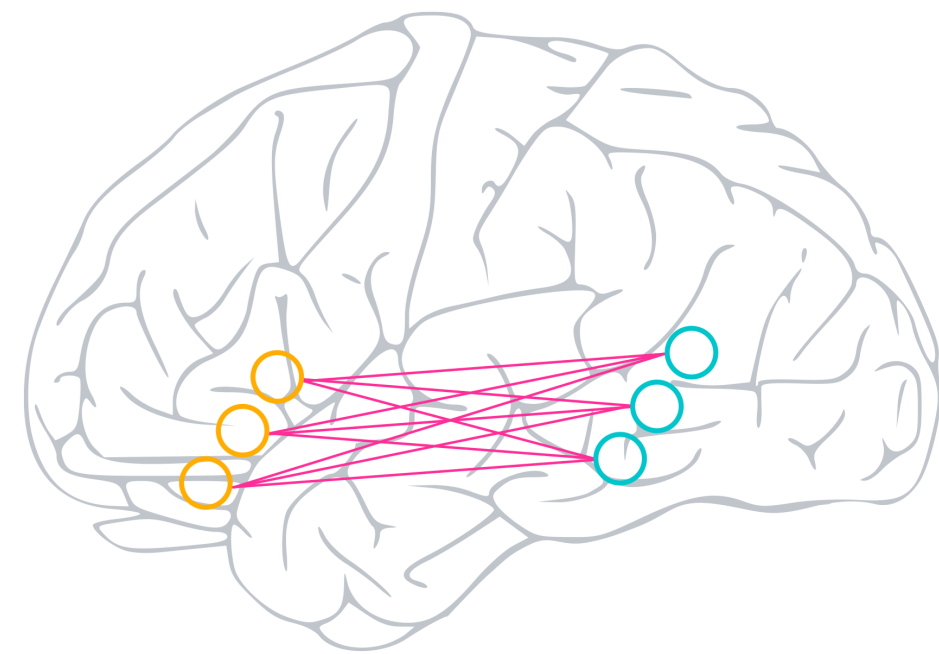




Neural Semantics

Bridging formal and probabilistic approaches to meaning

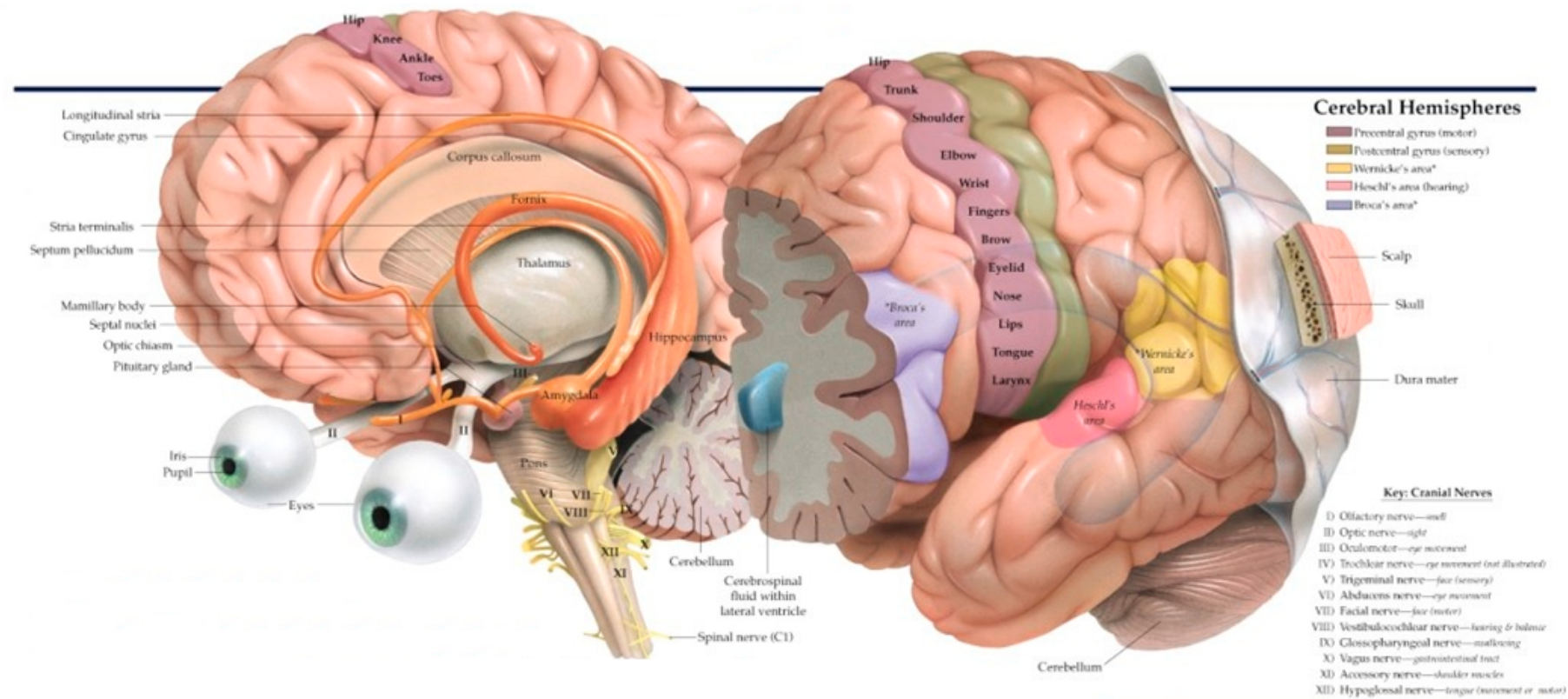
Semantic Theory — June 26th, 2018

Harm Brouwer



UNIVERSITÄT
DES
SAARLANDES

The Greatest Semanticist of them all ...



> Our **language comprehension system** is highly effective and accurate at attributing meaning to unfolding linguistic signal (~word-by-word)

>> This system's **representations** and **computational principles** are implemented in the **neural hardware** of the brain

>>> We should understand meaning construction and representation in terms of “brain-style computation” by identifying a **neural semantics**

Neural Semantics — Requirements

Neural Plausibility: assumed representations and computational principles should be implementable at the neural level [cf. Rumelhart, 1989]

Expressivity: representations should capture necessary dimensions of meaning, such as negation, quantification, and modality [cf. Frege, 1892]

Compositionality: the meaning of complex expressions should be derivable from the meaning of its parts [cf. Partee, 1984]

Gradedness: meaning representations are probabilistic, rather than discrete in nature [cf. Spivey, 2008]

Inferential: The derivation of utterance meaning entails (direct) inferences that go beyond literal propositional content [cf. Johnson-Laird, 1983]

Incrementality: As natural language unfolds over time, representations should allow for incremental construction [cf. Tanenhaus et al., 1995]

Today's lecture

I. A Framework for Neural Semantics

II. A Neural Model of Language Comprehension

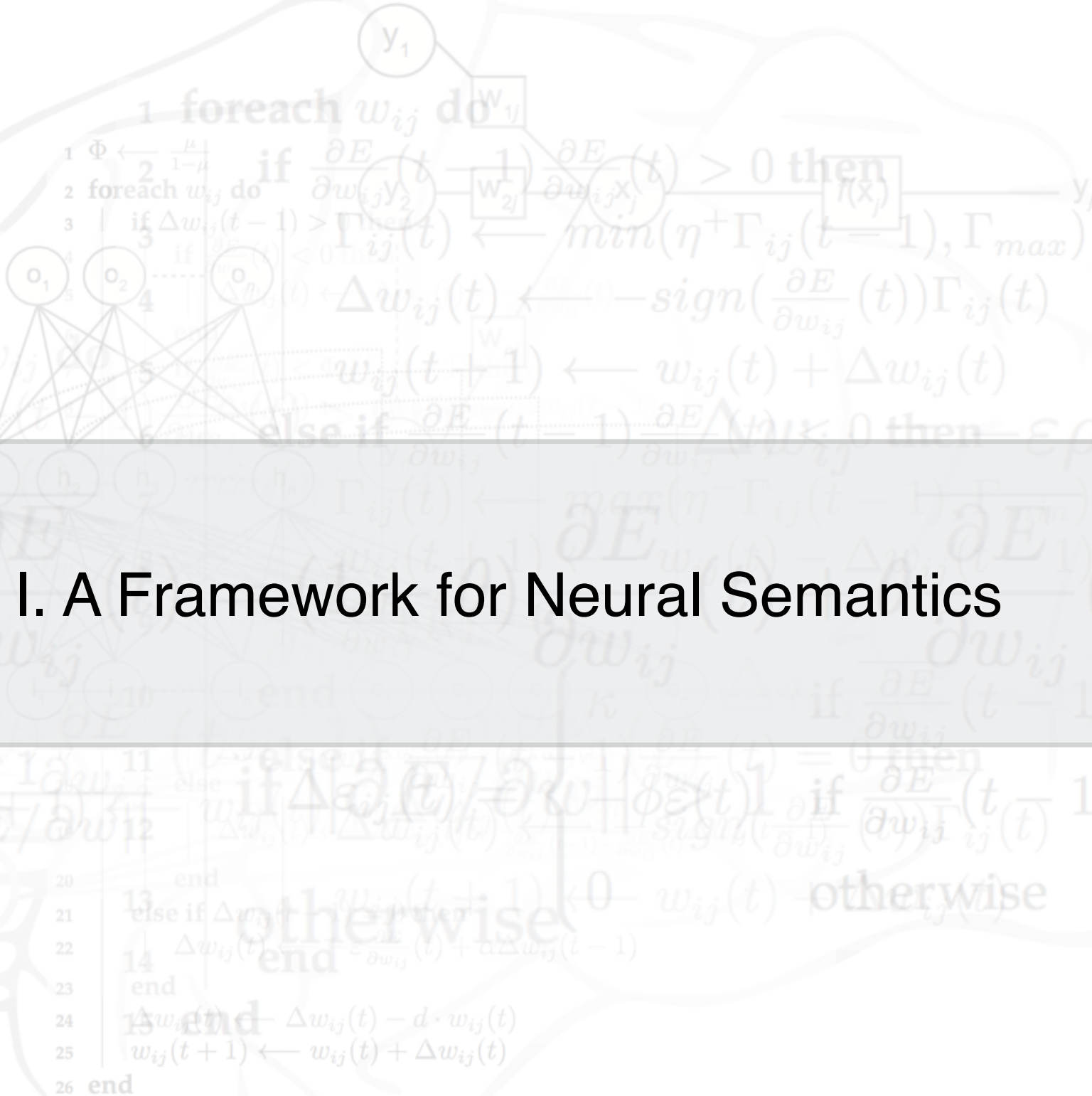
III. Neural Semantics — A fertile approach?

I. A Framework for Neural Semantics

output layer

hidden layer

input layer



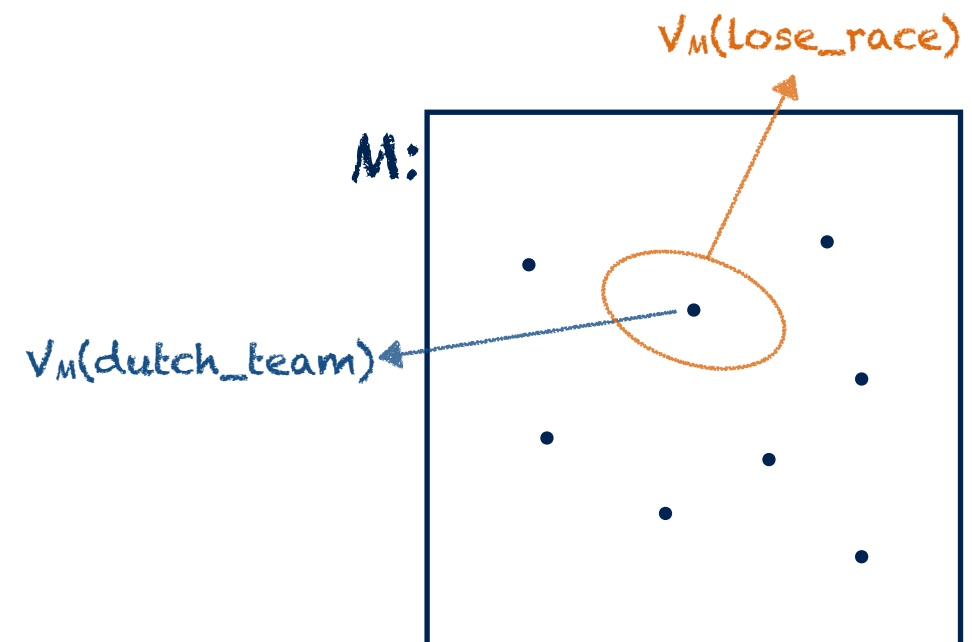
Meaning in formal semantics

Meaning is defined as **truth-conditions** over logical models

$\llbracket \text{The Dutch team lost the race} \rrbracket^{M,g} = 1$

iff $\llbracket \text{lose_race(dutch_team)} \rrbracket^{M,g} = 1$

iff $V_M(\text{dutch_team}) \in V_M(\text{lose_race})$



Intuition: Meaning is defined by the **situations** that satisfy an expression

Cognitively, truth-conditions derive from **experience** with the world

Idea: Model experience as **observations of states-of-affairs** in the world

Experiences as cues to meaning

Observations of SoAs can be captured by the set of logical models $\mathcal{M}$



M_1 : enter(beth,restaurant), order(beth,dinner), eat(beth,dinner),
enter(dave, cinema), order(dave,cola),



M_2 : enter(beth,cinema), order(beth,popcorn), order(dave,cola),
pay(dave), enter(thom, restaurant), leave(thom),



M_3 : enter(beth,cinema), order(beth,popcorn), eat(beth,popcorn),
order(dave,cola), pay(dave), order(thom, cola),

...



M_{150} : enter(beth,cinema), leave(beth), enter(dave,restaurant),
pay(dave), enter(thom, cinema), order(thom, water),

> Each $M \in \mathcal{M}$ provides a **cue** toward the underlying truth-conditions

Experiences as cues to meaning

Observations of SoAs can be captured by the set of logical models $\mathcal{M}$



M_1 : enter(beth,restaurant), order(beth,dinner), **eat(beth,dinner)**,
enter(dave, cinema), order(dave,cola),



M_2 : enter(beth,cinema), order(beth,popcorn), order(dave,cola),
pay(dave), enter(thom, restaurant), leave(thom),



M_3 : enter(beth,cinema), order(beth,popcorn), **eat(beth,popcorn)**,
order(dave,cola), pay(dave), order(thom, cola),

...

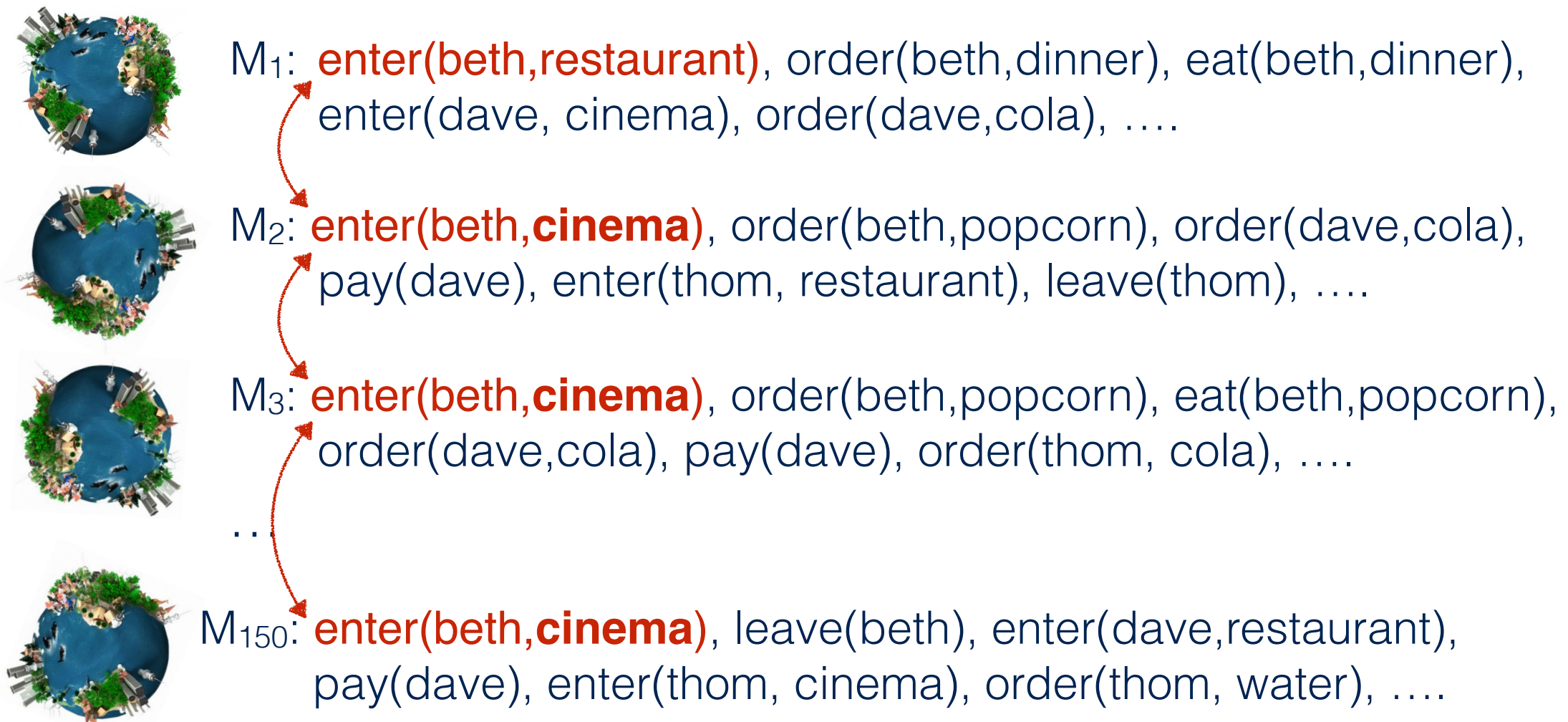


M_{150} : enter(beth,cinema), leave(beth), enter(dave,restaurant),
pay(dave), enter(thom, cinema), order(thom, water),

> Each $M \in \mathcal{M}$ provides a **cue** toward the underlying truth-conditions

Experiences as cues to meaning

Observations of SoAs can be captured by the set of logical models $\mathcal{M}$



> Each $M \in \mathcal{M}$ provides a **cue** toward the underlying truth-conditions

> Together, all $M \in \mathcal{M}$ reflect the world **truth-conditionally** and **probabilistically**

Neural Semantics — Meaning

Formally, we can define the model space $S_{\mathcal{M} \times \mathcal{P}}$ for a set of models $\mathcal{M}$ and propositions $\mathcal{P}$

Rows represent models that reflect observations of ‘states-of-affairs’ in the world

Columns represent meaning vectors of individual propositions

$$\mathbf{v}_i(p) = 1 \text{ iff } \llbracket p \rrbracket^{M_i} = 1$$

Meaning is defined in terms of **co-occurrence**

> Propositions with related meanings will be true in many of the same models

	$p_1 = \text{enter}(\text{beth}, \text{restaurant})$	$p_2 = \text{ask_menu}(\text{beth})$	$p_3 = \text{order}(\text{beth}, \text{cola})$		$p_n = \text{leave}(\text{dave})$
M_1	1	0	0		1
M_2	0	1	1		0
M_3	1	1	1		0
M_4	0	0	0		0
M_5	0	0	1		1
....					
M_m	0	1	1		1

$\mathbf{v}(p_2)$

Neural Semantics — Compositionality

Meaning vectors of **atomic situations** (propositions) are columns in $S_{\mathcal{M} \times \mathcal{P}}$

Meaning vectors of **complex situations** are derived compositionally:

$$\mathbf{v}_i(\neg p) = 1 \text{ iff } \mathbf{v}_i(p) = 0 \text{ for } 1 \leq i \leq m$$

$$\mathbf{v}_i(p \wedge q) = 1 \text{ iff } \mathbf{v}_i(p) = 1 \text{ and } \mathbf{v}_i(q) = 1, \text{ for } 1 \leq i \leq m$$

> which gives functional completeness [e.g., $\mathbf{v}_i(p \vee q) = \mathbf{v}_i(\neg(\neg p \wedge \neg q))$]

Quantification is defined over the combined universe of $\mathcal{M}$: $U_{\mathcal{M}} = \{u_1, \dots, u_k\}$

$$\mathbf{v}_i(\forall x \phi) = 1 \text{ iff } \mathbf{v}_i(\phi[x \setminus u_1] \wedge \dots \wedge \phi[x \setminus u_k]) = 1 \text{ for } 1 \leq i \leq m$$

$$\mathbf{v}_i(\exists x \phi) = 1 \text{ iff } \mathbf{v}_i(\phi[x \setminus u_1] \vee \dots \vee \phi[x \setminus u_k]) = 1 \text{ for } 1 \leq i \leq m$$

Neural Semantics — Probability

Meaning vectors inherently encode **(co-)occurrence probabilities** of situations—which may be atomic (~propositional) or complex (~conjunctive)

> The **prior probability** of situation **a**:

$$P(a) = \sum_i (\mathbf{v}_i(a)) / m$$

> The **conjunctive probability** of situations **a** and **b**:

$$P(a \wedge b) = \sum_i (\mathbf{v}_i(a)\mathbf{v}_i(b)) / m$$

> The **conditional probability** of situation **a** given **b**:

$$P(a | b) = P(a \wedge b) / P(b)$$

	$p_1 = \text{enter}(\text{beth}, \text{restaurant})$	$p_2 = \text{ask_menu}(\text{beth})$	$p_3 = \text{order}(\text{beth}, \text{cola})$	...	$p_n = \text{leave}(\text{dave})$
M_1	1	0	0		1
M_2	0	1	1		0
M_3	1	1	1		0
M_4	0	0	0	...	0
M_5	0	0	1		1
....					
M_m	0	1	1	...	1

Neural Semantics — Inference

How much is situation **a** understood to be the case from situation **b**?

$$\text{comprehension}(a,b) = \begin{cases} [P(a \mid b) - P(a)] / [1 - P(a)] & \text{if } P(a \mid b) > P(a) \\ [P(a \mid b) - P(a)] / P(a) & \text{otherwise} \end{cases}$$

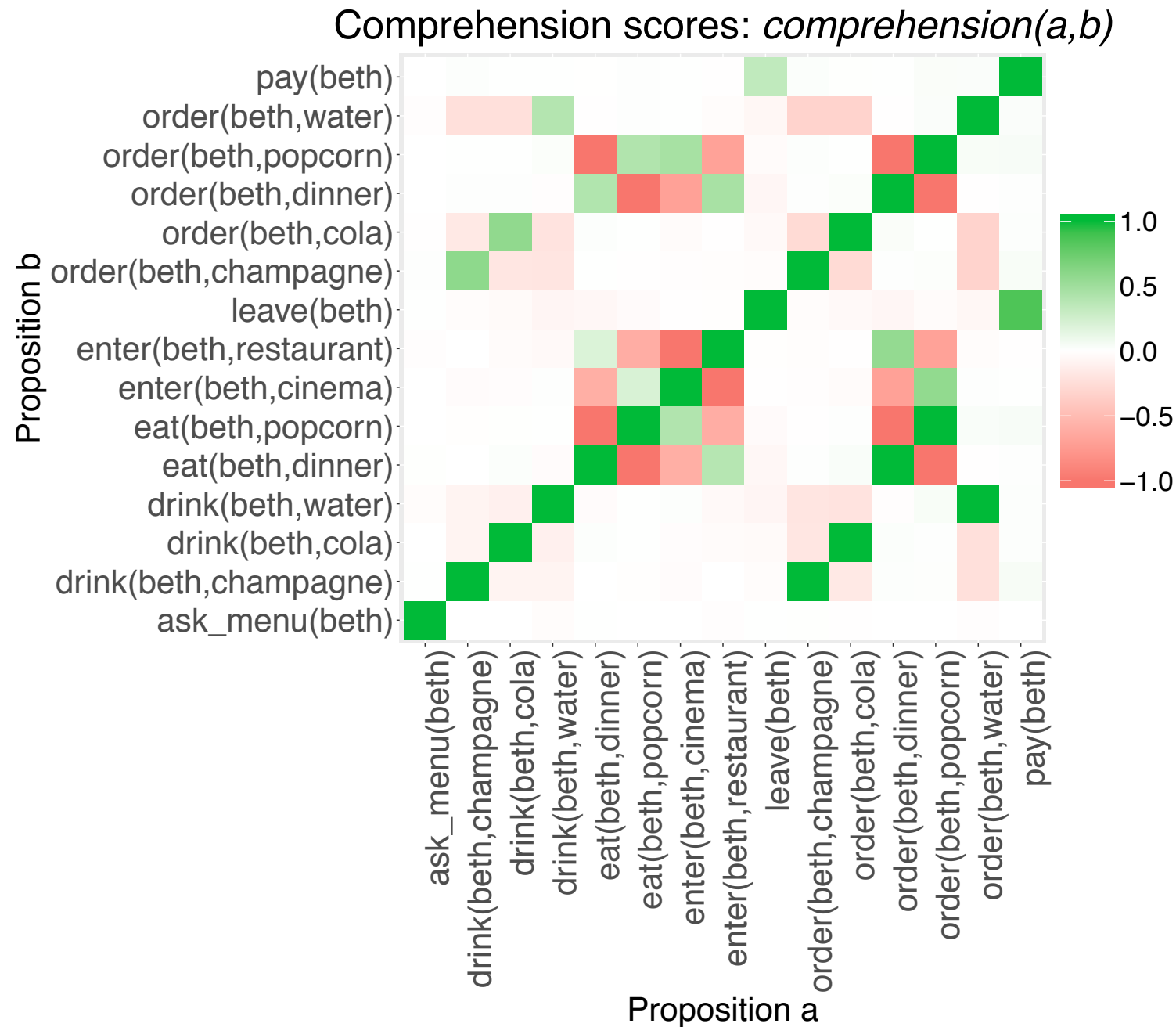
$P(a \mid b) > P(a)$ knowing **b** *increases* belief in **a**
→ **a** is inferred to be the case from **b**

$\text{comprehension}(a,b) = 1$: knowing **b** took away all ‘uncertainty’ in **a** → $b \models a$

$P(a \mid b) \leq P(a)$ knowing **b** *decreases* belief in **a**
→ **a** is inferred **not** to be the case from **b**

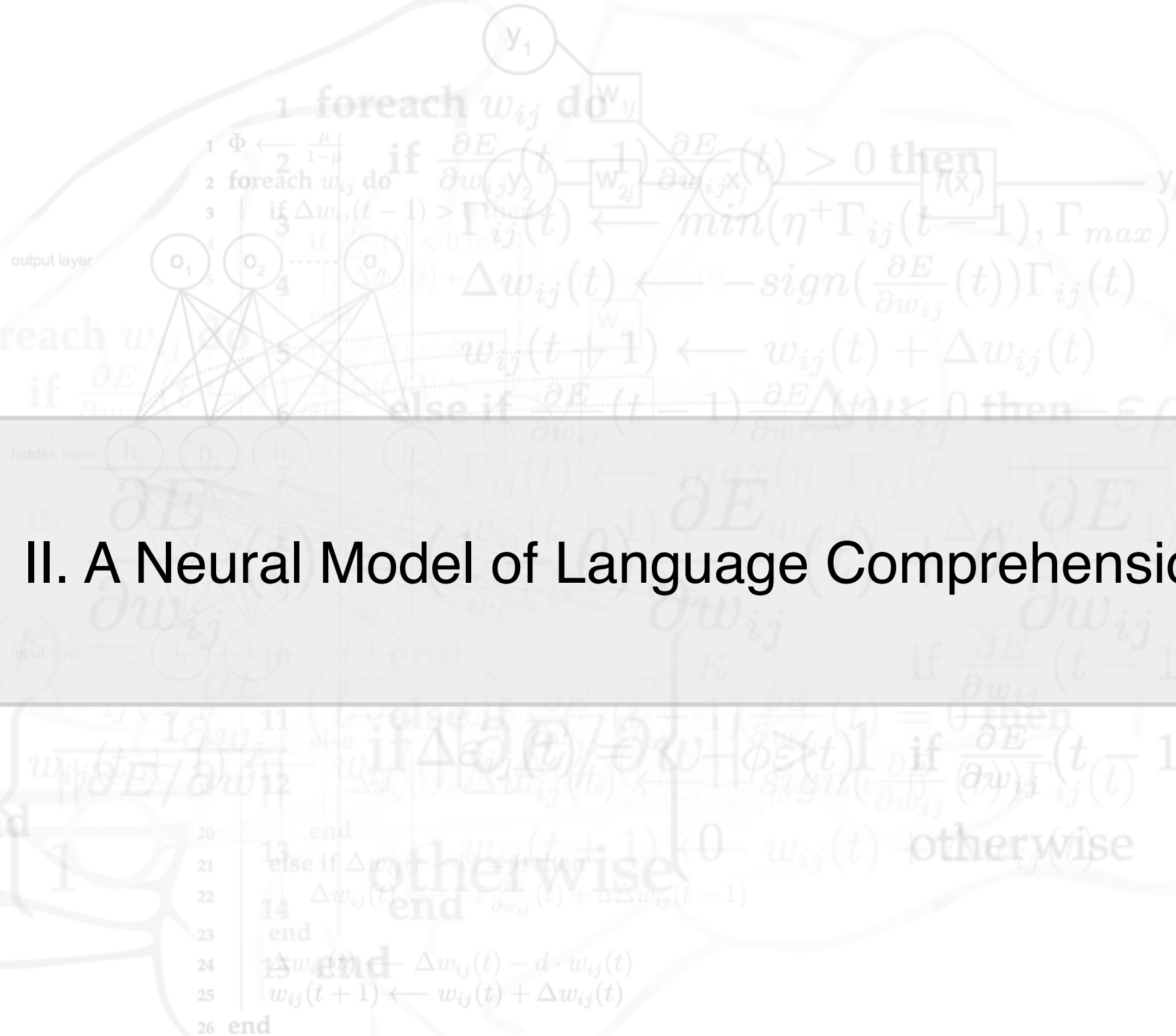
$\text{comprehension}(a,b) = -1$: knowing **b** took away all ‘certainty’ in **a** → $b \models \neg a$

World knowledge inferencing in $S_{\mathcal{M} \times \mathcal{P}}$

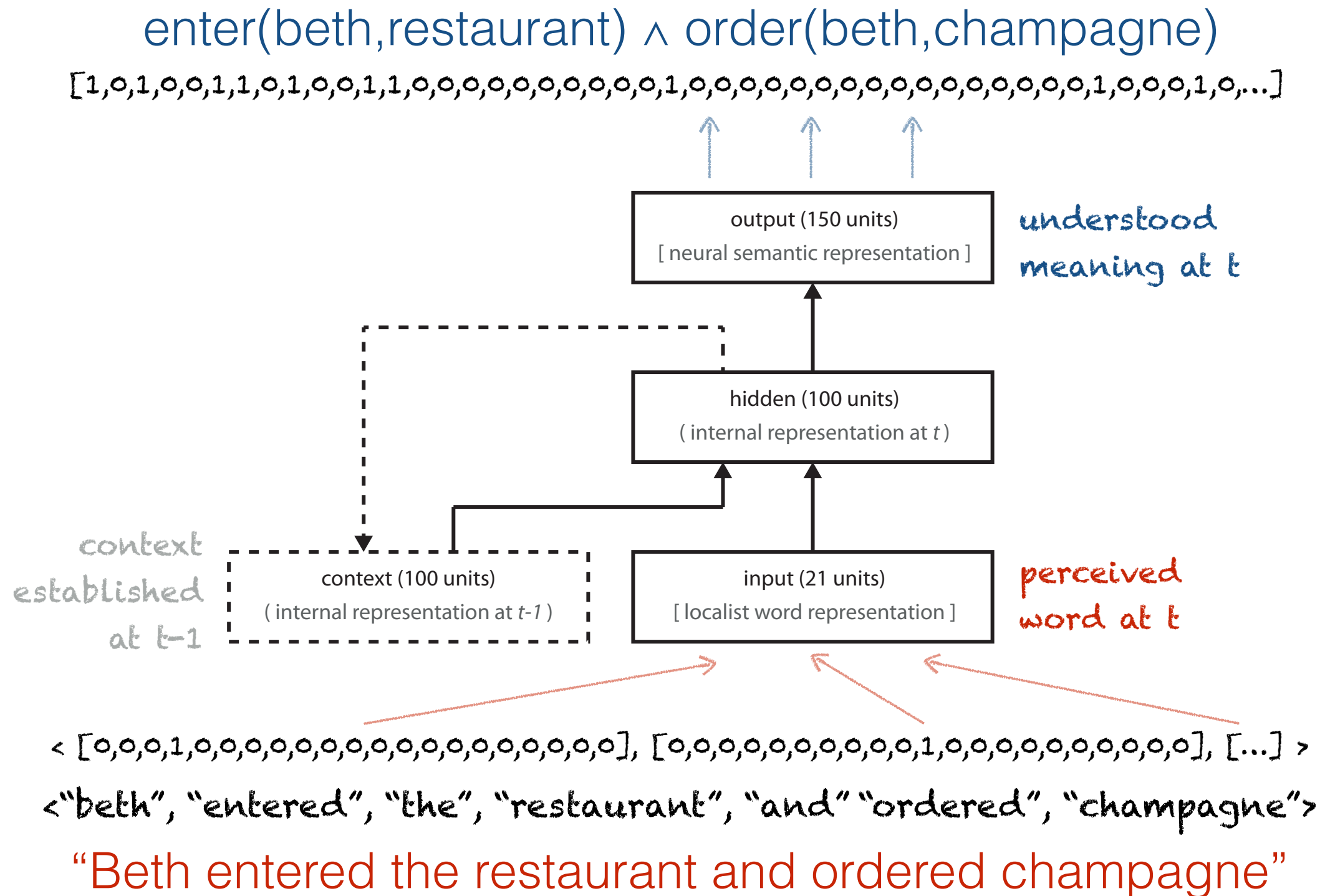


Next: Let's employ $S_{\mathcal{M} \times \mathcal{P}}$ in an incremental comprehension model

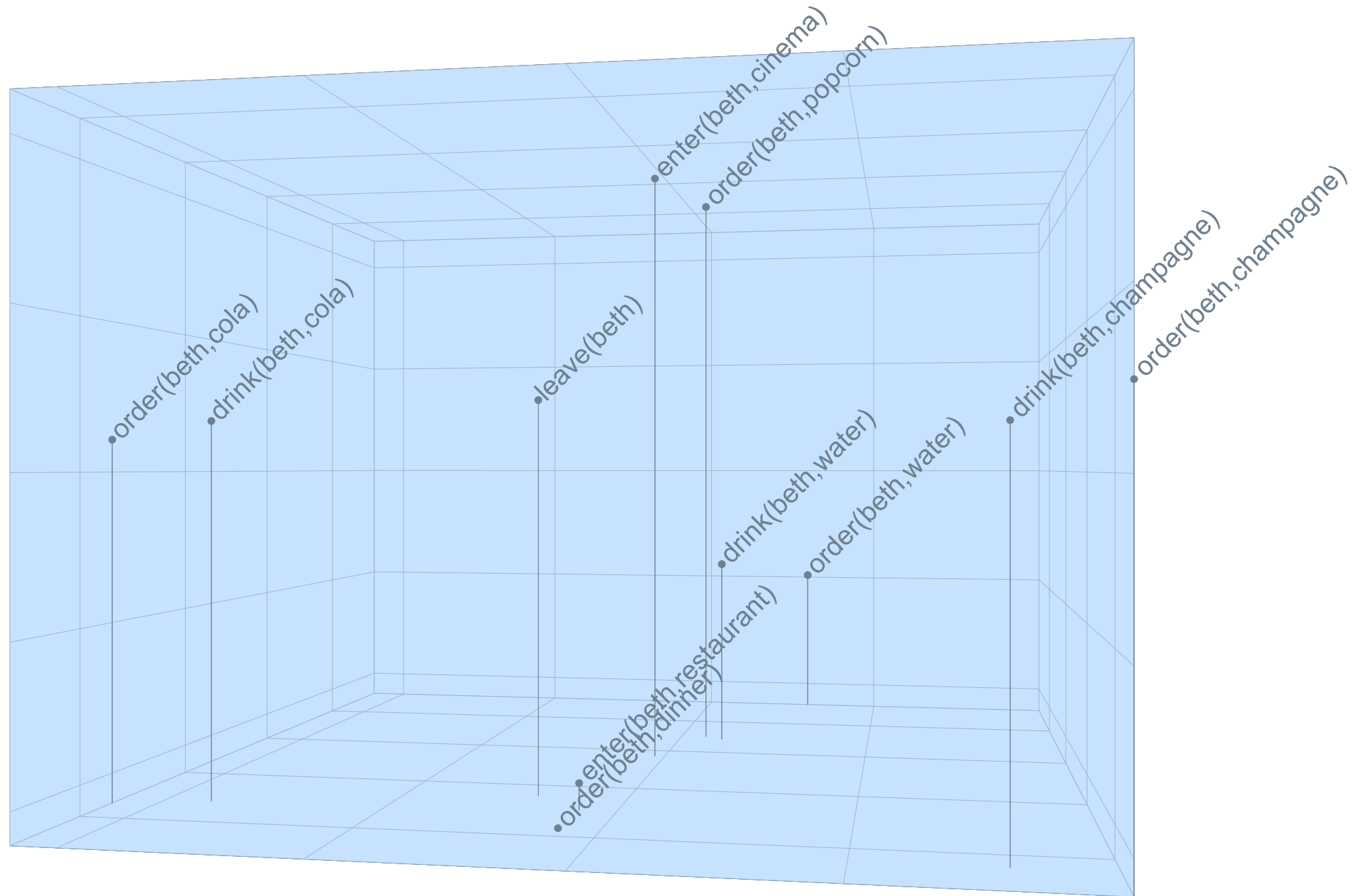
II. A Neural Model of Language Comprehension



A Neural Model of Language Comprehension

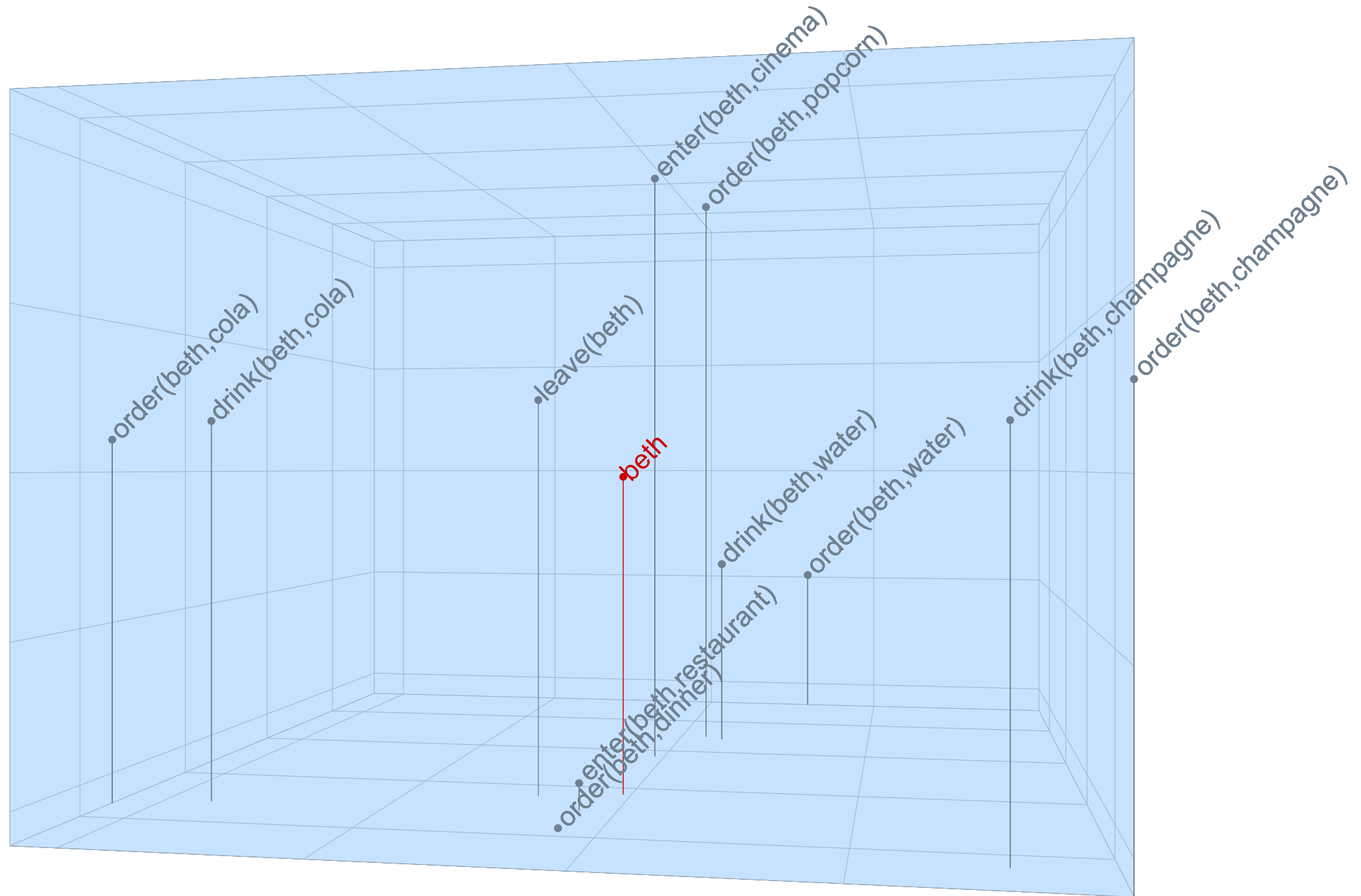


Comprehension is meaning-space navigation



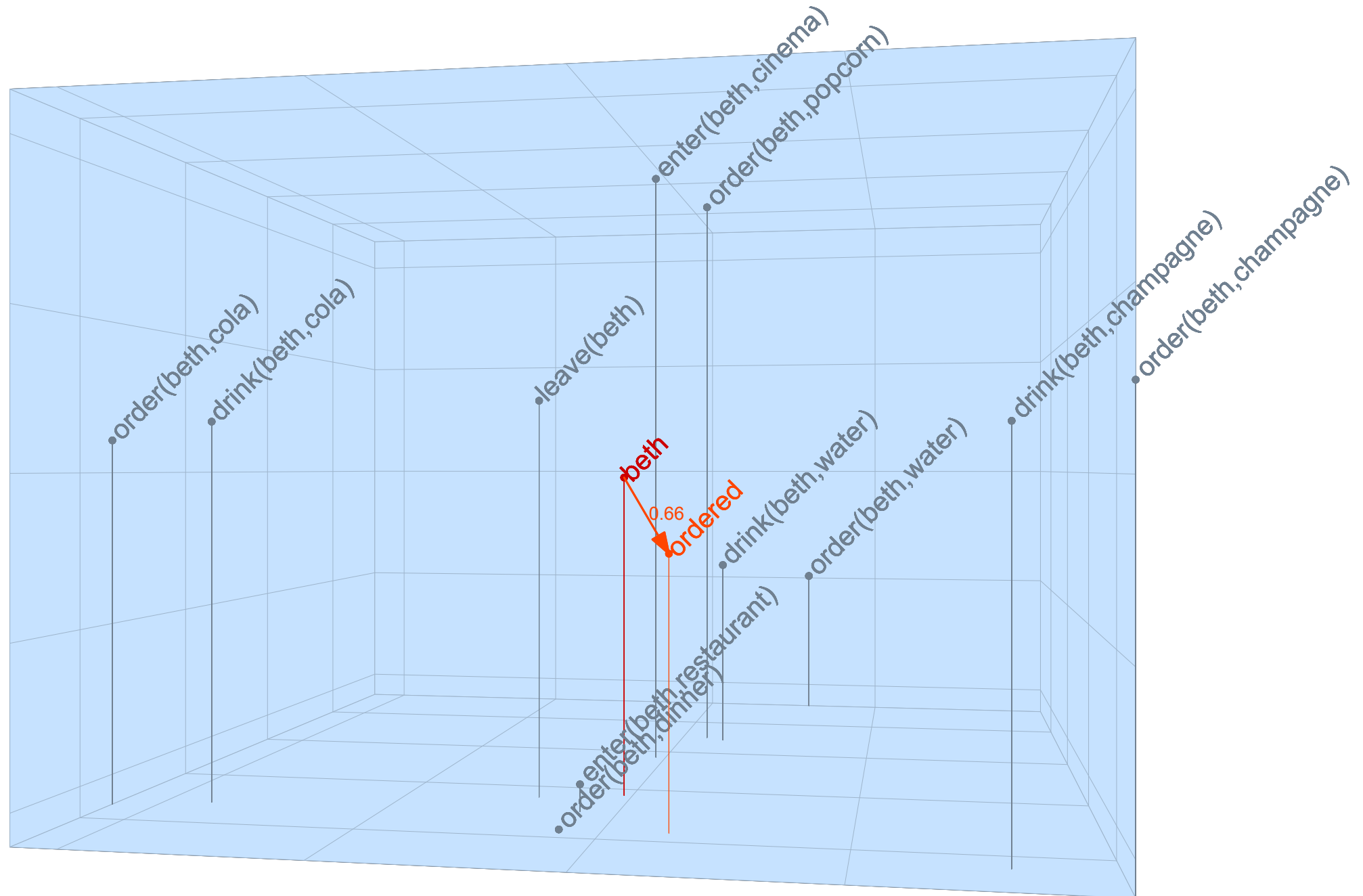
Multi-dimensional scaling: 150D $\mapsto$ 3D

Comprehension is meaning-space navigation



["beth"]

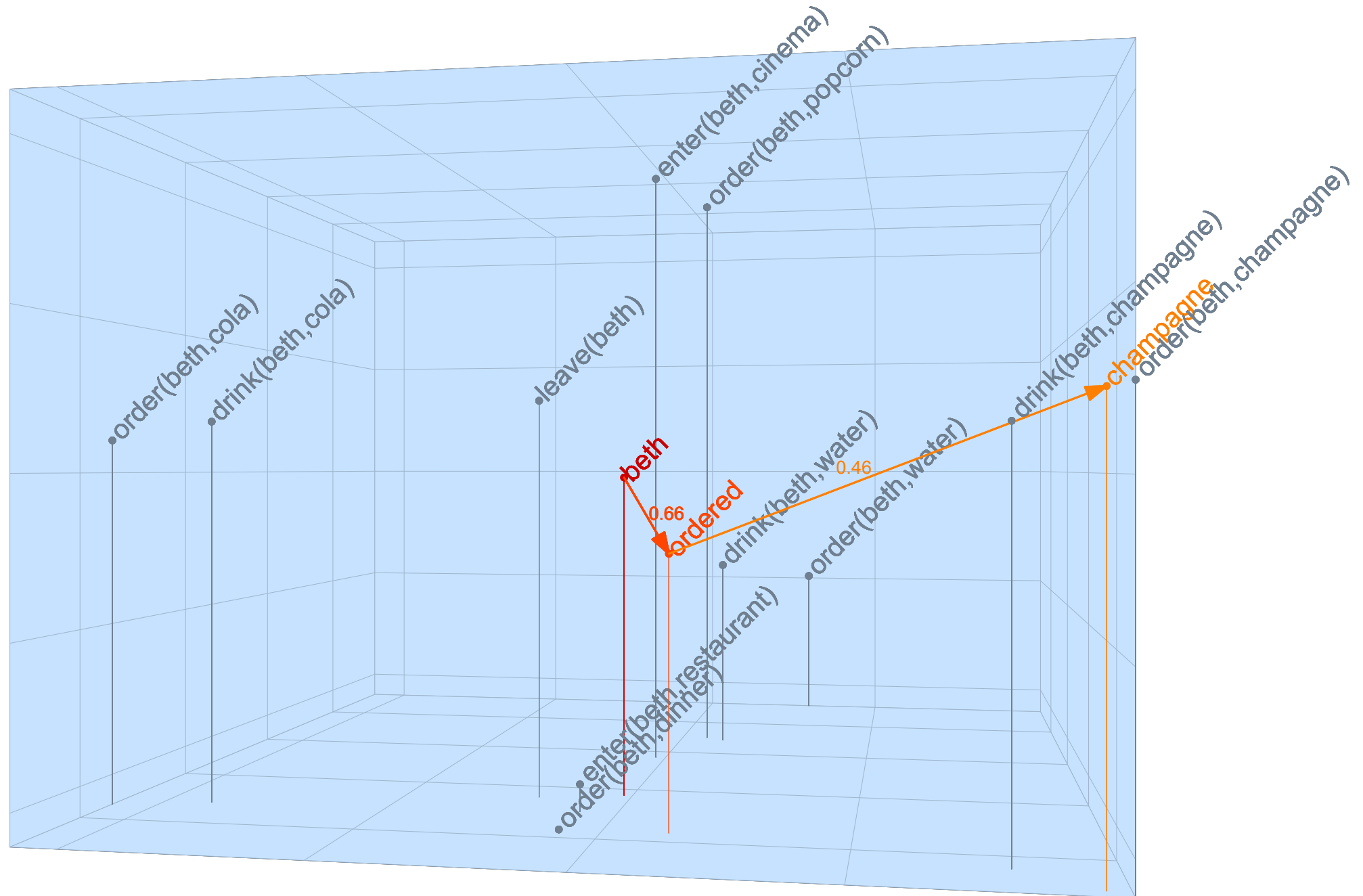
Comprehension is meaning-space navigation



["beth", "ordered"]

(scalars $\propto$ distance)

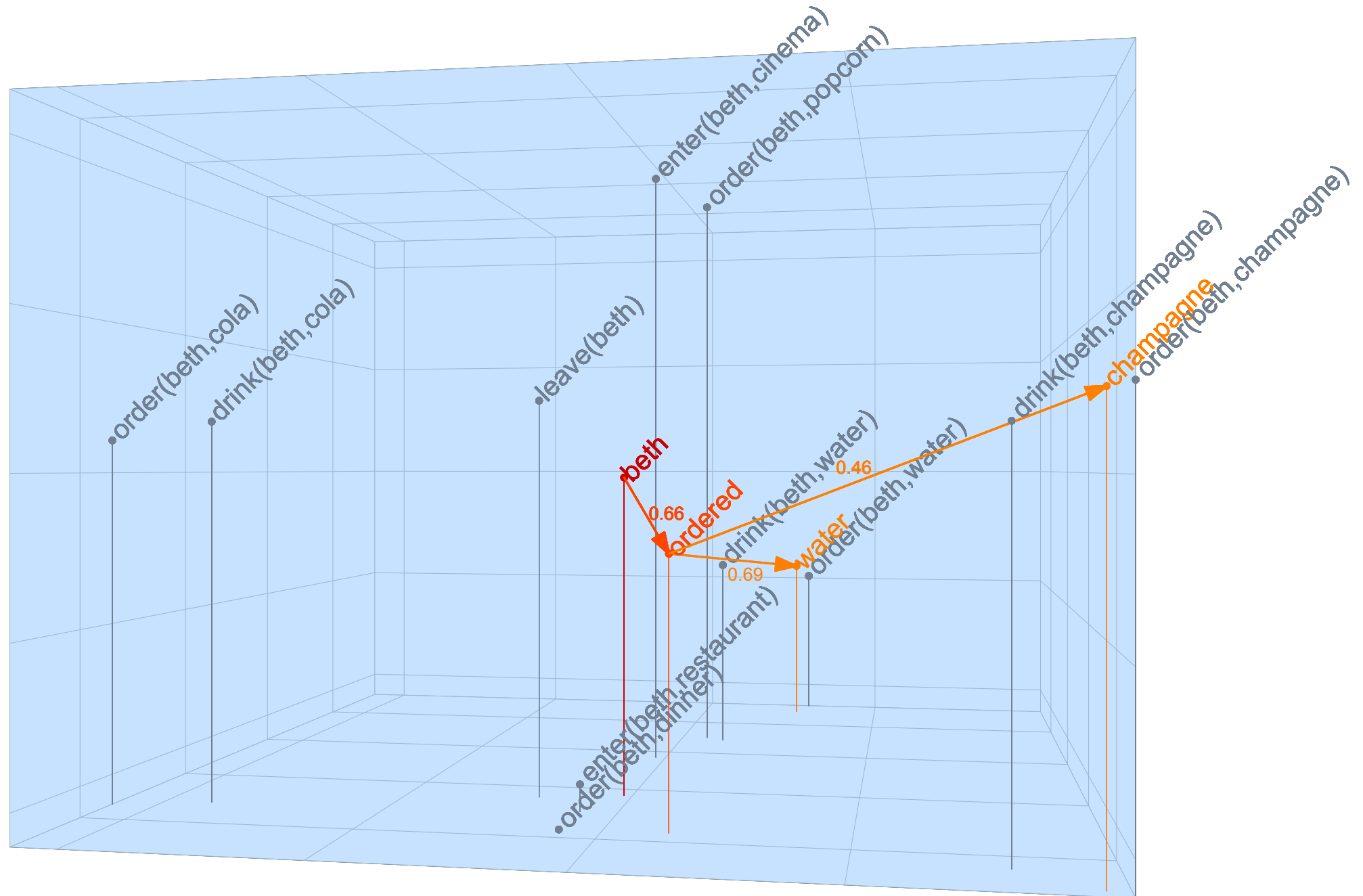
Comprehension is meaning-space navigation



["beth", "ordered", "champagne"]

(scalars $\propto$ distance)

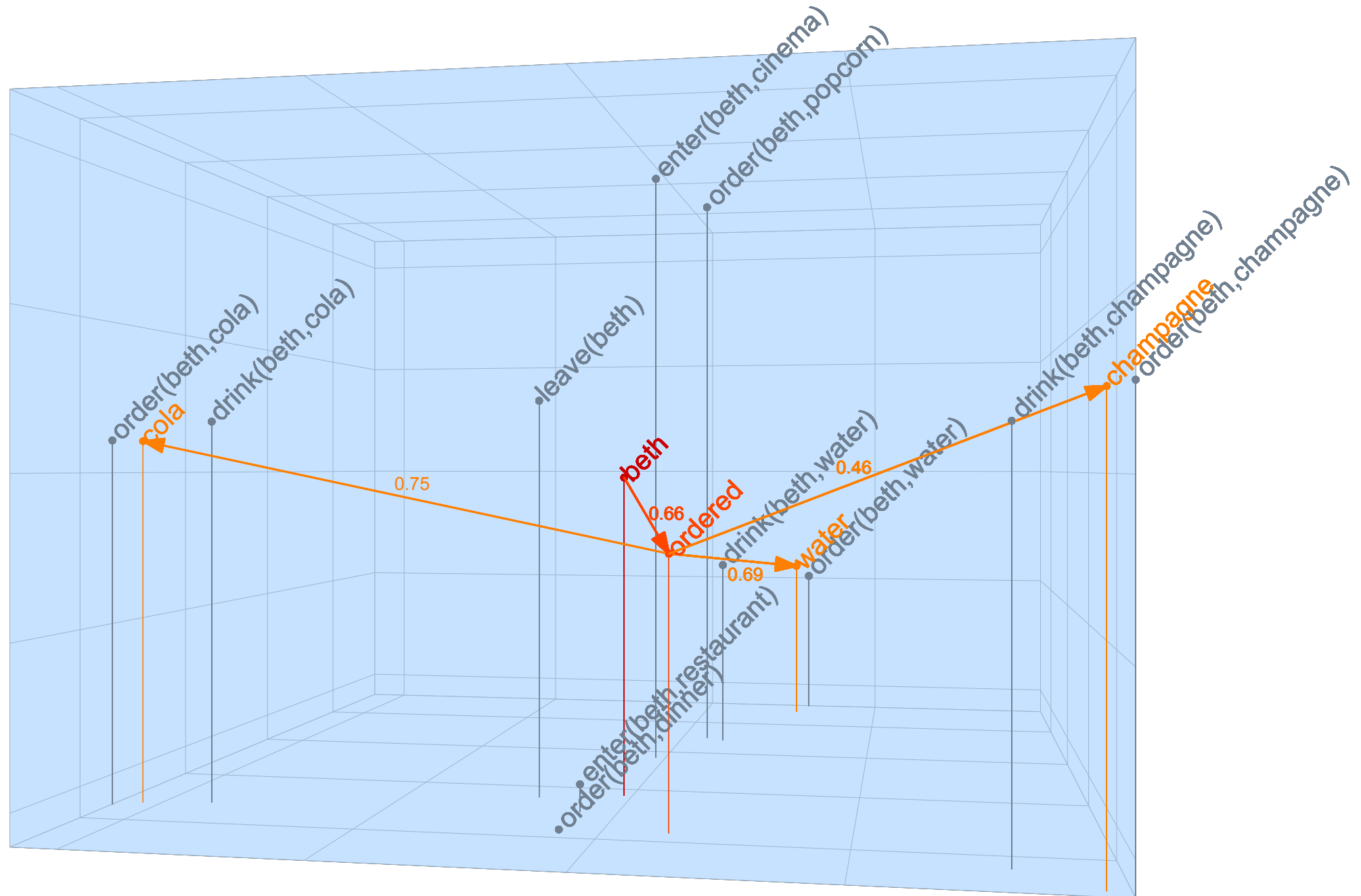
Comprehension is meaning-space navigation



["beth", "ordered", "water"]

(scalars $\propto$ distance)

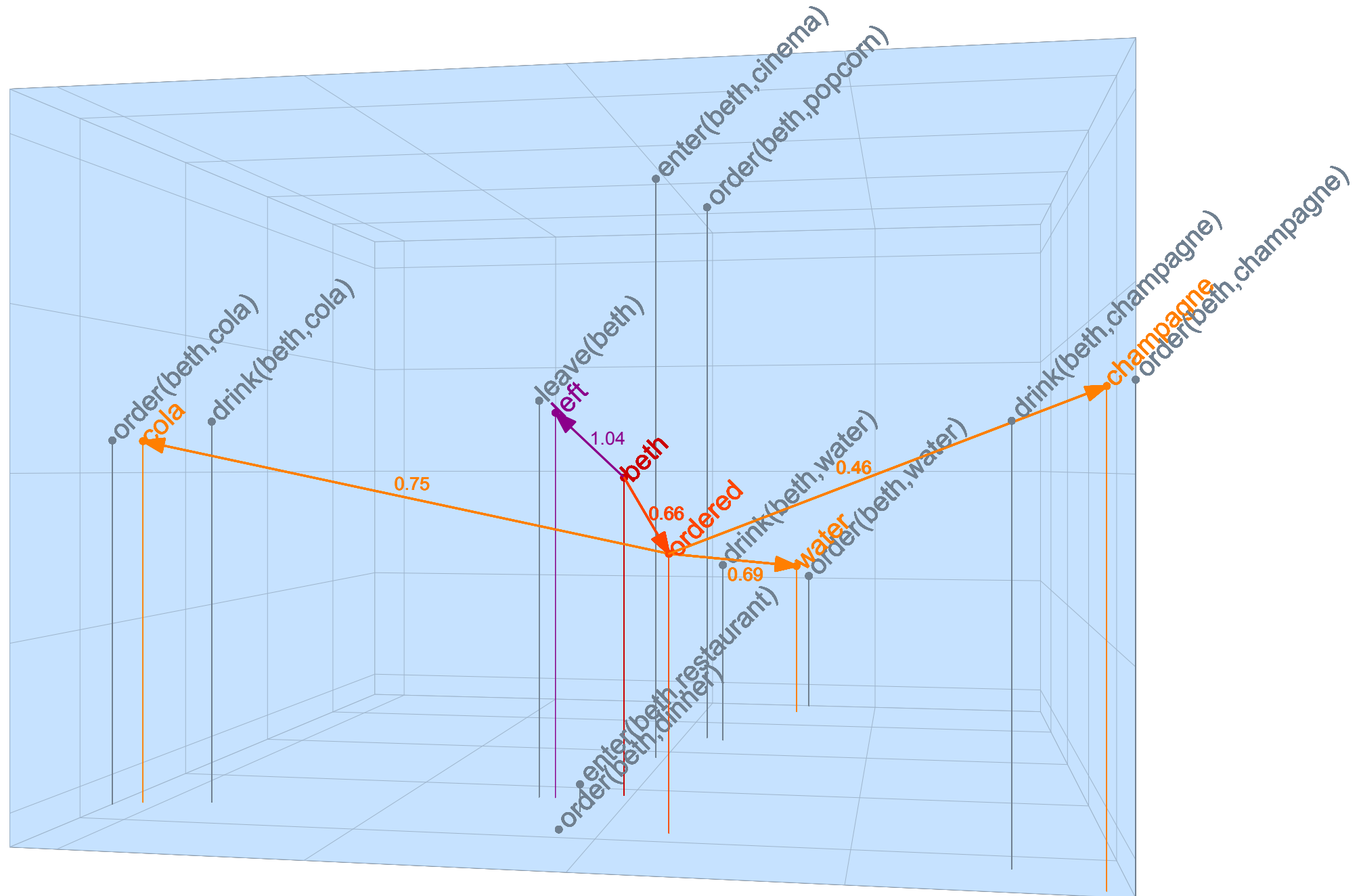
Comprehension is meaning-space navigation



["beth", "ordered", "cola"]

(scalars $\propto$ distance)

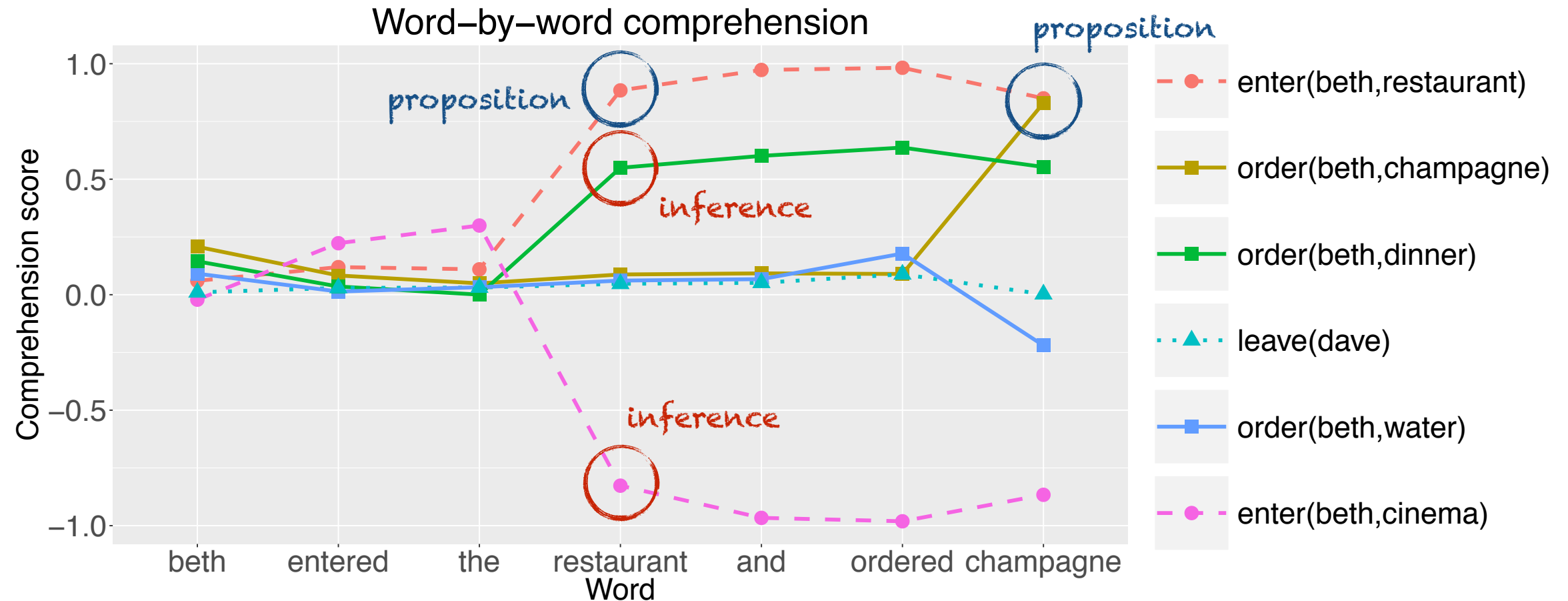
Comprehension is meaning-space navigation



["beth", "left"]

(scalars $\propto$ distance)

What does the model 'understand'?



> Points in meaning-space capture meaning beyond literal propositional content; i.e., model engages in **direct knowledge-driven inferencing**

What does the model 'understand'?

