and ontological classes

- Let $U = \{X_1, X_2, \dots, X_n\}$ be the finite set of individuals involved in the problem. Each individual X_i is assigned an ontological identity, for example: $\sigma(X_i) \in \{H, M\}$, given by: $\sigma : U \rightarrow \{H, M\}$
- H represents an Honest individual (always tells the truth). M represents a Liar individual (always lies).
- $\sigma(X_i) \Rightarrow \sigma \in M$ represents a complete assignment of identities to all individuals in U , which we call possible world.
- Define the set of all possible identity assignments (possible worlds) as the Universe $\mathcal{M}$, given by: $\mathcal{M} := \{\sigma \mid \sigma : U \rightarrow \{H, M\}\}$

2.2 Truth value according to identity and evaluation

- Definition 1: Truth value according to its Identity. For $\sigma \in M$ and $X_i \in U$, the truth value associated with the identity of X_i is: $v(\sigma(X_i)) = \{V \text{ if } \sigma(X_i) = H; F \text{ if } \sigma(X_i) = M\}$
- Definition 2: Truth value of the Statement issued by X_i . Let ϕ_{X_i} be the specific statement issued by individual X_i , and let $\sigma \in M$ be a possible world. The truth value of the statement evaluated at σ is: $v(\phi_{X_i}, \sigma) \in \{V, F\}$.

2.3 Semantic coherence criterion

- Semantic consistency rule: A possible world $\sigma \in M$ is coherent with the statement X_i if and only if: $v(\sigma(X_i)) = v(\varphi_{X_i}, \sigma)$
- That is, the truth value corresponding to the identity of X_i coincides with the truth value of its statement evaluated within the same possible world.

2.3.1 Definition of valid world

- Valid world: A possible world $\sigma \in M$ is a valid world for the given problem if it is consistent for all individuals who have made explicit statements.

2.4 Model operators

- $\sigma : U \rightarrow \{H, M\}$: Identity assignment (possible world). After the assignment, $\sigma(X_i)$ is the individual identity for a particular individual.
- $v(\sigma(X_i))$: Truth value assigned according to identity.
- $v(\varphi_{X_i}, \sigma)$: Truth value of the statement in the model.
- The model can be extended to k identities, e.g., $\sigma : U \rightarrow \{H, M, S, R\}$, producing k^n possible worlds and universe $\mathcal{M} = \{\sigma_j\}_{j=1}^{k^n}$.

2.5 General procedure

The general logical procedure for solving a problem is summarized in the following steps:

1. Enumerate all possible worlds $\sigma \in \mathcal{M}$, where each world assigns to each individual in the set U an identity (e.g., honest or liar).
2. For each statement φ_{X_i} uttered by some individual $X_i \in U$, evaluate its truth value $v(\varphi_{X_i}, \sigma)$ in every possible world σ . That is, ask yourself: "Would this statement be true if the world were σ ?"
3. Calculate in each world σ the value associated with the identity of the individual X_i , that is, $v(\sigma(X_i))$, where: $v(\sigma(X_i)) = \{ V, \text{ if } \sigma(X_i) = H; F, \text{ if } \sigma(X_i) = M \}$
4. Check whether the consistency rule is met for each individual who made a statement: $v(\varphi_{X_i}, \sigma) = v(\sigma(X_i))$. This ensures that X_i 's statement matches its identity in the world σ .
5. Keep only those worlds $\sigma \in \mathcal{M}$ in which all the speakers' statements are consistent with their identity. These worlds are the valid worlds of the problem. This consistency check can be done either simultaneously for all statements (evaluating in parallel), or step-by-step, evaluating one statement at a time and discarding inconsistent worlds at each step. In the sequential approach, one starts with all possible worlds and progressively filters according to the consistency of each individual statement.

3. Examples of Application of the Model

3.1 Example of use 1: Honest and lying agents

Consider a three-person universe: $U = \{A, B, C\}$, where each individual can be either honest (H) or liar (M). The assignment of identities is represented by a function:

$\sigma : U \rightarrow \{H, M\}$

and each distinct assignment constitutes a possible world. Since there are two options for each of the three individuals, there are $2^3 = 8$ possible worlds, and thus, the set of possible worlds M is constructed, which is presented in the following table. (Universe M table with $\sigma_1.. \sigma_8$ assignments H/M for A, B, C .)

Proposals issued

Suppose individuals make the following propositions:

- A says: "B is a liar"
- B says: "C is honest"
- C says: "A is a liar"

Each proposition φ_i is assigned a truth value according to the world σ , evaluating whether what it says is true or not. This is represented as: $v(\varphi_i, \sigma) \in \{V, F\}$.

Evaluation in a specific world

Let's look at the full procedure for the world σ_4 , where:

$$\sigma_4(A) = H, \quad \sigma_4(B) = M, \quad \sigma_4(C) = M$$

- A says: "B is a liar"
 - In σ_4 , B is really a liar: the proposition is true $\Rightarrow v(\varphi_A, \sigma_4) = V$
 - A is honest $\Rightarrow v(\sigma(A)) = V$
 - They match $\Rightarrow$ consistently
- B says: "C is honest"
 - In σ_4 , C is a liar: the proposition is false $\Rightarrow v(\varphi_B, \sigma_4) = F$
 - B is a liar $\Rightarrow v(\sigma(B)) = F$
 - They match $\Rightarrow$ consistently
- C says: "A is a liar"
 - In σ_4 , A is honest: the proposition is false $\Rightarrow v(\varphi_C, \sigma_4) = F$
 - C is a liar $\Rightarrow v(\sigma(C)) = F$
 - They match $\Rightarrow$ consistently

Since all statements are consistent with each individual's identity, the world σ_4 is a coherent world.

Filtering inconsistent worlds

To solve the problem, all worlds $\sigma \in M$ are evaluated. Those that contain any inconsistency are discarded, that is, for which:

$$v(\varphi_i, \sigma) \neq v(\sigma(X_i)) \quad \text{for some } i$$

The remaining worlds are the valid worlds. In many classical problems, only one survives, thus determining the identity of each individual.

Use Example 2: Reliable or faulty sensors

Suppose a set of digital sensors is given by $U = \{S1, S2\}$. Each sensor can have one of two possible functional identities:

- C: Reliable sensor (it tells the truth: its output reflects reality).
- D: Faulty sensor (it lies: its output is always opposite to reality).

Possible worlds

Each sensor can be either C or D, therefore there are $2^2 = 4$ possible worlds:

M	$\sigma(S1)$	$\sigma(S2)$
σ_1	D	D
σ_2	D	C
σ_3	C	D
σ_4	C	C

Assertions and sensor output

Suppose both sensors are trying to measure the same binary quantity (e.g., whether an object is present or not), and that:

- The real value of the world is 1 (present object).
- The observed outputs are:
 - S1 emits: 0
 - S2 emits: 1

So, for each world σ_j , we evaluate whether the sensor output matches (in the reliable case) or contradicts (in the faulty case) reality. We denote the truth value of the sensor's statement as:

$$v(\varphi\{S_i\}, \sigma_j) = \{1 \text{ if the sensor output matches its type (C/D) in } \sigma_j, 0 \text{ otherwise.}\}$$

For example:

- If $\sigma_j(S1) = D$, then your output is expected to be $\neg 1 = 0$, and since you output 0, then $v(\varphi_{\{S1\}}, \sigma_j) = 1$.
- If $\sigma_j(S2) = C$, output 1 is expected, and since it gave 1, then $v(\varphi_{\{S2\}}, \sigma_j) = 1$.



Let us evaluate in all worlds:

M	Inputs	Output	To Compare
	$\sigma(S1)$	$\sigma(S2)$	$v(\varphi\{S1\}, \sigma j) \mid v(\varphi\{S2\}, \sigma j)$
$\sigma1$	D	D	1 0
$\sigma2$	D	C	1 1
$\sigma3$	C	D	0 0
$\sigma4$	C	C	0 1

Consistency evaluation

We apply the consistency rule to keep only worlds where both outputs are consistent with their sensor type:

$$v(\varphi_{\{S_i\}}, \sigma) = v(\sigma(S_i))$$

Since we used 1 to mean "identity-consistent statement", we apply:

- $\sigma2$ is the only world where both statements are consistent.

Conclusion

The valid world is:

$$\sigma2 = \{S1 = D, S2 = C\}$$

This means that:

- S1 is faulty (it outputs 0, as opposed to the actual 1).
- S2 is reliable (it issued 1, which matches reality).

This example shows how the model can be applied to sensor diagnostics by filtering possible worlds based on the consistency between functional identity (C or D) and observable behavior.

4. Step-by-step practical example: Three inhabitants of an island of truth-tellers and liars

You are on an island where the inhabitants can be either honest (H) or liars (M). You meet three people: A, B, and C. Each makes a statement:

- A says, "I am..." but you can't hear it.
- B says: "A said he's a liar."
- C says: "B is lying; he is a liar."

We want to determine who is H and who is M using the logical model of consistency between identity and truthfulness.

Step 1: Table of possible worlds

Each person can be honest (H) or a liar (M), so there are $2^3 = 8$ possible worlds. We call each world σ_j .

M	$\sigma_j(A)$	$\sigma_j(B)$	$\sigma_j(C)$
σ_1	H	H	H
σ_2	H	H	M
σ_3	H	M	H
σ_4	H	M	M
σ_5	M	H	H
σ_6	M	H	M
σ_7	M	M	H
σ_8	M	M	M

Step 2: Evaluating A's proposition

From what is heard from A, and from the statement of B, it can be established that the statement of A must be modeled as a verbal self-assignment of identity, in the form:

$\phi_A := \text{"I am } \sigma_{\text{verbal}(A)} \text{"}, \text{ where } \sigma_{\text{verbal}(A)} \in \{H, M\}$

This statement is interpreted as a proposition about the identity of A that must be evaluated under each possible world $\sigma \in M_A$. Since in our system honest agents always tell the truth and liars always lie, we have the criteria:

- $v(\phi_A, \sigma_m)$: truth value of the specific proposition issued by A, evaluated in the possible world σ_m of MA .

$v(\phi_A, \sigma) = \{$
 $\quad V \text{ if } \sigma(A) = H \text{ and } \phi_A := \text{"I am H"}$
 $\quad F \text{ if } \sigma(A) = H \text{ and } \phi_A := \text{"I am M"}$
 $\quad F \text{ if } \sigma(A) = M \text{ and } \phi_A := \text{"I am H"}$
 $\quad V \text{ if } \sigma(A) = M \text{ and } \phi_A := \text{"I am M"}$
 $\quad \}$

(before_start section retained from earlier pages)

Before start

Section 2: Formal Definitions of the Model

2.1 Universe and ontological classes

- Let $U = \{X_1, X_2, \dots, X_n\}$ be the finite set of individuals involved in the problem. Each individual X_i is assigned an ontological identity, for example: $\sigma(X_i) \in \{H, M\}$, given by: $\sigma : U \rightarrow \{H, M\}$
- H represents an Honest individual (always tells the truth). M represents a Liar individual (always lies).
- $\sigma(X_i) \Rightarrow \sigma \in M$ represents a complete assignment of identities to all individuals in U , which we call possible world.
- Define the set of all possible identity assignments (possible worlds) as the Universe $\mathcal{M}$, given by:

$$\mathcal{M} := \{\sigma \mid \sigma : U \rightarrow \{H, M\}\}$$

2.2 Truth value according to identity and evaluation

- Definition 1: Truth value according to its Identity. For $\sigma \in \mathcal{M}$ and $X_i \in U$, the truth value associated with the identity of X_i is: $v(\sigma(X_i)) = \{V \text{ if } \sigma(X_i) = H; F \text{ if } \sigma(X_i) = M\}$
- Definition 2: Truth value of the Statement issued by X_i . Let ϕ_{X_i} be the specific statement issued by individual X_i , and let $\sigma \in \mathcal{M}$ be a possible world. The truth value of the statement evaluated at σ is: $v(\phi_{X_i}, \sigma) \in \{V, F\}$.

2.3 Semantic coherence criterion

- Semantic consistency rule: A possible world $\sigma \in \mathcal{M}$ is coherent with the statement X_i if and only if: $v(\sigma(X_i)) = v(\phi_{X_i}, \sigma)$
- That is, the truth value corresponding to the identity of X_i coincides with the truth value of its statement evaluated within the same possible world.

Step-by-step practical example: Three inhabitants of an island of truth-tellers and liars

- 1 We then write down the results of $v(\varphi_B, \sigma_j)$ and evaluate the consistency of B across worlds $\sigma_1.. \sigma_8$, comparing $v(\varphi_B, \sigma_j)$ (which is False in all worlds) with $v(\sigma_j(B))$ to determine Consistency B. The worlds consistent with A and B are: $\{\sigma_3, \sigma_4, \sigma_7, \sigma_8\}$.
- 2 Step 4: Evaluating C's proposition — C said: "B is a liar", that is $\varphi_C := (\sigma_C(B) = M)$. Evaluate φ_C in the remaining worlds:

- 3

$\mathcal{M}$	Inputs		Output	To Compare	Consistency C
	$\sigma_j(B)$	$\sigma_j(C)$	$v(\varphi_C, \sigma_j)$	$v(\sigma_j(C))$	
σ_3	<i>M</i>	<i>H</i>	<i>V</i>	<i>V</i>	✓
σ_4	<i>M</i>	<i>M</i>	<i>V</i>	<i>F</i>	×
σ_7	<i>M</i>	<i>H</i>	<i>V</i>	<i>V</i>	✓
σ_8	<i>M</i>	<i>M</i>	<i>V</i>	<i>F</i>	×

Fig. 1: Evaluation table for $v(\varphi_C, \sigma)$ across the remaining worlds $\{\sigma_3, \sigma_4, \sigma_7, \sigma_8\}$ (placeholder figure from the document).

- 4 but it is observed that the coherent worlds σ_3 and σ_7 are the same for B and C. So we finally get for A, B and C: $\sigma_{\{3,7\}} : A = \{H, M\}, B = M, C = H$
- 5 Result — The only possible world compatible with all the statements is: $A = \{H, M\}, B = M, C = H$. The model has filtered out the inconsistent worlds, leaving only one that satisfies the coherence between the identity of individuals and the value of their statements.

Protocol references

[1] Ronald Fagin, Joseph Y. Halpern, Yoram Moses, and Moshe Y. Vardi. Reasoning About Knowledge. The MIT Press, Cambridge, Massachusetts; London, England, First MIT Press paperback edition, 2003. Originally published in 1995.