

1 meanings write themselves [Keith]

When creating the meaning of a word, after choosing a type, you say that the meaning “writes itself.” Is there a more formal way of determining the meaning of a word? Saying that the meaning “writes itself” feels like hand-waving.

What “writes itself” is that part of the meaning of an item that follows from the general schema for meaning descriptions and from deciding on a semantic type for the item. So, once we have decided that *foggy* as in the sentence *London is a foggy town* is of type $\langle\langle e, t \rangle, \langle e, t \rangle\rangle$, the following writes itself:

- (1) For any context c and any world w :

$$\llbracket \text{foggy} \rrbracket^{c,w} = \lambda f_{e,t}. \lambda x_e. \dots$$

Part of what comes in the dots follows from the analysis that *foggy* is an intersective modifier (to be a foggy town, one has to be a town):

- (2) For any context c and any world w :

$$\llbracket \text{foggy} \rrbracket^{c,w} = \lambda f_{e,t}. \lambda x_e. f(x) = 1 \ \& \ \dots$$

There’s no handwaving so far.

But then we need to say what it takes for the function to return true for given arguments f and x in a given context c and world w . There we do often handwave to focus on whatever matters for our scientific inquiry at the moment. So, we often write just:

- (3) For any context c and any world w :

$$\llbracket \text{foggy} \rrbracket^{c,w} = \lambda f_{e,t}. \lambda x_e. f(x) = 1 \ \& \ x \text{ is foggy in } w.$$

If we were doing “lexical semantics” or “conceptual analysis”, we would investigate more what it takes for something to be foggy, but often we abstract

away from that and do engage in strategic handwaving, a time-honored practice in all sciences.

2 salience, relevance [Shuli]

In this class, we've often referred to things like "an individual relevant in the context" or "a salient individual". Is there work in semantics looking at how these notions of relevancy or saliency are defined in detail for different situations? What does that work look like?

There is work on that. Here's one place to start (a paper by a friend of mine): [Beaver 2004](#).

3 differences [Yuru, Karissa]

What are some notable aspects of the semantics we've learned that are markedly different in other languages?

I'm generally curious about how evaluating sentences works in languages with grammars which are very different from that of English since our model seems to go hand in hand with the syntactic structure of sentences. For example, when we looked at quantifier raising, I wonder if there are other languages do not have/need the process of QR because quantifiers already exist in the raised position. Basically, does the syntax of a language have a large effect on the study of semantics around that language, and do the fundamental ideas still apply despite the variation?

I tend to assume as a methodological principle that there is no difference between languages and then facing each case of divergence with the perspective of trying to figure out how much difference there really is. As far as semantics is concerned, I've written a paper with Lisa Matthewson (a pioneer in semantic fieldwork) on "semantic universals": [von Fintel & Matthewson 2008](#).

4 ignoring elements [Yuru]

how can one tell when a word can be ignored in a tree/calculation, like the "is" in "London is a foggy town"?

Our decision to ignore those was purely a pedagogical/strategic decision, not a scientific claim. In fact, it's clear that *is* and *a* are semantically meaningful. Whether there are any elements that are truly vacuous semantically, so-called “expletives” or “pleonastics”, is unclear. Candidates include:

- the *there* in existential sentences (*there is a horse in the garden*)
- the *that* in complementation (*Steph believes that the team will win*)
- the weather *it* (*it is raining*)

But for each of those, one might think that they do in fact contribute meaning.

5 ambiguity [Vignesh]

How would we write a meaning for an ambiguous sentence such as

(1) Flying planes can be dangerous.

This sentence seems to have two meanings: one meaning is “planes which are flying can be dangerous”. The other meaning is “(to be) flying planes can be dangerous” (hope this makes sense). Both meanings can arise but there is no distinction in syntax to rely on here that would allow us to distinguish between either. This seems to be highly context dependent but I would regard it as a well-formed sentence and I couldn't think of anyway to encode both meanings without relying heavily on the context. Is there a general way that ambiguities like above are handled? Can we encode the meaning of this sentence on its own like we have with other sentences in this class?

We would attribute two meanings to the sentence because it has two structures, depending on whether the subject is read as a description of an event (type) or of a plurality of individuals.

- (4) a. [PRO flying planes] can be dangerous.
 b. Flying planes can be dangerous.

There are discussions in the literature of cases where one might want to say that one and the same sentence structure is associated with a set of meanings, but then one has to say what it means to assert such a sentence etc. One example is my work with Thony Gillies on the meaning of *might*: [von Fintel & Gillies 2011](#).

6 might [Isabel]

How do we interpret existential intensional operators in sentences without propositional attitudes? We interpreted the sentence “Deborah believes that Bob might be in his office” to mean that there is some world w' consistent with Deborah’s beliefs such that Bob is in his office in w' , but in a sentence like “Bob might be in his office”, what must w' be consistent with?

It seems to me we can’t say that “there is some world w' consistent with the state of the world w such that Bob is in his office in w' ”, since Bob is either in his office or not in his office in w . Hence, there should be some context dependency, e.g. “there is some world w' consistent with the utterer’s knowledge of the state of the world w such that Bob is in his office in w' ”. Is this accurate? I’m skeptical, because if it is, then it seems to me that a sentence like “Bob is in his office” could just as easily be interpreted as “for all worlds w' consistent with the utterer’s knowledge of the state of the world w , Bob is in his office in w' ”. But earlier in the semester a similar idea was proposed, that every declarative sentence has an implicit “according to me” at the beginning, and this was invalid (although I don’t quite remember the argument for why).

I think that the best analysis for *might* for the cases you mention is that it is an existential quantifier over worlds compatible with the evidence available in the context c . More on this in [von Fintel & Gillies 2007](#).

7 50% chance [Isabel]

Also, at the very beginning of the semester, I asked about a sentence “There is a 50% chance of rain tomorrow”. Is it possible to come up with a meaning for this sentence using intensional operators? Something like “in half of all worlds w' consistent with the state of the world w , it will rain tomorrow” seems plausible, although that particular meaning is sketchy to me for the same reason as above (in w , it either rains tomorrow or it doesn’t).

Since there are very likely an uncountable infinity of possible worlds, notions like “half of” will not be well-defined. So, the usual approaches involve some kind of probability theory. More on this in [Yalcin 2010](#).

8 focus [Sofia]

How could we encode stress/emphasis?

I didn't say **she** stole my pizza. -> Someone else stole the pizza vs

I didn't say she stole **my pizza**. -> She stole something other than pizza

The standard approach is to say that focus makes salient a set of alternatives to the proposition uttered. The alternatives to *she stole my pizza* are propositions of the form *x stole my pizza* and in your example the speaker then says that they didn't say *she stole my pizza* but one of the alternatives. And something parallel would be said about the other example.

A classic paper is [Rooth 1992](#).

9 respectively [Brian]

how does "respectively" work in a sentence like e.g. "Kai and Sabine teach semantics and syntax, respectively."?

At first glance, it would look like "respectively" needs to match up elements of conjunctions.

But there are examples that show that whatever the matching process is can occur across sentences:

- (5) a. Where do John and Bill live?
b. They live in New York and Chicago respectively.
- (6) The people you have indicated are all New York residents except for the last two. They live in Chicago and Columbus respectively.

(See [Dalrymple & Kehler 1995](#) and for more recent work: use Google Scholar by finding this article and seeing what works it is cited in.)

10 semantics of math vs semantics of natural language [Daniel]

How does the semantics of the language of mathematical proofs differ from the semantics of natural language?

The first answer is that there isn't a difference. Mathematical proofs are written in a specialized subset of natural language, and so our semantics applies to it.

However, in metamathematics a proof-theoretic approach is more usual than a denotational/compositional typed lambda calculus approach (which is what we've developed).

However, the two kinds of systems are equivalent: the Curry-Howard Isomorphism. And in fact, some linguists do proof-theoretic syntax and semantics.

Of course, mathematics raised big questions about ontology and epistemology, but that is the job of philosophers of mathematics.

11 can computers understand the zombie sentence? [Keith]

what is the current state of the art in computers doing semantic calculations? do computers have difficulty understanding sentences like "if i were a zombie, i hope you would shoot me?"

Well ... the following is a conversation with an AI assistant (<https://beta.openai.com/playground/p/default-chat?model=text-davinci-002>).

(7) The assistant is helpful, creative, clever, and very friendly.

Human: Hello, who are you?

AI: I am an AI created by OpenAI. How can I help you today?

Human: If I become a zombie, I hope that you will shoot me.

AI: I would be very sad to shoot you if you became a zombie, but I would do it to protect others.

Follow up, what are your thoughts on GPT-3. I've heard some linguists express concern over GPT-3 because of its "cheap production of good, semantic artefacts." Do you have any insights on this?

Clearly, the performance of these systems after ingesting and processing terabytes of human language is very impressive. As a scientist of human language, we haven't learned anything from this so far, though, since we can't find out what the systems have learned: what structures they assign,

what meanings they assign to individual elements and to combinations. Even if we could learn what their analysis is, it's not obvious that it is the same analysis that allows human infants to develop language from much more limited data and with much more limited resources.

Also, the productions of GPT-3 are anything but cheap: [García-Martín et al. 2019](#).

12 what are semanticists working on [Xiaomeng/Miranda]

What are some interesting problems that semanticists are working on?

All of the problems we're working on are interesting. :) One way to find out what's happening at the moment is to browse the [Semantics Archive](#) and the semantics listings at [LingBuzz](#).

13 what's next [Karissa]

Are there any classes to take or other resources at MIT to pursue that you would recommend to further learn/build upon the material in this class?

At the undergraduate level, we do not teach semantics beyond 24.903, but some undergraduates with a strong interest in linguistics do take our first year graduate classes, including the semantics classes 24.970 and 24.973. The philosophy side of our department has relevant classes as well, including classes in logic and philosophy of language.

References

- Beaver, David I. 2004. The optimization of discourse anaphora. *Linguistics and Philosophy* 27(1). 3–56. DOI: [10.1023/B:LING.0000010796.76522.7a](#).
- Dalrymple, Mary & Andrew Kehler. 1995. On the constraints imposed by *respectively*. *Linguistic Inquiry* 26(3). 531–536. JSTOR: [4178909](#).
- von Fintel, Kai & Anthony S. Gillies. 2007. An opinionated guide to epistemic modality. In Tamar Szabó Gendler & John Hawthorne (eds.), *Oxford studies in epistemology: Volume 2*, 32–62. Oxford University Press. <http://mit.edu/fintel/fintel-gillies-2007-ose2.pdf>.

- von Fintel, Kai & Anthony S. Gillies. 2011. *Might* made right. In Andy Egan & Brian Weatherson (eds.), *Epistemic modality*, 108–130. Oxford: Oxford University Press. <http://mit.edu/fintel/fintel-gillies-2011-mmr.pdf>.
- von Fintel, Kai & Lisa Matthewson. 2008. Universals in semantics. *The Linguistic Review* 25(1-2). DOI: [10.1515/TLIR.2008.004](https://doi.org/10.1515/TLIR.2008.004).
- García-Martín, Eva, Crefeda Faviola Rodrigues, Graham Riley & Håkan Grahns. 2019. Estimation of energy consumption in machine learning. *Journal of Parallel and Distributed Computing* 134. 75–88. DOI: [10.1016/j.jpdc.2019.07.007](https://doi.org/10.1016/j.jpdc.2019.07.007).
- Rooth, Mats. 1992. A theory of focus interpretation. *Natural Language Semantics* 1(1). 75–116. DOI: [10.1007/BF02342617](https://doi.org/10.1007/BF02342617).
- Yalcin, Seth. 2010. Probability operators. *Philosophy Compass* 5(11). 916–937. DOI: [10.1111/j.1747-9991.2010.00360.x](https://doi.org/10.1111/j.1747-9991.2010.00360.x).