

1 Defining Pragmatics

What did they mean by that? It's a relatively common question, and it's precisely the subject of the field of pragmatics. In order to know what someone meant by what they said, it's not enough to know the meanings of the words (semantics) and how they have been strung together into a sentence (syntax); we also need to know who uttered the sentence and in what context, and to be able to make inferences regarding why they said it and what they intended us to understand. *There's one piece of pizza left* can be understood as an offer ("would you like it?") or a warning ("it's mine!") or a scolding ("you didn't finish your dinner"), depending on the situation, even if the follow-up comments in parentheses are never uttered. People commonly mean quite a lot more than they say explicitly, and it's up to their addressees to figure out what additional meaning they might have intended. A psychiatrist asking a patient *Can you express deep grief?* would not be taken to be asking the patient to engage in such a display immediately, but a movie director speaking to an actor might well mean exactly that. The literal meaning is a question about an ability ("are you able to do so?"); the additional meaning is a request ("please do so") that may be inferred in some contexts but not others. The literal meaning is the domain of semantics; the "additional meaning" is the domain of pragmatics.

This chapter will largely consider the difference between these two types of meaning – the literal meaning and the intended and/or inferred meaning of an utterance. We will begin with preliminary concepts and definitions, in order to develop a shared background and vocabulary for later discussions. A section on methodology will compare the corpus-based methodology favored by much current pragmatics research with the use of introspection, informants, and experimental methods. A brief survey of the philosophical background for pragmatics will be presented. Then, since no discussion of pragmatics can proceed without a basic understanding of semantics and the proposed theoretical bases for distinguishing between the two fields, the remainder of the chapter will be devoted to sketching the domains of semantics and pragmatics. A discussion of truth tables and truth-conditional semantics will both introduce the logical notation that will be used throughout the text and provide a jumping-off point for later discussions of the role of truth-conditions in drawing the semantics/pragmatics boundary. The discussion of the domain of semantics will be followed by a parallel discussion of the domain of pragmatics, including some of

the basic tenets of pragmatic theory, such as discourse model construction and mutual beliefs. The chapter will close with a comparison of competing models of the semantics/pragmatics boundary and an examination of some phenomena that challenge our understanding of this boundary.

1.1 Pragmatics and Natural Language

1.1.1 Introduction and preliminary definitions

Linguistics is the scientific study of language, and the study of linguistics typically includes, among other things, the study of our knowledge of sound systems (phonology), word structure (morphology), and sentence structure (syntax). It is also commonly pointed out that there is an important distinction to be made between our **competence** and our **performance**. Our competence is our (in principle flawless) knowledge of the rules of our **idiolect** – our own individual internalized system of language that has a great deal in common with the idiolects of other speakers in our community but almost certainly is not identical to any of them. (For example, it's unlikely that any two speakers share the same set of lexical items.) Our performance, on the other hand, is what we actually do linguistically – including all of our hems and haws, false starts, interrupted sentences, and speech errors, as well as our frequently imperfect comprehension: Linguists commonly point to sentences like *The horse raced past the barn fell* as cases in which our competence allows us – eventually – to recognize the sentence as grammatical (having the same structure as *The soldiers injured on the battlefield died*), even though our imperfect performance in this instance initially causes us to mis-parse the sentence. (Such sentences are known as **garden-path** sentences, since we are led “down the garden path” toward an incorrect interpretation and have to retrace our steps in order to get to the right one.)

Pragmatics may be roughly defined as the study of language use in context – as compared with semantics, which is the study of literal meaning independent of context (although these definitions will be revised below). If I'm having a hard day, I may tell you that my day has been a nightmare – but of course I don't intend you to take that literally; that is, the day hasn't in fact been something I've had a bad dream about. In this case the semantic meaning of *nightmare* (a bad dream) differs from its pragmatic meaning – that is, the meaning I intended in the context of my utterance. Given this difference, it might appear at first glance as though semantic meaning is a matter of competence, while pragmatic meaning is a matter of performance. However, our knowledge of pragmatics, like all of our linguistic knowledge, is **rule-governed**. The bulk of this book is devoted to describing some of the principles we follow in producing and interpreting language in light of the context, our intentions, and our beliefs about our interlocutors and their intentions. Because speakers within a language community share these pragmatic principles concerning language production and interpretation in context, they constitute part of our linguistic competence, not merely matters of performance. That is to say, pragmatic knowledge is part of our knowledge of how to use language appropriately. And as with other areas of linguistic competence, our pragmatic competence is generally **implicit** – known at some

level, but not usually available for explicit examination. For example, it would be difficult for most people to explain how they know that *My day was a nightmare* means that my day (like a nightmare) was very unpleasant, and not, for example, that I slept through it. Nightmares have both properties – the property of being very unpleasant and the property of being experienced by someone who is asleep – and yet only one of these properties is understood to have been intended by the speaker of the utterance *My day was a nightmare*. The study of pragmatics looks at such interpretive regularities and tries to make explicit the implicit knowledge that guides us in selecting interpretations.

Because this meaning is implicit, it can be tricky to study – and people don't even agree on what is and isn't implicit. One could make a strong argument that *a nightmare* in *My day was a nightmare* is actually explicit, that this metaphorical meaning has been fully incorporated into the language, and that it should be considered literal, not inferential (i.e., semantic rather than pragmatic). This in itself is a very interesting question: Every figure of speech began as a brand-new but perfectly interpretable utterance – one could say *My day was one long, painful slide down an endless sheet of coarse-grain sandpaper* – that eventually became commonplace. Upon their first utterance, such figures of speech require pragmatic inference for their interpretation; the hearer must (whether consciously or subconsciously) work out what was intended. It's possible that this is still what's done when the figure of speech becomes commonplace; it's also possible that it becomes more like a regular word, whose meaning is simply conventionally attached to that string of sounds. If the latter is the case, it's obviously impossible to say precisely when its status changed, since there was no single point at which that happened – which is to say, the shift from pragmatic meaning to semantic meaning, if and when it occurs, is a continuum rather than a point.

One might ask why it matters – but in fact there are a great many reasons why it matters. Just consider a court of law: It matters enormously what counts as “the truth, the whole truth, and nothing but the truth.” Does inferential meaning count as part of that truth? Courts have frequently found that for legal purposes, only literal truth matters; that is, in saying *There's one piece of pizza left*, you can be held responsible for the number of pieces of pizza left, but not for any additional meaning (such as “offer” vs. “scolding”). On the other hand, we'll see in Chapter 9 that the courts haven't been entirely consistent on this issue. More generally, most people can think of cases within their own relationships in which what the speaker intended by an utterance and what the hearer took it to mean have been two entirely different things; rather sizeable arguments are sometimes due to a difference in pragmatic interpretation, with each party insisting that their interpretation constitutes what was “said.”

Pragmatics, then, has to do with a rather slippery type of meaning, one that isn't found in dictionaries and which may vary from context to context. The same utterance will mean different things in different contexts, and will even mean different things to different people. The same noun phrase can pick out different things in the world at different times, as evidenced by the phrase *this clause* in *This clause contains five words; this clause contains four*. All of this falls under the rubric of pragmatics. In general terms, pragmatics typically has to do with meaning that is:

- nonliteral,
- nonconventional,
- context-dependent,

- inferential, and/or
- not truth-conditional.

We'll talk a lot more about that last one ("not truth-conditional") later on; for now, it's enough to notice that when I say *There's one piece of pizza left*, the truth of that statement has everything to do with how many pieces of pizza are left, and nothing to do with whether I intend the statement as an offer or a scolding. Thus, the conditions under which the statement is true don't depend on its pragmatic meaning; that's what we mean when we say that the pragmatic meaning is generally not truth-conditional.

The "and/or" in that bulleted list is the real problem. Linguists disagree on which of these are actually defining properties of pragmatics. A prototypical case of pragmatic meaning is indeed nonliteral, nonconventional, context-dependent, inferential, and not truth-conditional. However, there are other cases in which it's not so clear. The case of *this clause* is a good example: Many linguists would say that determining which clause is being referred to requires a pragmatic inference, even though it affects the truth conditions of the utterance. (That is, which clause is being referred to crucially affects the question of whether *This clause contains four words* is true.) Others would say that any piece of meaning that affects truth is by definition semantic. Thus, the boundary between what counts as semantics and what counts as pragmatics is still a matter of open debate among linguists, and it will recur throughout this book as an important theme.

1.1.2 *Situating pragmatics within the discipline of linguistics*

Language use involves a relationship between **form** and **meaning**. As noted above, the study of linguistic form involves the study of a number of different levels of linguistic units: **Phonetics** deals with individual speech sounds, **phonology** deals with how these sounds pattern systematically within a language, **morphology** deals with the structure of words, and **syntax** deals with the structure of sentences. At each level, these forms may be correlated with meaning. At the phonetic/phonological level, individual sounds are not typically meaningful in themselves. However, intonational contours are associated with certain meanings; these associations are the subject of the study of **prosody**. At the morphological level, individual words and morphemes are associated with meanings; this is the purview of **lexical semantics** and **lexical pragmatics**. And at the sentence level, certain structures are conventionally associated with certain meanings (e.g., when two true sentences are joined by *and*, as in *I like pizza and I eat it frequently*, we take the resulting conjunction to be true as well); this is the purview of **sentential semantics**. Above the level of the sentence, we are dealing with pragmatics, including meaning that is inferred based on contextual factors rather than being conventionally associated with a particular expression.

Pragmatics is closely related to the field of **discourse analysis**. Whereas morphology restricts its purview to the individual word, and syntax focuses on individual sentences, discourse analysis studies strings of sentences produced in a connected discourse. Because pragmatics concentrates on the use of language in context, and the surrounding discourse is part of the context, the concerns of the two fields overlap significantly. Broadly speaking,

however, the two differ in focus: Pragmatics uses discourse as data and seeks to draw generalizations that have predictive power concerning our linguistic competence, whereas discourse analysis focuses on the individual discourse, using the findings of pragmatic theory to shed light on how a particular set of interlocutors use and interpret language in a specific context. In short (and far too simplistically), discourse analysis may be thought of as asking the question “What’s happening in this discourse?,” whereas pragmatics asks the question “What happens in discourse?” Pragmatics draws on natural language data to develop generalizations concerning linguistic behavior, whereas discourse analysis draws on these generalizations in order to more closely investigate natural language data.

1.1.3 *Methodological considerations*

It should be noted that (like all of linguistics) the study of pragmatics is inherently **descriptive**, describing language as it is actually used, rather than **prescriptive**, prescribing how people “ought” to use it according to some standard. A linguist will never tell you not to split your infinitives; they will simply observe that people do indeed split their infinitives, and will include this in their descriptive observations of language use.

Although it may seem obvious that we as scientists are interested in describing language use rather than in telling language users how they should speak, the terminology of the field can sometimes confuse the issue. For example, the Cooperative Principle to be discussed in Chapter 2 presents a series of maxims phrased as imperatives – “say enough,” “don’t say too much,” and so on. In truth, however, these are not rules that language users are being told to follow, but rather descriptions of principles that they generally **do** follow, and which they expect each other to follow. Nobody has to be explicitly taught to follow these guidelines; instead, they are part of what we implicitly know as rational users of language. Therefore, it is important to keep in mind that although some of the principles described in this book are phrased in imperative form, they actually describe what speakers do automatically in using language. Rather than “speakers should do X,” what is really meant is “speakers (consistently and reliably are observed to) do X.”

In order to determine what it is that speakers do, linguists have traditionally used one of three basic methods to study language use and variation:

1. Native-speaker intuitions
 - a. Your own (introspection)
 - b. Someone else’s (informants)
 - questionnaires
 - interviews
2. Psycholinguistic experimentation
 - a. lexical decision, eye tracking, etc.
3. Naturally occurring data
 - a. Elicitation
 - b. Natural observation
 - c. Corpus data

The first of these, the researcher's own **intuition**, is valuable during the initial stage of research, during the process of forming a hypothesis. It helps to guide the researcher toward a reasonable hypothesis and away from hypotheses that are clearly untenable. But once you have a hypothesis, your intuition becomes unreliable, since it may be biased toward confirming your own hypothesis. A better option is to use the intuitions of a group of informants via questionnaires or interviews, but here too you must be careful: Subjects may (consciously or not) try to please or impress you by reporting their speech as more prescriptively "correct" than it actually is. This is the **observer's paradox** (Labov 1972): The presence of the observer affects the behavior of those being observed. Moreover, people often don't have accurate knowledge of how they speak when they're not paying attention.

Psycholinguistic experimentation is able to eliminate some of these difficulties by testing people's actual linguistic knowledge and behavior outside of their ability to manipulate this behavior. For example, a **lexical decision** task might ask subjects to read a text and then present them with either a common word of the language or a nonsense word; their task is to determine whether the word shown is real or not. Words made salient or cognitively "accessible" by the prior text are more quickly identified as real words than are unrelated words.

Similarly, an **eye-tracker** can determine precisely where someone is looking at a given instant (to determine, for example, what the individual takes to be the referent of a particular pronoun in a presented text, or what part of a sentence takes the longest to understand). But again, very careful set-up and control of the experiment are required in order to eliminate the observer's paradox. Typically, care is taken to ensure that the subject is unaware of what is actually being tested.

The use of **naturally occurring data** gets around these difficulties by observing language in actual use under natural conditions. **Elicitation** (in which the researcher creates a context that's conducive to getting the subject to utter the desired form) is only an improvement over intuitions if the subject is unaware that they're being observed. William Labov is famous for (among other things) a dialect study in which he asked department-store workers about the location of various items; in truth, he was merely eliciting the words *fourth floor* in order to determine which individuals dropped the [r] sound from each of the words (Labov 1966). **Natural observation** is like elicitation, except that rather than setting up a context to compel your subject to utter the desired form, you simply wait in some natural setting and watch, hoping that they will do so – and that they will do so with sufficient frequency to give you enough data to be useful. However, depending on the frequency of the desired form, one could wait quite a long time before collecting enough data to do a proper study.

The use of **corpus data** circumvents many of the above problems, in that it involves a pre-existing collection of raw language data, typically consisting of millions of words, which have been naturally produced and which can be scoured for instances of the forms under investigation. In the past, such corpora have been extremely difficult to compile, but with the computer age has come the ability to store a virtually unlimited amount of text in an easy-to-search format. The use of corpora avoids the observer's paradox, as well as sparing the researcher the trouble of waiting for a form to be produced or trying to elicit it. The use of corpus data does, however, have its own drawbacks. For example, you must

take care in selecting your data sample. If your data are skewed, so will your results be. If you only look at men's speech, your results are only valid for men's speech. If you do a corpus study but use as your corpus only romance novels from the 1990s, your results will only be valid for that group of works, and you cannot generalize them to English as a whole. Less obviously, if your corpus is entirely written, it may not accurately tell you what spoken English is like. If Labov had only conducted his experiment in a single department store, he would have gotten a skewed impression of what English is like in New York City as a whole. Thus, it is important to be certain that your data are appropriate to the hypothesis that you plan to test. Second, be aware that some of the utterances encountered in corpora will contain performance errors – all those hems, haws, false starts, and so on that do not accurately reflect the language user's linguistic competence. Thus, in interpreting the results of a corpus study, researchers inevitably make reference once again to their own imperfect intuitions in order to interpret the data they are confronted with. The best insurance is to collect as many tokens as possible (since the more data one has, the less likely it is that a performance error here or there will pose a serious threat of corrupting one's findings) and to have any necessary coding of the data performed by more than one researcher and then checked for inter-rater reliability.

Because of the nature of the field of pragmatics, it is especially important for researchers in this field to look at spontaneous language use in a naturally occurring context. Intuitions are notoriously unreliable for pragmatic research. Some ingenious psycholinguistic studies have been devised to test pragmatic theories, but much of the current research in pragmatics is based on the study of naturally occurring data.

Finally, the type of hypothesis you are testing should be both **falsifiable** and **predictive**. To say it should be falsifiable is not the same as saying it should be false; rather, there should be some way of testing whether it is true or false, which entails that the test allow for the possibility of its being false and present a clear answer to the question, "If my claim is false, how will this test demonstrate that it's false?" For example, consider the following claims:

- A discourse sometimes begins with a greeting.
- A discourse typically begins with a greeting.
- A discourse always begins with a greeting.

The first claim is not falsifiable, because there is no way to show that it is false (even though it's trivially easy to show that it's true). Suppose we check 100,000 discourses and find that none begins with a greeting; we will not know for sure that our claim is false, because it's always possible that the next discourse we look at will begin with a greeting and our claim will be vindicated. The second claim appears stronger, yet it too is unfalsifiable: First, the term *typically* is vague; second (and less obviously), here again we find the possibility (however unlikely) that we've just been unlucky in our selection of data and that the next 300,000 discourses will in fact begin with a greeting and will open up the possibility that our claim was correct after all. Only the third claim is falsifiable: Discovery of a single discourse that does not begin with a greeting (under some specific definition of the word *greeting*) irrevocably and irrefutably falsifies our claim. Because only the third claim is falsifiable, it is also the only one of the three that constitutes an **empirical**

(i.e., testable) claim. A claim is only empirical if you can imagine a circumstance that would show that it is false. And only empirical claims are scientifically interesting.

In order to be interesting, the claim must also be predictive, in the sense of being general or generalizable. That is, the claim must not simply be about a single instance of language use; instead, it must make a general claim about an entire class of uses, and therefore also predict how speakers will behave in the future. It's not interesting to present an example of a business letter and observe that it presents a problem and offers a solution, unless you can generalize this into a claim that business letters in general are constructed in such a way as to present a problem and offer a solution. Only by showing that your pragmatic theory applies to an entire definable class of data can you argue that the knowledge that it represents constitutes part of a native speaker's linguistic competence.

1.2 The Philosophical Background

The early insights of pragmatics originated within the field of philosophy, and came from philosophers who struggled with a range of interconnected questions concerning the nature of meaning, reference, and truth. An early identification of semantic meaning with philosophical logic led philosophers like H.P. Grice to address the relationship of literal semantic meaning and the chasm between this sort of "meaning" and the actual intention of a given speaker in a given instance. But even these questions arose earlier, in the wrestling over the place of use and intention in linguistic communication. To try to untangle some of these issues, let's consider the views of some of the major philosophers who have influenced today's pragmatic theories.

1.2.1 *Meaning and truth*

In approaching the notion of what a sentence "means," one might start from a notion of **compositionality** – the idea that the meaning of a larger unit (like a sentence) is made up of the meanings of the smaller units that it's made up of (like words and phrases). So we might look to the meanings of words such as nouns and verbs and adjectives and decide that it seems reasonable to say that they refer to objects in the real world. This is a **referentialist** view of meaning, as opposed to the **mentalist** view of meaning, which holds that meanings are objects in the human mind. Gottlob Frege (1848–1925) held a referentialist view of meaning, and moreover held a compositionalist view. So for him, the meaning of a sentence is constructed from the meanings of its parts – and since the meanings of its parts are their referents, a sentence has a referent built up from those parts ("Frege's Principle"). According to Frege, if you replace one phrase in the sentence with some other coreferential phrase, the entire sentence should continue to mean the same thing in the sense that it continues to be either true or false. Any sentence that is true should remain true if you swap out one of its phrases and replace it with a coreferential phrase.

The phrase has to be coreferential, though; you can't simply swap out a phrase for another one with the same semantic meaning. If I say *My Tuesday afternoon class has 35 students*, that may be true. But if I utter that same sentence two years from now, when I have a different group of students in a different Tuesday afternoon class, the sentence may be false at that point. So Frege distinguished between the **sense** of a phrase (which is essentially its semantic meaning) and the **reference** of the phrase (what it picks out in the world). For Frege, the referent of the sentence is either "the True" or "the False," and that will remain constant as long as the referents of its components are held constant. The sense of a phrase remains more or less constant, but its reference can change; only if the reference remains constant can you expect the reference (i.e., the truth or falsity) of the whole sentence to remain constant. (We'll return to this distinction later in this chapter.)

Because Frege is a referentialist, under his view a definite description (such as *the dog* or *the ceiling* or *my cousin*) has something in the world as its referent, and if no such referent exists, then there can be no referent for the sentence as a whole (that is, no True or False). So if I say *My sister is visiting me tomorrow* but I have no sister, then what I have said is neither true nor false.

Bertrand Russell (1872–1970) disagreed. He said that a sentence like *The present King of France is wise* is just plain false, because it entails that there is a present King of France; if there is no such entity, the sentence cannot be true. And since Russell held to the **Law of the Excluded Middle**, under which a given proposition is necessarily either true or false, there's no middle ground. On Russell's view, the meaning of a phrase is no more or less than what it contributes to the meaning of a sentence, and in this case what *The King of France* contributes includes existence and uniqueness. In short, for Russell (1.1) entails each of the propositions in (1.2):

- (1.1) The King of France is wise.
- (1.2) a. There is a King of France.
- b. There is exactly one such King.
- c. He is wise.

If any one of these three propositions is false, the entire sentence is false. This tidily preserves the Law of the Excluded Middle and spares us the trouble of tossing out our bivalent logic and replacing it with something else. The problem, however, is that these three propositions intuitively don't carry the same weight in the sentence. If I tell you *The King of France is wise*, and you deny its truth (saying, for example, *I disagree* or *That's false*), nobody would take you to be disagreeing with propositions (1.2a) or (1.2b) above; they would only take you to be quarrelling with the assertion that he is wise.

P.F. Strawson (1919–2006) criticized Russell's account on these grounds, arguing that if there's no King of France, then it's odd to say that (1.1) is false OR that it's true; according to Strawson, without a King in existence, the question of his wisdom simply doesn't arise. For Strawson, "Mentioning, or referring to, something is a characteristic of a use of an expression, just as 'being about' something, and truth-or-falsity, are characteristics of a use of a sentence" (1950: 326). So there have been points in history when there was in fact a King of France, and at such times (1.1) could have been used to posit something of that person, and it might have been either true or false. But at the current moment, there is no

such person, so in using that sentence, the speaker fails to refer, and so (1.1) has no truth-value. For Strawson, the use of (1.1) **presupposes** the existence of a King, and without him, its truth cannot be evaluated. In short, he argues that it's not a **sentence** that has a truth-value, but rather the **use** of a sentence.

But then came Keith Donnellan (1931–2015), who argued that it's more complicated than that, because even a single use of a single definite description can have two distinct meanings, which he called the **referential** and **attributive** meanings, and these meanings can have distinct truth-values. So to take his example, imagine you're a detective at a crime scene, where a guy named Smith has been murdered. You utter (1.3):

(1.3) Smith's murderer is insane.

This can mean one of two things: Either you're asserting your belief that whoever murdered Smith must be insane (the attributive reading), or else you believe you know who murdered Smith, and you're asserting of that person that they are insane (the referential reading). So if Sam Smith has been murdered by Jenny Jones, and you know Jones to be insane on independent grounds, you could felicitously utter (1.3) referentially to mean that Jones is insane.

Now, according to Donnellan, these two readings could have different truth-values under the same set of circumstances. Suppose Jones is in fact insane, but she didn't actually kill Smith. (Let's say Smith slipped on a banana peel in such a way that it looked like a murder.) For Donnellan, (1.3) is false under the attributive reading, but it's true under the referential reading, because your intended referent was Jones and it's true of Jones that she's insane.

But this brings us back to the question of referentialism vs. mentalism: If there's no real-world referent that can be semantically picked out by the phrase "Smith's murderer", does (1.3) really have a referent? If there's no real-world referent, it can't be true under a referentialist view – but if what matters isn't the real world but rather the **discourse model**, why should it matter whether Jones is insane in the real world, since she has the property "Smith's murderer" only in the speaker's discourse model – the same discourse model in which she's insane? For Donnellan's critics, to say that Smith's murderer is ambiguous between "the person who murdered Smith" and "the person I intend to refer to" is a confusion of semantic and pragmatic meaning, or between semantic reference and speaker reference (Kripke 1977).

Ludwig Wittgenstein (1889–1951) argued against any account of word meaning that gave a set of necessary and sufficient conditions, such as defining a table in terms of its construction or shape or use. He argued that the meaning of a word couldn't simply be described; instead, its meaning was an amalgam of all of the instances of its use. He introduced the notion of a **language game**, and the meaning of an expression was its use within such a game. So it wouldn't have bothered him to know that people argue about whether a hot dog is or is not a sandwich; the meaning of *sandwich* is to a large extent its appearance in the activity of using language, and it has a "family resemblance" to other related things with which it overlaps to a greater or lesser extent (think Philly cheesesteaks and tacos and pizzas and hamburgers); this view is similar to the notion of Prototype theory which we'll consider below (Rosch 1973, 1975), though it lacks a central prototype

for the concept in question, and it depends much more on the ever-evolving status of a word's meaning as it is used in the world. He makes the same argument for sentence meaning, contending that a sentence has no meaning whatsoever apart from its use. If I say *I'm cold* to my husband, he has to determine the meaning of that sentence in that instance of use; it could mean anything from a strict description of my state of comfort to a request that he start a fire in the fireplace. Its meaning will vary from one context to another. Pragmatics is everything.

1.2.2 Language use and context

At this point it would seem clear that language use in context is at least as important as the conventional meaning of the words. Philosopher John L. Austin (1911–1960) observed that in some cases, if the context is right, we can use a particular sentence to perform an activity in the world (think of *I apologize* or *I christen you Boaty McBoatface* (yes, that's a real boat in the British fleet!)). As Austin continued his investigations, he found that not only could language be used to perform the acts explicitly described, like apologies and christenings, but also such acts as asserting, asking, requesting, and much more. In fact, he realized, every act of speaking is, by definition, an act; and as he developed a taxonomy of the different ways in which language use performed acts in the world, the field of **speech acts** was created. Most interesting among these speech acts are the **indirect speech acts**, which perform an act that differs from the act they semantically purport to perform, as when I ask *Can you pass the salt?* – asking you about your ability to pass the salt as a way of indirectly requesting you to do so. John Searle (1932–) expanded on this theory of indirect speech acts, describing the **felicity conditions** that must be satisfied in order for a speech act to succeed, and developing the distinctions between **locutionary**, **illocutionary**, and **perlocutionary** acts – roughly, what was said, what was intended, and the effect it had. (We'll discuss these issues further in Chapter 6.)

Nowhere is the role of context more important than when we consider inference. And in the philosophical debate over the relative roles of literal vs. inferential meaning, nobody stands taller than H. Paul Grice (1913–1988). Whereas Wittgenstein argued that a sentence had no meaning outside of use, Grice, ironically enough, was trying to preserve the notion that a sentence did in fact have a literal, truth-conditional meaning – that is, a meaning outside of its use, in which its truth or falsity depended on the possible worlds in which it holds. This is perhaps ironic given that Grice's argument showing how we could retain such an account ended up being the seminal theory of language use in context.

We'll devote an entire chapter to Grice (and then another chapter to later refinements to his theory), but in essence he was dealing with the problem of how we can preserve a truth-conditional approach to the **logical connectives** (“connecting” words like *and* and *or*) when their actual use in natural language differs from their logical meanings. So take the word *and*. In standard logic, given two propositions *p* and *q*, the conjoined proposition *p and q* is true if and only if *p* is true and *q* is true, and in no other circumstance. If either of the two conjoined propositions is false (or both of them), then the whole thing (*p and q*) is false. Clear enough.

But logically speaking, that is the **entire** meaning of *p and q*. In natural language, however, we frequently take *p and q* to convey that *p* occurred before *q*, or that *p* caused *q*, as in (1.4) and (1.5):

- (1.4) a. I'm going to get a sweater and go out to the store.
- b. I'm going to go out to the store and get a sweater.
- (1.5) a. I complained to the manager, and they gave me reheated food.
- b. They gave me reheated food, and I complained to the manager.

In (1.4a), the most natural reading is that the speaker is going to grab a sweater to put on before going out to the store; in (1.4b) the opposite ordering is more natural, in which the speaker is planning to go to a store and buy a sweater there. In each case the connective *and* seems to have an additional meaning of “and then”, imposing an ordering on the events. In (1.5), ordering is again conveyed, but more importantly, there's a sense of causation: In (1.5a) the most natural interpretation is that the speaker complained to the manager (presumably about the food being cold), and because of that complaint, the food was reheated. In (1.5b), the opposite holds: The most natural reading is one in which the speaker was given reheated food and (because of that) complained to the manager.

Grice's problem was: Where did these additional meanings come from? Could philosophers retain the logical meaning of a word like *and* if the use of the word in context didn't match its logical meaning? He argued that they could, and should; and as we'll see in Chapter 2, he developed an account of how hearers arrive at these additional interpretations. Again, as with issues of reference above, the stage was set for the field of pragmatics to step into the limelight; and Grice was a major player – perhaps THE major player – in the early development and ultimate prominence of pragmatics as an approach to linguistic meaning.

1.2.3 Summary

In short, the field of pragmatics arose in response to a long train of conversation among philosophers on the nature of meaning and its relation to concepts like truth, reference, context, and the relationship between convention and context in human communication. Because Frege viewed the referent of a sentence as being “the True” or “the False”, he believed that if some element of that sentence had no referent, the entire sentence was without a referent – that is, it was neither true nor false. Russell countered that such a sentence was simply false, but this led to counterintuitive results, such as a sentence like *The King of France is wise* being false if there is no King of France. Strawson argued instead that such a sentence in such a situation isn't false but instead simply odd: Why are we asserting wisdom of a nonexistent entity? Moreover, such a sentence might have been true at a point in history when France was a monarchy; hence, truth and falsity must apply to the use of a sentence in some context rather than to the sentence itself. Donnellan argued that even within a single use, a sentence can be false on one reading (e.g., *Smith's murderer is insane* if the person who murdered Smith is perfectly sane)

while simultaneously being true on another (if the speaker incorrectly believes Jones killed Smith and wishes to assert, accurately, that Jones is insane). But at this point, we need to take into account the distinction between semantic meaning and pragmatic meaning, and address the question of what sort of “meaning” the question of truth and falsity applies to – which will be a central focus of this book. Indeed, we’ve seen that Wittgenstein argued that the meaning of an expression is nothing more or less than its use; and later philosophers such as Austin and Searle expanded the discussion into language use as an act, performed either directly or indirectly. And Grice, in many ways the father of modern pragmatics, developed the seminal theory showing how some basic principles of cooperation can explain a great deal about the relationship between conventional meaning and what we mean when we use language in context.

We will talk about most of these philosophers and their contributions to pragmatics later in the book, but this is a snapshot of how the conversation among philosophers about the nature of meaning has led to the linguistic subfield of pragmatics, and this is where we will pick up the conversation in Chapter 2.

1.3 The Boundary Between Semantics and Pragmatics

No discussion of pragmatics can proceed very far without a basic understanding of semantics and the proposed theoretical bases for distinguishing between the fields. Both deal with meaning, so there is an intuitive sense in which the two fields are closely related. There is also an intuitive sense in which the two are distinct: Most people feel they have an understanding of the “literal” meaning of a word or sentence as opposed to what it might be used to convey in a certain context. Upon trying to disentangle these two types of meaning from each other, however, things get considerably more difficult. We will spend the remainder of this chapter attempting to both describe and circumscribe the domains of semantics and pragmatics, ending with a discussion of some important phenomena that challenge traditional conceptions of the boundary between the two. We will begin with a brief survey of the field of semantics and the issues with which it concerns itself.

1.3.1 *The domain of semantics*

1.3.1.1 Word meaning

Semantic meaning is typically thought of as literal meaning of the sort one would find in the dictionary. Thus, perhaps the most straightforward place to begin a discussion of semantics is in the area of word meaning. The study of word meaning is called **lexical semantics**, as opposed to **sentential semantics**, which is the study of sentence meaning (discussed below). The meaning of a word has often been described in terms of the features necessary for a thing to count as an instance of the category described by the word; for example, the meaning of the word *dog* is that set of features by which something is

known to be a dog. Most word meanings are characterized by more than one such feature, so that we can talk about **lexical relations** between words, by which is meant relationships of overlap (or lack thereof) in the words' semantic features. Thus, two words that overlap in all of their semantic features are said to be **synonyms**, as in the case of *car* and *automobile* or *pail* and *bucket*. **Antonyms**, on the other hand, share all of their features except for one – and on that one, they differ in choosing either opposing ends of a continuum (**gradable antonyms**, like *hot* and *cold*) or different choices from a set of exactly two options (**complementary antonyms**, like *dead* and *alive*). Contrary to what one might expect, then, antonyms are actually very much alike: *Hot* and *cold* have a great deal in common semantically, since both are adjectives describing temperature; they differ only in which end of the temperature scale they pick out. Gradable antonyms are easy to distinguish from complementary antonyms, since gradable antonyms can be modified to represent various points on the scale: Food can be *very hot* or *somewhat hot*, and some foods can be *hotter* than others. This is not true for complementary antonyms. While it's possible to say that a party is *really dead* or that an individual is *very alive*, these are metaphorical and relatively uncommon uses; aside from very esoteric medical discussions of, perhaps, brain death vs. heartbeat, one cannot speak in any literal way of one person being more alive than another. In the case of complementary antonyms, to not be in the category described by one word is to be in the category described by the other, assuming the categories can be appropriately applied at all. That is, as long as the entity in question is the sort of thing to which terms like *alive* and *dead* may be applied (e.g., it's a rosebush or a goldfish, not a house or a coffee mug), it is necessarily either alive or dead; if it is not alive, it is necessarily dead, and vice versa. This is not the case with gradable antonyms; if one is not *cold*, it is not necessarily the case that one is *hot*. In short, gradable antonyms permit variance along a continuum, whereas complementary antonyms present an either-or situation.

Hyponymy is also a case of feature-sharing, but in this case one word (the **hyponym**) shares all of the features of another (the **superordinate**) as well as others. For example, *poodle* incorporates all of the meaning of the word *dog*, plus more. This results in a taxonomic relationship that can be drawn in tree form, as shown in Fig. 1.1.

While *poodle* and *collie* are hyponyms of *dog* (their superordinate), *dog* is in turn a hyponym of *mammal*, sharing all of the semantic features of *mammal* (fur, milk production, etc.) and more. That is, a word can simultaneously be a hyponym of one word and a superordinate of another, just as *dog* is a hyponym of *mammal* while being a superordinate of *poodle*.

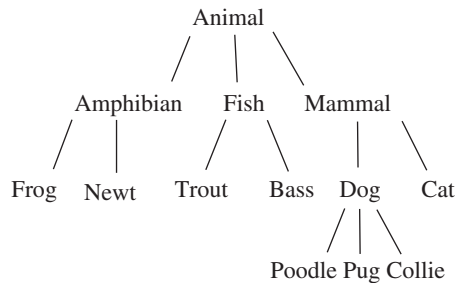


Fig. 1.1 Taxonomy

Homonyms result from two distinct words having the same form, as with *light* (meaning “not heavy”) and *light* (meaning “illumination”). Such a situation results in **lexical ambiguity** – that is, a case of a single lexical form having two distinct meanings. An **ambiguous** word, phrase, or sentence is simply one that has two or more distinct meanings. Ambiguity is to be distinguished from **vagueness**, in which the boundaries of what the term applies to are indistinct. The word *pleasant* is vague, in that there’s no clearly defined cut-off between what is and isn’t pleasant, whereas the word *present* is ambiguous, in that it can mean, for example, either “gift” or “current time”, but neither of the two meanings of *present* is particularly ill-defined in terms of what it applies to.

It might seem intuitively correct to describe homonyms as a single word with more than one meaning, but it’s important to recognize that *light* and *light* under the different meanings described above are actually two distinct words that happen to have the same form. This situation is to be distinguished from the case of **polysemy**, in which a single word has two related meanings, as with *nickel* (the coin) and *nickel* (the metal). This is a subtle but important distinction. In the case of polysemy, the two meanings are clearly related, and the fact that the two meanings are expressed via the same lexical form is not accidental. Most dictionaries acknowledge the distinction in the way that they list words; *bat* (the mammal) and *bat* (the baseball implement) will have separate entries in recognition of their status as homonyms, while *diamond* (the geometric shape) and *diamond* (the baseball field) will be listed as subentries under a single main entry. There are, however, very tricky cases. For example, should *ruler* (a monarch) and *ruler* (a measuring stick) be considered a case of homonymy or polysemy? The answer may differ from person to person; some people recognize the relationship between the two meanings (either historically, in that measuring sticks originally used monarchs’ hand and foot lengths for measurement standards, or synchronically, in that both monarchs and measuring sticks “govern” some domain), whereas others don’t. If our goal in linguistics is to describe linguistic competence, that competence will vary from person to person; one person’s homonymy may well be another’s polysemy.

As noted above, the meaning of a word is often taken to be that set of features by which we know that the object in question is an instance of the category described by the word; thus, the meaning of the word *woman* might be composed of the features +female and +adult, which would distinguish the potential referents of the word *woman* from those of, say, *girl*, which would be +female, –adult.

This is the approach of **componential semantics**, which attempts to boil down the meanings of words to a set of **primitive features**. But now we have a number of problems. First is the binary nature of the features. To define *woman* as +female and +adult ignores the existence of nonbinary people, and more generally the rarity of true binaries. Allowing a feature marker of +/- (e.g., the word *child* would be marked +/-adult) helps in some cases, but not for features that are continuous or have multiple degrees or levels.

In addition to the difficulty of finding actual binaries is the problem of finding words that are actually definable by a small number of features. To return to the case of *woman*, even if we were to take “+/-female” to be a useful binary category distinguishing females from all other genders, it doesn’t solve our problem. For instance, what about the meaning of the word *mare*? Using only the features “female” and “adult”, it will be identical to *woman* – i.e., +female, +adult. But obviously the words *woman* and *mare* don’t mean the

same thing; they don't pick out the same entities in the world. So we'll need to add features to distinguish them – say, *equine* and *human*:

	female	adult	human	equine
woman	+	+	+	–
girl	+	–	+	–
mare	+	+	–	+

So far, so good. But now what happens when the words *hen* (an adult female bird) and *doe* (an adult female deer) come along? Using the features listed above, they will be indistinguishable from each other; we will need to add *avian* and *cervine* (the parallel adjective for “deer”) as features. And no sooner will we decide that things are now in order than *sow* (adult female porcine) will come along to disturb the works, requiring yet another feature:

	female	adult	human	equine	avian	cervine	porcine
woman	+	+	+	–	–	–	–
girl	+	–	+	–	–	–	–
mare	+	+	–	+	–	–	–
hen	+	+	–	–	+	–	–
doe	+	+	–	–	–	+	–
sow	+	+	–	–	–	–	+

Clearly this could go on for a very long time, with a new feature required for every species in which a female adult has a lexicalized form (and there's a lot of them!).

Yet another difficulty with componential semantics is that for many lexical items, it's not at all simple to determine what the correct set of semantic features would be. For example, what are the features that constitute the meaning of the word *sandwich*? Does an object have to include two slices of bread to count as a sandwich? Apparently not, since open-face sandwiches exist. Does bread have to be involved at all? What about a pita sandwich? What about a taco? This precise question has real-world consequences: In 2006, a Massachusetts judge ruled that a burrito is not a sandwich. A Panera Bread cafe had a stipulation in its lease preventing the opening of another sandwich shop in the same shopping center. At issue was the opening of a Qdoba outlet, which sold burritos. Panera argued that a burrito is a sandwich; the judge disagreed. What set of primitive features would determine that a meat-filled pita is a sandwich while a meat-filled tortilla is not?

As an alternative to componential semantics, **Prototype theory** offers a way of dealing with such issues. According to this approach (Rosch 1973, 1975), the meaning of a word is a **fuzzy set**, that is, a set whose boundaries are indistinct, or “fuzzy”. The set contains a central member, or **prototype**, that constitutes the “best” example of the set in question; for example, the prototypical sandwich might consist of two slices of bread with sliced meat and cheese between them, and a condiment such as mustard. Other combinations will be more or less sandwich-like depending on their resemblance to this prototype, and toward the fuzzy boundary of the set there will be cases whose membership in the class is debatable, including stuffed pitas, tacos, and burritos.

This view accommodates the existence of gradient categories, as well as the instinct that for some categories (and therefore the words that express them), there simply isn't a firm boundary between category members and non-category members. This, in turn, explains the ongoing battles over whether a hot dog does or does not count as a sandwich: It's not a prototypical sandwich but rather a borderline member of the category; and indeed whether it counts as a category member at all will depend on how your own personal fuzzy set for *sandwich* is constructed. (For one person, a ham and cheese is the prototypical sandwich; for another, a PBJ.) Under this approach, there's no one set of criterial features that determine what's in and what's out of the set. A child acquiring language bases their prototype for a given word, and ultimately the set characterizing that word, on what they hear the word applied to. All of which is to say that you and I will most likely have very similar sets and prototypes for the word *sandwich*, but they won't be identical, and the fate of the hot dog – for each of us – hangs in the balance.

1.3.1.2 Sentence meaning

It is intuitive to think of the meaning of a sentence as the sum of its parts – that is, that determining the meaning of *Sheila won the tournament* is simply a matter of combining the meanings of the words *Sheila*, *won*, *the*, and *tournament*. And to a great extent, this is the case. A **compositional** semantics is one that takes the meaning of a sentence to be essentially the sum of its parts, in combination with a set of rules governing the way in which the meaning of the sentence is built up from the meanings of its components in light of the syntactic structures in which they are placed; that is, *Mary loves frogs* does not mean the same thing as *Frogs love Mary*, and our linguistic theory must be able to explain why. Thus, the fields of syntax and semantics overlap significantly in their areas of concern.

Just as the meanings of words can overlap partially (hyponymy) or completely (synonymy) or can be in opposition (antonymy), these semantic relations have analogs at the sentence level. For instance, **redundancy** is a case of partial repetition of meaning, as in *The child plodded slowly across the yard* (where *plod* entails *slowly*) or *My female sister is tall* (where *sister* entails *female*). As these examples illustrate, the effect of the redundancy can range from the hardly noticeable to the patently ridiculous. Notice also that hyponymy within a sentence can give rise to redundancy: *Sister* is a hyponym of *female* (i.e., *sister* includes the meaning of *female* plus more), which is what makes the sentence *my female sister is tall* redundant. Complete overlap of meaning results in **paraphrase**; for example, *My brother is older than me* is a paraphrase of *I am younger than my brother*. In this case, the paraphrase relationship is due to the lexical relationship between *older* and *younger*, but here again, the paraphrase can be due to synonymy at the lexical level: *My couch needs to be cleaned* and *My sofa needs to be cleaned* are paraphrases due to the synonymy of *couch* and *sofa*. As we will see in the next section, paraphrases are characterized by the fact that the two sentences are true under the same set of conditions; that is, if one is true, the other is necessarily true, and if one is false, the other is necessarily false as well.

Similarly, **antonymy** at the lexical level can give rise to **anomaly** – a clash of semantic meaning – at the sentence level, as with *?The water is quite hot, and very cold*. (Throughout this text, a question mark before a sentence or clause will indicate that it is anomalous.) Not all anomaly is attributable to antonymy; consider, for example, Noam Chomsky's

famous sentence *Colorless green ideas sleep furiously* (Chomsky 1957). Here, it seems that virtually every pair of words in the sentence clash with each other: Nothing can be both green and colorless, ideas by their nature can be neither green nor colorless, ideas can neither sleep nor do anything furiously, and it is hard to imagine what it would be to sleep furiously. Thus, the sentence is wildly anomalous. Nonetheless, it is syntactically flawless, i.e. grammatical, and this was precisely Chomsky's point: He used this sentence to show that syntax and semantics are distinct, and specifically that our knowledge of the rules of syntax is autonomous – independent of the meaning of any particular sentence. The syntactic correlate of semantic anomaly is **ungrammaticality**, as in **Dog the small slept the red rug on*. (Ungrammaticality will be indicated in this text with an asterisk.)

Finally, lexical ambiguity can give rise to ambiguity at the sentence level, as with *George walked down to the bank* (where *bank* could mean “river bank” or “financial institution”). But sentences may also exhibit **structural ambiguity**, due to the existence of two distinct syntactic analyses for the sentence, as in *Jenny ate the pizza on the table*, in which either Jenny or the pizza might be on the table, depending on the structure assigned to the sentence, specifically how much of the postverbal material is taken to be part of the direct object: *Jenny ate [the pizza on the table]* vs. *Jenny ate [the pizza] on the table*.

1.3.1.3 Formal logic and truth conditions

Semantic meaning is often represented using formal notation borrowed from the study of formal logic. It's important to understand the analysis of certain English connectives in formal logic, because some of the seminal works in pragmatic theory take these analyses as their starting point.

First, it is useful to distinguish between **deductive** and **inductive** logic. Deductive logic involves rules for drawing necessarily valid inferences from a set of propositions. These propositions are called **premises**, and a valid inference we can draw from a set of premises is called the **conclusion**. For example:

Premises: All students love linguistics.
 Hinkelmeyer is a student.
 Conclusion: Hinkelmeyer loves linguistics.

The conclusion is **entailed** by the premises. This means that there is no situation in which the premises could be true and the conclusion false. But notice that the validity of the deduction is totally independent of the actual truth of the premises and conclusion. It could be the case, in reality, that NOT all students love linguistics, and even that Hinkelmeyer herself despises linguistics. Nonetheless, the deduction above is valid: There is no situation in which the premises could be true and the conclusion false. This is not altered by the fact that the premises themselves may not actually be true.

Inductive logic, on the other hand, is a matter of probability. Inductive inferences are not **necessarily** true, as deductive inferences are. Here's an example of an inductive inference:

Premises: The sun has risen every day of this century.
 Tomorrow will be a day of this century.
 Conclusion: The sun will rise tomorrow.

This conclusion is very likely to be true, but it is not necessarily true by virtue of the premises. That is, the fact that the sun has risen every day of this century thus far does not in itself guarantee that it will rise again tomorrow.

Formal logic concerns itself with deductive inferences – that is, with flawlessly valid inferences. It's interesting to note that scientific experiments, on the other hand, are generally designed to lead to inductive inferences – inferences that are not necessarily true. Let's say we form a hypothesis – say, that if I hold a book three feet above the floor and let go, it will fall to the floor. And let's say I perform the experiment of releasing a book from three feet above the floor 10,000 times, and each time that I let go of the book, it falls to the floor. Based on these experiments, I may confidently infer that a book held three feet above the floor and released will always drop to the floor. But notice that this is an inductive inference; it leaves open the possibility that on the 10,001st trial, the book will fail to fall to the floor. This may be unlikely, but it is a logical possibility. And indeed, if on the 10,001st trial my friend walks in and catches the falling book before it hits the floor, my hypothesis will have been falsified and will need to be revised.

For this reason, the results of scientific experiments are typically reported along with a numerical value indicating the degree of confidence in the study's conclusions, expressed in terms of probability, or a *p-value*: " $p < 0.01$ " indicates that there is less than a 1-in-100 chance that the conclusion is wrong and that the results are due to chance. Put another way, this *p-value* indicates a 99 percent confidence in the reliability of the findings. This is one reason why it's so important that a scientific hypothesis be in principle falsifiable: Since it's impossible to confirm beyond a doubt that the claim is true (10,000 instances of dropping a book on the floor are insufficient for certainty), it is necessary to at least know what sort of circumstance would confirm that it is necessarily false (a single instance of my friend catching it as it falls).

As noted above, formal semantics employs the notation of formal logic, which it uses as a neutral, connotation-free language for expressing the meanings of **sentences**. A sentence is a sequence of words, that is, an abstract linguistic object. An **utterance** is a sentence that's produced in some actual context (whether oral, written, or signed, as in American Sign Language). There are many sentences that have never been uttered and never will be; it's quite likely, for example, that nobody has ever before uttered the sentence *My chihuahua's favorite lampshade is submerged in the meringue*, even though it's perfectly interpretable. A **proposition** is what a sentence expresses. Thus, the sentence *I read the assignment today* can be used to express very different propositions depending on who utters it and when. And just as a single sentence can be used to express many different propositions, a single proposition can be expressed in a variety of sentences; *Mary spoke to Jane* and *Jane was spoken to by Mary*, for example, express the same proposition.

A proposition will be true in some **possible worlds** and false in others. A possible world is precisely what it sounds like: a way that the world could be, or could have been. The idea is that the world we happen to be living in isn't the only possible world. So the proposition "all dogs are blue" happens to be false in the real world, but there's another possible world – another way the world could have happened to be – in which it's true. On the other hand, the proposition "if a dog is blue, it is blue" is true in all possible worlds. There is no possible world in which this proposition could be false; it is necessarily true. An **analytic** sentence is one whose truth is independent of what the world is like; it's either necessarily true (as in *if a dog is blue, it is blue*) or necessarily false (as in *if a dog is blue, it is not blue*). A sentence that is true in all

possible worlds (such as *if a dog is blue, it is blue*) is a **tautology**. A sentence that is not true in any possible world (such as *if a dog is blue, it is not blue*) is a **contradiction**. A sentence whose truth depends on the condition of the world (such as *some dogs are blue*) is **synthetic**. In order to know whether a synthetic sentence is true in a given world, it is necessary to see what that world is like (for example, whether it contains any blue dogs).

The **truth conditions** of a sentence are the conditions under which it would be true – that is, what the world (or any possible world) would have to be like in order for that sentence to be true. The truth conditions of a sentence are independent of what the world actually is like; they're just a specification of what the world **would** be like if the sentence were true. On the other hand, the **truth value** of a sentence in some particular world is a specification of whether the sentence is in fact true in that world. Thus, the truth conditions of the sentence *A blue dog exists* are essentially that the world contains a blue dog, while the truth value of the sentence is T (true) in a world that does contain a blue dog and F (false) in a world that does not. **Truth-conditional** meaning is any piece of meaning that affects the conditions under which a sentence would be true. Thus, the difference between *and* and *or* is truth-conditional, since the sentences in (1.6) and (1.7) are true in different sets of circumstances:

(1.6) All women are tall and all women are smart.

(1.7) All women are tall or all women are smart.

In a world in which all women are smart but not all women are tall, (1.6) would be false while (1.7) would be true. However, the difference between *moreover* and *nonetheless* is not truth-conditional:

(1.8) All women are tall; moreover, all women are smart.

(1.9) All women are tall; nonetheless, all women are smart.

The sentences in (1.8) and (1.9) will be true under the same set of circumstances; there is no possible world in which one of them is true and the other false. There is, of course, an additional piece of meaning that's conveyed in (1.9); here you understand the speaker to be suggesting that in the context of all women being tall, there is something unexpected about their also being smart. By saying that this piece of meaning is non-truth-conditional in (1.9), we don't mean that the sentence *There is something unexpected about all women being smart* has no truth conditions; it obviously does. Rather, we mean that its truth conditions play no role in determining the truth conditions of (1.9), and likewise that its truth value (i.e., whether it is in fact the case that this is unexpected) plays no role in determining the truth value of (1.9) when it's uttered.

The study of logical relationships between sentences is called **propositional calculus** or **propositional logic**. In propositional calculus, *p*, *q*, and *r* stand for propositions, and they are connected by various **logical connectives** such as *and* and *or*. The logical connectives can be viewed as functions that map truth values (or sets of truth values) onto truth values. For example, take **negation**:

<i>p</i>	$\neg p$
t	f
f	t

This is called a **truth table**. What it tells us is that anytime p is true, $\neg p$ (“not- p ”) is false, and anytime p is false, $\neg p$ is true. Thus, negation is a function that maps **t** in the first column onto **f** in the second, and vice versa. In each row, the values to the left of the double line give us the truth value(s) of the given proposition(s) in some world, and the values to the right of the double line tell us what that means for the values of the propositions in combination with the given connectives. In the little truth table above, for example, the first line under p and $\neg p$ represents any world in which p is true; in such a world, $\neg p$ is necessarily false. The second line represents any world in which p is false; in such a world, $\neg p$ is necessarily true. Thus, if *All fish have fins* is true, then *It’s not the case that all fish have fins* must be false, and vice versa. While negation isn’t technically a connective (since it doesn’t connect two propositions), it is typically grouped with the logical connectives because, like the logical connectives, its meaning is defined as a function from truth values to truth values. Notice that it doesn’t matter what the proposition in question (p) is; the effect of negation will be the same regardless of the particular meaning of p .

The truth table for **conjunction** (‘and’, symbolized $\wedge$) is slightly more complicated, since it involves two propositions:

p	q	$p \wedge q$
t	t	t
t	f	f
f	t	f
f	f	f

What this table tells us is that $p \wedge q$ is only true when both p and q are true (the first line). In all other cases, $p \wedge q$ is false. That is to say, *All monkeys are mean and all buffalo are brave* is false if either *All monkeys are mean* is false or *all buffalo are brave* is false, regardless of the truth of the other conjunct.

Here’s the truth table for **disjunction** (‘or’, symbolized $\vee$):

p	q	$p \vee q$
t	t	t
t	f	t
f	t	t
f	f	f

What this table tells us is that $p \vee q$ is false only when both p and q are false (the fourth line); in all other cases, it’s true. This is the truth table for what’s known as **inclusive ‘or’**, meaning “one or the other or both.” On this reading of *or*, the sentence *All monkeys are mean or all buffalo are brave* is true if either all monkeys are mean or all buffalo are brave, regardless of the truth of the other conjunct.

The truth table for **exclusive ‘or’**, meaning “one or the other, but not both,” would be:

p	q	$p \vee q$
t	t	f
t	f	t
f	t	t
f	f	f

Here, if both propositions are true, the entire disjunction is false (line 1). This would be the meaning generally intended in the utterance of a sentence such as *I'll pay you tomorrow or the day after* (where the speaker doesn't intend to leave open the possibility of paying on both days). Exclusive "or" is usually assumed to be derived via a pragmatic inference; that is, truth-conditionally "or" is assumed to have only the inclusive meaning, but in many contexts hearers infer that it's not the case that both conjuncts are true, because if they were (and if the speaker knew they were), the speaker should have used "and".

Here's the truth table for **logical implication** (aka the **conditional**, or "if . . . then", symbolized $\rightarrow$):

p	q	$p \rightarrow q$
t	t	t
t	f	f
f	t	t
f	f	t

This one is highly counterintuitive, and tends to trip people up. Notice that the only case in which $p \rightarrow q$ is false is the case in which p is true and q is false. That is, what implication ($\rightarrow$) says is that the truth of p guarantees the truth of q . If p is false, however, q can be anything and $p \rightarrow q$ will still be true. Think of what's meant by the statement *If you're a genius, then I'm a monkey's uncle*. This is a statement of the form $p \rightarrow q$, where q is clearly false (since I'm not a monkey's uncle). Since q is false, the only way for the statement as a whole to be true is for p to also be false (check the chart!). So this is a roundabout way of conveying that p ("you're a genius") is false ("if you're a genius, then I'm a monkey's uncle – but since I'm not a monkey's uncle, you must not be a genius"). That verifies the fourth row.

Now consider the counterintuitive third row. Suppose you're a great baker, and I'm a big fan of chocolate. On my birthday, you bring me something in a cake box. I say *If this is a chocolate cake, it will be delicious*. Now suppose we're in the world described by line 3: It's not actually a chocolate cake, it's a blueberry pie, and it is in fact delicious. Was my statement true or false? Or suppose a friend of mine, knowing I live in the Chicagoland area, sends me a text saying *If you're home next week, I'll be in Chicago*. Again assume we're in row 3: I'm not going to be home, but my friend will indeed be in Chicago. Have they said something false? It would seem not; rather, it seems that in both of the worlds in which my friend will actually be in Chicago (i.e., rows 1 and 3), they've said something true in saying *If you're home next week, I'll be in Chicago*. But don't worry – if this still feels wildly counterintuitive, just remember that logic and natural language are very different things.

Finally, here's the truth table for **equivalence**, or **bidirectional implication** ("if and only if", also known as "iff", symbolized as $\leftrightarrow$ or $\equiv$):

p	q	$p \leftrightarrow q$
t	t	t
t	f	f
f	t	f
f	f	t

This means that $p \leftrightarrow q$ is true in exactly those situations in which p and q have the same truth value; otherwise it's false. The last line shows that a bidirectional that joins two false propositions is true, as in *The Sun is smaller than the Earth if and only if the Earth is larger than the Sun*. Although both of the smaller propositions are false, the statement as a whole is certainly true.

With more complicated truth tables, you can check whether quite complicated formulae are true under various sets of conditions. For example, take $(p \wedge q) \rightarrow (\neg r \vee p)$:

p	q	r	$(p \wedge q)$	$\neg r$	$(\neg r \vee p)$	$(p \wedge q) \rightarrow (\neg r \vee p)$
t	t	t	t	f	t	t
t	t	f	t	t	t	t
t	f	t	f	f	t	t
t	f	f	f	t	t	t
f	t	t	f	f	f	t
f	t	f	f	t	t	t
f	f	t	f	f	f	t
f	f	f	f	t	t	t

Here we start out with every possible combination of t/f values for p , q , and r , as seen in the first three columns. Each horizontal row, then, corresponds to one possible way the world could be. So, in the first row, we're asking what happens if, in some possible world, p , q , and r are all true. Well, in that world, $(p \wedge q)$ is also true, so we put that in the fourth column. And since r is true, $\neg r$ is false, so we put that in the fifth column. Since $\neg r$ is false and p is true, $(\neg r \vee p)$ is true, which goes in the sixth column. And since $p \wedge q$ is true (fourth column) and $(\neg r \vee p)$ is also true (sixth column), $(p \wedge q) \rightarrow (\neg r \vee p)$ is also true (last column). Notice that the parentheses serve to group things together, as in math – so that $p \wedge q$ will be taken as a sub-unit – a constituent – of the larger formula, but $q \rightarrow \neg r$ will not.

Notice also that in this particular example, it turns out that in **every case**, the final formula turns out to be true. This means that there is no possible world in which this formula could be false; it's true regardless of what the world looks like – regardless of whether its component propositions (p , q , and r) are true. Such a formula constitutes a **tautology**. Any sentence of this form will be necessarily true. For example:

(1.10) If a man is tall and he is smart, then either he is not young or he is tall.

And in fact, this sentence is indeed true in all possible worlds; all tall, smart men either are not young or are tall (since they're all tall).

A sentence that is necessarily **false** in all possible worlds is a **contradiction**. The truth table for a contradiction will have all f's in the final column. For example, $p \wedge \neg p$ is an obvious contradiction. As noted above, tautologies and contradictions are analytic, meaning their truth value is independent of what a particular world is like; all other sentences are synthetic, meaning that they depend for their truth value on what the world is like. The truth table for a synthetic sentence will have a mix of t's and f's in the final column.

Truth tables and the calculation of truth values for complex propositions are part of **propositional calculus**. Whereas propositional calculus deals with relationships among propositions, **predicate logic** looks at truth-conditional meaning within an individual sentence. For example, the sentence *Sally is a plumber* might be formalized as $P(s)$, where P stands for “plumber” and s stands for “Sally”, and the whole formula states that we are predicating plumber-hood of Sally. The predicates are capitalized, and the **terms** (individuals) are lowercased. *Sally* here is represented by a **constant** (s); each constant represents a specific individual. Alternatively, you can also have **variables**, as with $P(x)$, which means (roughly) “ x is a plumber,” where x is some unspecified entity. You can also have more than one term, or **argument**: $L(s,p)$ might stand for *Sally likes Paul*. In this case, *Sally* and *Paul* are the arguments of *likes*.

Where it all gets interesting is when you bring in **quantifiers**. Quantifiers tell us something about the quantity of entities that a predicate applies to. The two most basic quantifiers are the **universal quantifier** (which specifies that the predicate applies to all entities) and the **existential quantifier** (which specifies that the predicate applies to some entity or entities):

$\forall$ – The universal quantifier, roughly paraphrased as “for all”

$\exists$ – The existential quantifier, roughly paraphrased as “there exists”

Here are some examples of the universal and existential quantifiers at work. (For ease of exposition, these examples make the simplifying assumption that all of the entities in the universe of discourse are people.)

$\forall x(P(x))$ – “For all x , x is a plumber,” or “Everyone is a plumber”

$\exists x(P(x))$ – “There exists an x such that x is a plumber,” or “Someone is a plumber”

You can mix and match predicate logic and propositional calculus:

$\forall x(P(x) \rightarrow T(x))$ – “For all x , if x is a plumber, then x is tall,” or “All plumbers are tall”

Notice that some subtle meaning differences, known as **scope** differences, may depend on the ordering of the quantifiers. The outside quantifier (that is, the leftmost one) is said to have scope over the inside quantifier (as well as having scope over everything in the set of parentheses following it). So if you have, for example, “ $\forall x\exists y$ ”, that universal quantifier has scope over the existential, and you’d read it as “for all x , there exists some y . . .”, whereas if you have “ $\exists x\forall y$ ”, the existential has scope over the universal (said another way, the universal is within the scope of the existential), and so it’s read as “there exists some x such that for all y . . .”. In short, everything with “narrow” scope (in this case, the second or inside quantifier) is interpreted with respect to the quantifier with “wide” scope (the outside one). So let’s look at an interesting set of examples:

$\forall x\exists y(L(y,x))$ – “For all x , there exists a y such that y loves x ,” or “Everyone is loved by someone”

$\forall x\exists y(L(x,y))$ – “For all x , there exists a y such that x loves y ,” or “Everyone loves someone”

$\exists x \forall y (L(x,y))$ – “There exists an x such that for all y , x loves y ,” or “There is someone who loves everyone”

$\exists x \forall y (L(y,x))$ – “There exists an x such that for all y , y loves x ,” or “There is someone who is loved by everyone”

Notice that each of these describes a different world: In the first case, the formula says that for any given person, someone loves them, whereas the second says that for any given person, there’s someone they love, and so on. Notice also that the English gloss given for the first case, “everyone is loved by someone”, is strictly speaking ambiguous: It could mean either that every individual has someone who loves them, or (on a less common reading) that there’s some individual who loves everyone. The formula, however, does not share this ambiguity; it can only mean that each individual has someone who loves them. (The alternative reading is captured by the third formula.) Thus, the use of logical notation has the advantage that it is unambiguous, which makes it useful for expressing precise meanings where natural language utterances might be subject to ambiguity or misinterpretation. There is a great deal more complexity to the study of logic and linguistics (the treatment of tense, for example, or modals such as *might* and *could*), but this will be sufficient for our purposes.

1.3.2 The domain of pragmatics

The word *meaning* is notoriously imprecise, both in the sense of being ambiguous and in the sense of being vague. We have seen that the field of semantics deals with one sort of meaning – the conventional, context-independent meaning of a word or sentence, such as a dictionary might try to capture in a definition or a logician might try to capture in predicate logic notation and truth tables. While it might seem straightforward to state that pragmatics simply covers whatever aspects of meaning are left over, the issue turns out to be far more complex. We will begin by describing the various uses of the word *meaning*, after which we will consider possible ways of delimiting the boundary between semantics and pragmatics.

1.3.2.1 Nonnatural meaning

If we’re going to discuss meaning, it makes sense to first have a notion of what the word *meaning* itself means. Philosopher H.P. Grice, possibly the most influential figure in the field of pragmatics, observed that **meaning** is far from a unitary notion. Consider the following sentences:

- (1.11) a. That clap of thunder means rain is coming.
 b. *Supercilious* means “arrogant and disdainful”.

Each of these contains the word *means*, but the word is being used in two very different ways in the two cases. In the first case, the meaning in question is what Grice (1957) calls **natural meaning** – an indication that is independent of anybody’s intent. A clap of thunder indicates that rain is coming independently of whether anybody intends for that indication to be present, either on this particular occasion or in general: Nobody has arranged for

this particular clap of thunder to have this particular meaning, and more generally, the correspondence between claps of thunder and subsequent rain was not set up with the intent that the presence of the former should convey the imminent presence of the latter. The type of “indication” present here is merely a matter of our having noticed, after years of observation, that there is a correlation between the two events.

In the case of (1.11b), on the other hand, there clearly is an intent that the word *supercilious* be taken to mean “arrogant and disdainful”. Someone who uses this word intends that the word/meaning correlation be recognized by their interlocutor. This meaning is **nonnatural**, in Grice’s terms; there is no automatic, natural correlation between the word and its meaning. Instead, the word/meaning correlation is arbitrary; this meaning could just as easily have ended up being attached to another string of sounds, had the history of the language worked out differently. With the exception of onomatopoeia (the phenomenon of words “sounding like” what they stand for, as with *crash* and *tweet*), the vast majority of words in a language exemplify nonnatural meaning. We understand them not because of a natural relationship between the sound and the meaning, but because we as a society have agreed to arbitrarily correlate the sound with the meaning in order to use the former intentionally to evoke the latter.

1.3.2.2 Sense and reference

Within linguistic (hence nonnatural) meaning, there is another important distinction to be made between the “dictionary” sort of meaning of a word and what it is used to refer to in the world. Consider (1.12a–b):

- (1.12) a. *Supercilious* means “arrogant and disdainful”.
 b. When the judge asks the defendant to rise, she means you.

In (1.12a), it is the **sense** of the word that is at issue – that is, the sort of meaning that a dictionary would give for the word. The sense of the word *chair* is what one must have access to in order to answer the question “is this a chair?”; similarly, the sense of the word *supercilious* is what one must have access to in order to answer the question “is this person being supercilious?” This meaning is more or less invariant; that is, *supercilious* means “arrogant and disdainful,” regardless of who utters it, when, and under what circumstances. In (1.12b), on the other hand, the meaning in question is a matter of what particular entity is being picked out, or referred to – that is, the **referent** of the expression. Frege (1892) developed the distinction between sense and reference (in his native German, *Sinn* and *Bedeutung*) using the example of the phrases *the morning star* and *the evening star*, which have the same referent – the planet Venus – but obviously different senses, since *morning* and *evening* have different senses. (Strictly speaking, he used the two names for these senses, Phosphorus and Hesperus, respectively, but it’s vastly easier to make the point without them.)

Unlike sense, it’s possible for reference to vary in different contexts: On one occasion, a judge may use the phrase *the defendant* to refer to John Doe, and on another, to refer to Jane Snow, depending on the trial in question. A week after (1.12b) is first uttered, it might be uttered again with a different referent, whereas a week after (1.12a) is uttered,

supercilious will still mean “arrogant and disdainful.” Thus, sense is a context-independent, purely semantic notion, whereas determination of reference may require access to pragmatic information.

1.3.2.3 Speaker meaning vs. sentence meaning

The distinction between sense and reference described in the previous section is related to the distinction between **sentence meaning** and **speaker meaning**. Sentence meaning is the literal meaning of a sentence, derivable from the sense of its words and the syntax that combines them. Sentence meaning is “sense” as applied to entire clauses rather than individual words and phrases. Speaker meaning, on the other hand, is the meaning that a speaker intends, which usually includes the literal meaning of the sentence but may extend well beyond it. Thus, consider (1.13):

- (1.13) A: I’m cold.
 B: Does that mean I should build a fire? (In conversation, 2022)

The **sentence meaning** in the first line here is straightforward: The speaker is cold. The **speaker’s meaning** in using this utterance in a particular context, however, could be any of a number of things, including:

- (1.14) a. Close the window.
 b. Bring me a blanket.
 c. Turn off the air conditioner.
 d. Snuggle up closer.
 e. The heater is broken again.
 f. Let’s go home. [uttered, say, at the beach]
 g. Please build a fire.

The possibilities are limited only by one’s imagination. (One could imagine, for example, a rather dull crime novel in which the phrase *I’m cold* is used as a code to mean *We steal the jewels at midnight* – a case in which the sentence meaning is not, in fact, part of the speaker meaning.) Speaker meaning is also sometimes called **utterance meaning**; if you recall the difference between a sentence (which is an abstract entity) and an utterance (an instance in which a sentence is actually used), you will see that the meaning of a sentence is context-independent, whereas the meaning of an utterance is context-dependent and depends in particular on the intentions of the speaker. Speaker meaning, therefore, is a pragmatic notion, while sentence meaning is semantic.

1.3.2.4 Possible worlds and discourse models

Although we talk about linguistic communication as though it involved a straightforward transfer of information – saying things like *I got my ideas across* or *Let me give you my thoughts on that* or *He conveyed several notions to us in his talk* – this is actually a misleading way of thinking about language, as observed by Reddy (1979). My thoughts

never leave my head and travel to yours; instead, Reddy points out, the hearer must attempt to reconstruct the speaker's intended meaning, and this is a process that is fraught with the possibility of miscommunication and misunderstanding. Linguistic communication, far from being effortless and unidirectional, is essentially **collaborative** in nature; as a speaker, my goal is typically to help my hearer develop an internal representation of the discourse that matches my own, while the hearer's goal is correspondingly to develop such a representation. This representation of the discourse is called a **discourse model**. Consider, for example, the beginning of the second paragraph of Trevor Noah's *Born a Crime*:

- (1.15) I was nine years old when my mother threw me out of a moving car. It happened on a Sunday. I know it was on a Sunday because we were coming home from church, and every Sunday in my childhood meant church. (Noah 2016)

Upon encountering this text, the reader already has a discourse model that includes Trevor Noah, the author. On reading the first sentence, they will add to this model Noah's mother, a moving car, and the fact that at some point his mother threw him out of the car. The second sentence adds the fact that it was a Sunday. At this point the reader has three "discourse entities" – Noah, his mother, and the car – along with an event and the day of the event; in addition, the discourse entities for Noah and his mother are tagged with the property of always attending church on Sunday, Noah is tagged with the property of being nine years old, and his mother is tagged with the property of having thrown him from a moving car. And so on. As the reader continues, their model of the discourse continues to grow, with more entities and events being added, more properties being attached to them, and more cross-references among them (entity A is Noah's mother; entity B is entity A's son).

Each of the interlocutors has a distinct discourse model, and at the moment of any given utterance, one of the speaker's goals (assuming the speaker has no deceptive intent) is usually to increase the similarity between their own discourse model and that of their interlocutor(s), by, for example, expressing information, asking for information, or making a request (which in turn causes the addressee's discourse model to include this desire on the part of the speaker).

It was noted above that a **possible world** is some way the world could have been (and of course one possible world is the actual world – that is, the way the world in fact is). A discourse model maps onto a set of possible worlds – a set of worlds in which the information in the discourse model holds true. It's important to realize that a discourse model does not necessarily represent reality – and that even if a speaker or hearer believes their discourse model represents reality, they may be mistaken; the discourse model may be inaccurate as a representation of reality. Consider again the discourse model a hearer constructs upon reading (1.15) above. This model matches the real world in many ways, but might not match it in others; for example, the reader might construct a mental image of the scene as involving a red car, and might even be surprised to learn later that no such detail was mentioned. People's memories even for events that involved them are notoriously unreliable.

We can also construct discourse models for possible worlds we share but which are not the real world; for example:

- (1.16) Han Solo loves Leia.

Nearly everybody familiar with the *Star Wars* movies would immediately judge this sentence to be true, even though the entities in question don't exist in the real world. But we share a discourse model containing certain entities with certain attributes, and thus can make judgments about the set of possible worlds it represents. In those possible worlds, (1.16) is true. (Of course, there are also possible worlds in which (1.16) is false, but for me to make statements based on one of those worlds would be uncooperative if I'm in conversation with someone who, based on their status as a member of modern American society, would reasonably expect me to share the standard *Star Wars* discourse model.)

We can also construct possible worlds on the fly for the purposes of the current discussion, as with conditionals:

(1.17) If we have a pop quiz in class today, I'm going to fail.

Here, the speaker evokes a set of possible worlds in which there is a pop quiz in class today, in order to state that in any of these possible worlds, the speaker will fail. And again, it's also common for a discourse model to be inaccurate; a speaker may utter a sentence such as (1.18) in error:

(1.18) We're having a pop quiz in class today.

This may be spoken either with the intent to deceive or via an innocent error (where the speaker for some reason mistakenly believes that there is a pop quiz in class today); in either case, the hearer is likely to be unwittingly left with a discourse model that does not conform to reality, and in the latter case both interlocutors' models will fail to conform to reality. It's impossible to know whether one's discourse model truly reflects reality, since even our perceptions may be in error, as may our interpretations of those perceptions. Thus, in any ongoing discourse, our discourse model reflects only our **beliefs** concerning the possible world under discussion – which in turn may be the real world or some fictional or hypothetical world. And it is entirely possible for both interlocutors to be similarly misinformed about the state of reality, such that an entire conversation is held successfully about some object that in fact does not exist – and it is entirely possible that neither of the interlocutors, and in fact nobody in the world, ever becomes aware of the error.

A discourse model, then, is a mental model whose correlation with reality can be believed in (and even supported by a great deal of evidence), but never definitively established. Nonetheless, when we utter a sentence such as *Sally is a plumber*, we feel that we are positing a property of an actual individual in the world, not of a concept in our minds; that is, we are not trying to say that a concept is a plumber. Thus, there are problems with both the **mental-ist** point of view, which holds that the referents of linguistic expressions are mental entities, and the **referential** point of view, which holds that the referents of linguistic expressions are real-world entities (de Swart 2003). The relationship among linguistic expressions, mental constructs, and the real world is still a topic of debate within semantics and pragmatics.

1.3.2.5 Mutual belief

Because my discourse model is distinct from yours, and because we can never check the extent to which my model and your model agree, the best we can do is to operate on the assumption that we share beliefs. Thus, for me to refer **feliculously** (that is, in a pragmatically

appropriate way) to *today's pop quiz*, I must believe that you believe there is a pop quiz today. And for the reference to be successful, you must believe, first, that there's a pop quiz, and second, that I also believe there's a pop quiz, and third, that I believe that *you* believe that there's a pop quiz. Imagine what would happen if even just the last step were missing: Suppose the pop quiz is a secret (as they generally are), but my friend Judy, who saw the quiz on the instructor's desk, has told me about it. So (a) is true:

- a. I believe there's a quiz.

Now suppose Judy has also told you about the quiz. Now both of the following are true:

- a. I believe there's a quiz.
- b. You believe there's a quiz.

This is insufficient for me to felicitously ask you *How do you feel about the quiz?* Since I don't know that you know there's a quiz, I wouldn't refer to it, since I'd assume you wouldn't know what quiz I'm talking about. Now suppose Judy has also told me that she told you. Now we have:

- a. I believe there's a quiz.
- b. You believe there's a quiz.
- c. I believe that you believe there's a quiz.

In theory, this still isn't enough for the communication to go through successfully, because you still don't know (c). Let's say, for example, that the pop quiz is in Linguistics, but there's also a less salient quiz that was given yesterday in our History class. So if I ask you *How do you feel about the quiz?*, you might plausibly reason, "Well, she doesn't know that I know about the Linguistics quiz. So she can't mean *that* quiz, because she'd assume I wouldn't know about it. She must mean the History quiz."

Finally, suppose Judy has mentioned to you that she told me that she told you. Now:

- a. I believe there's a quiz.
- b. You believe there's a quiz.
- c. I believe that you believe there's a quiz. (That is, I believe (b).)
- d. You believe that I believe that you believe there's a quiz. (That is, you believe (c).)

Is this enough? Well, no. When I ask you *How do you feel about the quiz?*, you might now reason, "Well, she does know that I know about the quiz, but she doesn't know that *I* know that. That is to say, she doesn't know (d). So she would expect me to reason just as I did in the case where (d) didn't hold. Therefore, she must mean the History quiz." You can see how this quickly becomes an infinite regress, with an infinite number of increasingly embedded beliefs being necessary for even this fairly simple utterance. This argument was first and most clearly set forth by Clark and Marshall (1981), who argued that the apparent impossibility of linguistic communication is resolved through a number of **co-presence heuristics**. For example, if you and I are co-present in the room when the instructor announces that there will be a quiz next Tuesday, I can then felicitously utter the noun phrase *the quiz* in conversation with you

and fully expect that you will assume that I'm referring to the quiz that the instructor told us about – on the grounds that we were mutually co-present at the time the instructor made the announcement. Similarly, we can appeal to community co-presence (our being part of the same community and thus sharing certain cultural knowledge) in interpreting phrases like *the sun* or *the President*. Since the overwhelming majority of linguistic interactions do not involve situations as complex as the Judy/quiz situation above, Clark and Marshall argue, we are generally able to bridge the infinitely deep chasm of mutual knowledge through the use of such co-presence heuristics.

1.3.3 Delimiting the boundary

Since both semantics and pragmatics deal with issues of linguistic meaning, it would seem to be crucial to distinguish between the two. However, drawing the boundary is not as straightforward as it might appear. For example, semantic meaning is sometimes identified as **context-independent**, whereas pragmatic meaning is said to be **context-dependent**. Alternatively, semantic meaning is often identified as **truth-conditional** meaning, while pragmatic meaning is often identified as meaning that does not affect the truth conditions of the utterance. While both are true most of the time (that is, that semantic meaning is both context-independent and truth-conditional while pragmatic meaning is context-dependent and non-truth-conditional), there are some cases where the two distinctions do not align perfectly, as we will see below.

1.3.3.1 Context-dependence

The question of context-dependence has to do with whether the meaning of a linguistic form changes with the context in which it is uttered. One commonly used test to see whether some piece of meaning is semantic or pragmatic is to see whether it remains constant regardless of context. For example, consider (1.19):

(1.19) This weather is too cold.

There are a number of elements in (1.19) whose meaning is constant, regardless of the context of utterance. The word *weather*, for example, means something along the lines of “atmospheric conditions, including temperature, wind, and precipitation” regardless of when or where the word is used. Likewise, although the word *cold* is vague (in the sense that what is cold to one person might not be cold to another), it is consistently used to refer to the low end of some scale of temperature. These meanings, then, are context-independent and semantic. On the other hand, the meaning of *this* depends entirely on the context in which the sentence is uttered. If (1.19) is uttered by the mayor of San Diego on a July afternoon, the noun phrase *this weather* will evoke a very different set of atmospheric conditions than if it is uttered by the mayor of Chicago on a February afternoon.

And similarly, what is meant by *too cold* – that is, how cold is too cold – depends on the speaker and the context. For example, it might be interpreted as *too cold for my tastes*, if for instance it is uttered by someone first stepping outside. But in a different context, it

might be interpreted with respect to some potential activity, if for instance it is uttered by someone who is considering having a picnic outdoors on an autumn day. And in either case, what counts as *too cold* will be relative to the speaker; what is “too cold” for me might be ideal for you. So whether (1.19) means “the temperature in San Diego on July 15 is insufficiently high for Sam Johnson’s third-grade class to feel comfortable holding a picnic” or “the temperature in Chicago on February 15 is insufficiently high for the mayor to feel comfortable walking outside” or any number of other things will depend on who has uttered it and where and when.

In fact, even if you hold these factors constant, the meaning can vary depending on the intent of the speaker: (1.19) could be used by a particular speaker at a particular time and place to mean either “the temperature is insufficiently high for me to feel comfortable outside” or “I need my sweater” or even “please put your arm around me”; and whether the utterance will succeed depends upon the addressee’s powers of inference, as well as the speaker’s ability to correctly predict the addressee’s powers of inference (and the addressee’s ability to infer the speaker’s predictions regarding the addressee’s inferences, etc.). All of these variations in meaning are context-dependent, where the “context” includes not only who uttered the sentence and when and where and to whom, but also the assumed mutual beliefs of the interlocutors; and because these aspects of the meaning are context-dependent, they are pragmatic.

1.3.3.2 Truth conditions

As noted above, the truth conditions of a sentence are the conditions under which it would be true, whereas truth-conditional meaning is any piece of meaning that affects the conditions under which a sentence would be true. Thus, consider (1.20):

(1.20) John is a real genius.

Truth-conditionally, this means that John is extraordinarily intelligent; thus, the sentence is true only under the condition that John is in fact extraordinarily intelligent. If John is actually not at all smart, the sentence is false.

However, notice that such a sentence can be uttered with the intention of achieving quite the opposite effect – that is, to convey that John is not at all intelligent. Suppose John is known to both the speaker and the hearer to be not at all smart. And further suppose that the speaker and hearer have been discussing some particularly foolish comment John made earlier in the day. At this point, the speaker in uttering (1.20) would be taken to mean the opposite of what has been said; in effect, the speaker would convey the belief that John is not smart. Notice, however, that this does not make sentence (1.20) true; rather, it plays off the obvious falsity of (1.20) for ironic effect. (This strategy will be discussed at length in Chapter 2.) The sentence itself remains false in the conditions under discussion, and the hearer’s interpretation of the speaker’s intent relies on the hearer’s belief that it is obvious to both parties that the statement is false. In fact, if the hearer has no reason to believe the speaker takes (1.20) to be false, the hearer will interpret it at face value, and assume that the speaker thinks highly of John’s intellect.

In short, the truth-conditional meaning of (1.20) as a **sentence** is that John is in fact a person of high intelligence; however, the non-truth-conditional meaning of the **utterance**

in the described context is that he is not. Under a truth-conditional view of the semantics/pragmatics boundary, the truth-conditional meaning is semantic, while the non-truth-conditional meaning is pragmatic. And notice that these meanings directly fit the discussion of context-dependent and context-independent meaning above. That is, the aspect of the meaning of (1.20) that is truth-conditional is exactly that aspect that is context-independent; regardless of context, the semantic meaning of (1.20) remains the same. Correspondingly, the non-truth-conditional meaning of (1.20) – here, that John isn’t very bright – is entirely context-dependent; as noted above, if it’s uttered by someone who perhaps thinks highly of John’s intellect, the pragmatic meaning will be entirely different. In this case, therefore, context-independent meaning aligns with truth-conditional meaning, while context-dependent meaning aligns with non-truth-conditional meaning. The former can easily be identified as semantic, and the latter can easily be identified as pragmatic.

1.3.4 *Some boundary phenomena*

Given the amount of overlap between truth-conditional and context-independent meaning, it might seem appealing to collapse the two categories, define semantic meaning as that which is truth-conditional and context-independent and pragmatic meaning as that which is non-truth-conditional and context-dependent, and leave it at that. Indeed, this is precisely what is often implicitly assumed when the difference between semantics and pragmatics is under discussion. Moreover, such an assumption typically leads to another oversimplification: the notion that semantics is done “first” – that is, that in interpreting someone’s utterance, the addressee first determines the context-independent meaning of the sentence and its truth conditions, and then feeds this information into the context to determine the non-truth-conditional meaning that the speaker actually intended in this context. In fact, in discussing (1.20) above, that very assumption was implicit – that the addressee must first determine the context-independent, truth-conditional meaning of the sentence (“John is very bright”), and only then can they notice that this semantic meaning does not appear to be true in the current context (i.e., that the speaker clearly does not believe John to be bright, and neither do I, and the speaker knows this); given these two conflicting facts, the addressee can then go on to calculate the speaker’s intended meaning (“John is not very bright”). Unfortunately, there are significant difficulties with this model. First, the assumed ordering of semantics before pragmatics leads to problems. Second, the overlap between truth-conditional and context-independent meaning is imperfect. Each of the next two sections will discuss a particular example to illustrate these problems.

1.3.4.1 Anaphoric pronouns

One problem with the assumption that semantic, truth-conditional meaning is interpreted first and then fed into the context to yield the pragmatic, non-truth-conditional meaning is that it assumes that there is no context-dependent input into truth conditions. This assumption, however, turns out to be incorrect in the case of pronoun resolution.

A given pronoun may be either used either **anaphorically** or **deictically**. **Anaphora** is the use of a linguistic expression coreferentially with some other linguistic expression used earlier in the discourse (where **coreferential** means “having the same referent”), as in (1.21):

(1.21) My uncle told me that he was a war hero.

Here, *he* is anaphoric, on the reading in which *he* is interpreted as coreferential with *my uncle*. Notice that there is another reading in which *he* is interpreted as someone other than my uncle; for example, if the speaker and hearer have previously been discussing a particular presidential candidate, the uncle may be taken to be referring anaphorically to this individual, who would then be the referent of *he*. On the other hand, if someone has just walked into the room and is highly salient, that person might well be taken to be the referent of *he*; in that case, the use of the pronoun is no longer anaphoric (since it is not coreferential with something in the prior discourse) but rather **deictic** (being interpreted with respect to the context of utterance; see Chapter 4 for a detailed discussion of both types of pronouns).

While both types of pronouns pose similar problems for a “semantics first” theory, let’s focus on the problems posed by the first reading of the pronoun in (1.21). On this reading, the sentence is true if and only if, in the world under discussion, the speaker’s uncle has told the speaker that he, the uncle, was a war hero. The problem, however, is that in order to determine that *he* has the speaker’s uncle as its referent requires access to contextual information – in this case, the earlier part of the sentence. That is, the lexical, conventional, invariant, context-independent meaning of the word *he* says nothing about uncles. You might counter that this can be handled syntactically – that perhaps the structure of the sentence tells us that the two are coreferential, and therefore we can have syntax, rather than pragmatics, provide the referent of *he*. Unfortunately, the problem remains when the two references occur in separate sentences:

(1.22) My uncle was a war hero. He fought in major battles.

Now there is no syntactic connection between *my uncle* and *he*; yet the determination of truth conditions for the second sentence is dependent on our making the connection between the two – a connection that is made pragmatically, due to the salience of the uncle at the time the word *he* is uttered. Undaunted, you might then protest that we could build salience and gender, and even animacy, into the semantic meaning of the word *he*, such that *he* means something like “the most salient animate male in the current context”. In that case, the second sentence in (1.22) would semantically mean “the most salient animate male in the current context fought in major battles”. But that account suffers from two problems: First, it essentially builds the pragmatics into the semantics, muddying the distinction between the two. Second, it’s not supported by our intuitions. Consider (1.23a–b):

- (1.23) a. A: My dad was an officer in the Navy.
 B: Yeah? My uncle was a war hero.
 A: He fought in major battles.
 B: So did my uncle.

- b. A: My dad was an officer in the Navy.
 B: Yeah? My uncle was a war hero.
 A: He may have fought in major battles, but my dad actually saved a guy's life.

In (1.23a), by using *he*, speaker A is not referring to the most recently mentioned, and therefore arguably the most salient, male. Nonetheless, while we might find the conversation a bit awkward, we certainly would not accuse speaker A of saying something false by using *he* to refer to A's own father if the father but not the uncle had fought in major battles. And while one might argue that the uncle is perhaps not the most salient male in the context despite being most recently mentioned, that would leave open the question of why *he* in (1.23b) does seem to be used in reference to the uncle, despite the similarity of the prior contexts in the two discourses. In both cases, the truth or falsity of the utterance depends on the intended referent of *he* and whether that individual fought in major battles, and not on a determination of the referent on the basis of purely semantic and/or syntactic factors. Since the truth conditions of the utterance appear to be dependent on pronoun resolution, which in turn is dependent on the inferred intent of the speaker, which is a pragmatic issue, our model of language comprehension cannot require that truth conditions be resolved prior to pragmatic interpretation.

1.3.4.2 Conventional implicature

A more direct problem for the identification of truth-conditional meaning with context-independent meaning is the fact that there are some aspects of meaning that are context-independent but not truth-conditional. Consider (1.24):

- (1.24) Clover is a labrador retriever, but she's very friendly.

This sentence would be judged to be true just in case Clover is a labrador retriever and is very friendly – that is, in precisely those cases in which the two conjuncts (*Clover is a labrador retriever* and *she's very friendly*) are true. That is to say, it has the exact same truth conditions as (1.25):

- (1.25) Clover is a labrador retriever, and she's very friendly.

The two sentences are not identical in meaning, however. In (1.24), the word *but* conveys a strong sense that the speaker believes (or thinks that the hearer believes) that there is some contrast between being a labrador retriever and being friendly. (In Chapter 2, we will apply the term *conventional implicature* to this aspect of the meaning of *but*.) This contrast is absent in the otherwise parallel (1.25). However, although this contrast is indisputably present, and clearly attached to the use of the word *but*, it arguably does not affect the truth conditions of the sentence. For example, consider a case in which the speaker and hearer both realize that labradors are typically quite friendly. In such a context, (1.24) becomes perhaps an odd thing to say, but one would not want to say that it is false.

Notice, moreover, that this sense of contrast is conventionally attached to the word *but*; it is impossible to use the word *but* without invoking a contrast of some sort between the conjuncts. It is an interesting exercise, in fact, to try to determine precisely how the contrast is interpreted in a range of utterances containing *but*:

- (1.26) a. Clover is a labrador retriever, but she's very friendly.
 b. Clover is very friendly, but she's a labrador retriever.
 c. Clover is very friendly – but (then again) she's a labrador retriever.
 d. Mary's wrong to think labradors are unfriendly; Clover is a labrador retriever, but she's very friendly.

In (1.26a), we get the contrast already noted – that is, that there is some conflict between being a labrador retriever and being friendly, specifically that labradors tend not to be friendly. In (1.26b), there is a subtle difference; here, what is conveyed is that friendly dogs tend not to be labradors. But there's another, possibly even more natural, reading of (1.26b) that corresponds to the reading in (1.26c); interestingly, here there is no contrast at all between being a labrador and being friendly. Instead, the contrast is between Clover being a labrador and some prior sense of surprise regarding her friendliness. Finally, in (1.26d) it may well be that both speaker and hearer are aware that labradors are friendly; the contrast here is between some third party's belief that labradors are unfriendly and the fact of Clover's friendliness.

In all cases, however, the use of *but* conveys contrast, regardless of the context. Thus we can see that the contrast associated with *but* is context-independent yet non-truth-conditional (since, as shown above, the presence or absence of this contrast does not render an otherwise true sentence false). Therefore, context-independent meaning and truth-conditional meaning are not identical. It follows that the dividing line between semantics and pragmatics might in theory be drawn either on the basis of context-dependence or truth-conditional status, but not both. And these are not the only two options. The question of precisely how and where to draw the line will follow us throughout this book. Notice, however, that in a very real sense, there is no right or wrong answer; after all, *semantics* and *pragmatics* are merely lexical items, and like all lexical items, they relate to their meanings in a way that has been arbitrarily but conventionally established. However, because the community using the terms – that is, the community of linguists – is not in complete agreement on their meanings, these meanings are less conventional, less universally shared, than we would like. The question that will follow us throughout the book, then, is not which definition for the term *pragmatics* is correct, but rather which is more helpful.

1.4 Summary

This chapter has presented the fundamental concepts, methodological considerations, and background upon which the remainder of the material in the book will be built. The chapter began with a discussion of basic linguistic principles such as the difference between competence and performance, the difference between prescriptive and descriptive attitudes

toward language, and the rule-governed nature of language. Pragmatics was then situated within the field of linguistics, with a brief description of each of the core areas covered in the study of language structure. Three broad classes of evidence that linguists use in support of their claims – intuitions, experimentation, and naturally occurring data – were discussed, along with the pros and cons of each. The discussion of methodological considerations emphasized the importance in scientific research of formulating a falsifiable claim. From there we went on to a brief consideration of the philosophical background of pragmatics, showing how philosophers' concerns regarding meaning, truth, reference, and context led to the development of the linguistic field of pragmatics.

The remainder of the chapter was devoted to the distinction between the domains of semantics and pragmatics, an issue that will recur throughout this book. Semantics encompasses both lexical and sentential semantics. The discussion of lexical semantics included a range of lexical relations and a comparison of componential semantics and Prototype theory. The discussion of lexical semantics led into a discussion of sentential semantics, including semantic relations at the sentential level that parallel those at the lexical level. The basic principles and notation of formal logic provided a way to formalize the semantic meaning of sentences, as well as a way of thinking about different types of inference. The introduction of concepts such as possible worlds, truth conditions, and truth values laid the groundwork for later discussion of problems concerning the semantics/pragmatics boundary. There was a brief synopsis of propositional calculus and predicate logic, both of which are crucial to the study of semantics and also to the understanding of the need for a field of pragmatics.

The discussion of semantics was followed by a discussion of some basic principles of pragmatics. Distinctions were made between natural and nonnatural meaning, between sense and reference, and between speaker meaning and sentence meaning. The essentially collaborative nature of communication gave rise to the introduction of discourse models as a way of representing sets of possible worlds and keeping track of interlocutors' beliefs (and beliefs about each others' beliefs) concerning these worlds. The concept of mutual belief was examined, with particular focus on the apparent need for an infinite number of beliefs in order to process a discourse, and co-presence heuristics were offered as a way in which interlocutors bridge the gap.

In a theme that will recur throughout the book, a comparison was made between two different possible ways of drawing the boundary between semantics and pragmatics – either on the grounds of context-dependence, with context-independent meaning being semantic and context-dependent meaning being pragmatic, or on the grounds of truth conditions, with truth-conditional meaning being semantic and non-truth-conditional meaning being pragmatic. However, certain boundary phenomena challenge a model that equates these two perspectives and views semantic meaning as truth-conditional, context-independent meaning that serves as the input to contextual considerations and results in additional, non-truth-conditional pragmatic meanings. In particular, the need to resolve anaphora as part of the process of establishing truth conditions challenges the sequentiality of a model that treats semantics as the input to pragmatics, while conventional implicatures such as the contrastive meaning associated with *but* present an instance of meaning that is context-independent yet nonetheless non-truth-conditional.

1.5 Exercises

1. Which of the following claims are falsifiable? For those that are not, is there a way to change them so that they become falsifiable?
 - a. Speakers use the word *please* in order to be polite.
 - b. The word *please* is only used in the context of a request.
 - c. Men and women use language differently.
 - d. Men interrupt more often than women do.
 - e. Women are more concerned with discussing relationships than men are.
 - f. Women spend more time discussing relationships than men do.
 - g. On average, a group of women will spend more time discussing relationships than a group of men will.
 - h. The word *but* signals a contrast between the meanings of the conjuncts.
 - i. The word *and* serves at least three distinct functions in discourse.
 - j. Pauses consistently signal a change in sentence topic.

2. An interesting but little-recognized distinction exists between **homonyms**, which share both sound and spelling (as in *light* “not heavy” and *light* “illuminating device”); **homophones**, which share their sound (“phones”) but may or may not be spelled alike (as in *see* and *sea*); and **homographs**, which share their spelling (“graph”) but may or may not sound alike (as in present-tense *read* and past-tense *read*). Homonyms are, in essence, homophones which are also homographs. In fact, one can correctly say that *homonym* is a hyponym of *homophone*. Explain why.

3. For each pair of words in a–l, determine whether the two words are (a) synonyms, (b) gradable antonyms, (c) complementary antonyms, or (d) none of the above, and tell why. What do your decisions suggest to you about these two categories of antonyms? Can they be improved on, and if so, how?
 - a. uncle/aunt
 - b. sister/sibling
 - c. hide/conceal
 - d. animate/inanimate
 - e. rich/poor
 - f. single/married
 - g. comfortable/uncomfortable
 - h. book/magazine
 - i. magazine/journal
 - j. brother/sister
 - k. true/false
 - l. top/bottom

4. The word *unbuttonable* is ambiguous, in that it can mean either “not able to be buttoned” (as might hold of a jacket that has a zipper instead of buttons) or “able to be unbuttoned” (as might hold of a jacket that does have buttons). Is this a case of lexical ambiguity or structural ambiguity, and why?

5. Give a componential analysis of the words *father*, *mother*, *sister*, *brother*, *son*, and *daughter*.
6. For each of the following, use a truth table to determine whether the formula is analytic or synthetic. If it is analytic, tell whether it is a tautology or a contradiction; if it is synthetic, tell what the world must be like in order for it to be true (i.e., give its truth conditions).
 - a. $(p \wedge q)$
 - b. $(p \wedge q) \rightarrow (q \wedge p)$
 - c. $(p \wedge q) \vee (q \wedge p)$
 - d. $(p \wedge q) \wedge \neg q$
 - e. $(p \wedge q) \vee \neg q$
 - f. $(p \rightarrow q)$
7. Express each of the following in the notation of predicate logic:
 - a. *Everybody loves linguistics.*
 - b. *Either everybody loves linguistics or somebody is crazy.*
 - c. *If Mary is a linguist, John loves her.*
 - d. *If someone's a linguist, they're loved by everybody.*
8. A syntactically unacceptable sentence is said to be **ungrammatical**, and is marked with an asterisk (*), whereas a semantically inappropriate sentence is said to be **anomalous**, and is marked with a question mark (?), and a pragmatically inappropriate utterance is said to be **infelicitous**, and is marked with a crosshatch (#). Mark each of the following to indicate the nature of the unacceptability. Briefly discuss any particularly problematic cases.
 - a. *Yesterday I saw a spider. I chased a spider with the baseball bat.*
 - b. *Spider I saw with bat baseball.*
 - c. *Sarah sped slowly down the stairs.*
 - d. *I saw a gorgeous jacket yesterday; in a store was the jacket.*
 - e. *My dog is the smartest cat I know.*
 - f. *I got 100 percent on the last test; nonetheless, it was a passing grade.*

1.6 Discussion Questions

1. Compare the following claims:
 - a. Discourse markers can be categorized into at least three categories.
 - b. All discourse markers can be categorized into one of three categories.
 Which is the more interesting claim? Why?
2. Discuss the pros and cons of using an internet search engine to construct a corpus, or of using the web itself as a corpus. Attempt to construct a corpus of 50 passive utterances (e.g., *Martha was elected*; *Jeremy was hit by the ball*) with the use of a search

engine. How easy or difficult is it? Why? What types of corpus could you develop straightforwardly with a search engine, and what types would pose problems?

3. Groucho Marx said *Time flies like an arrow; fruit flies like a banana*. Assuming that the first clause is making a statement about the passage of time and the second about the dining habits of certain insects, discuss the contributions of homonymy, polysemy, and syntactic structure to the humor in Groucho's statement.
4. For each of the following sentences, discuss whether the "meaning" in question is natural or nonnatural, and tell why. Do any of these raise difficulties for the distinction?
 - a. *Smoke means fire.*
 - b. *That broken vase means trouble.*
 - c. *The teacher's stern look means trouble.*
 - d. *A red light means stop.*
 - e. *That light on the computer means its battery is charging.*
 - f. *In German, Tisch means "table".*
 - g. *When I yawn, it usually means I'm bored.*
 - h. *When I yawn, it means I want you to finish your drink so we can leave.*
5. Frege's discussion of sense and reference used as an example the phrases *the morning star* and *the evening star*, which have the same referent but different senses. Give examples from your own life of each of the following:
 - a. two expressions with different senses, but the same referent
 - b. a single expression (with a single sense) that can be used for different referents on different occasions
 - c. an expression with sense but no real-world referent
 - d. an expression with reference but no sense (consider the difference between, say, *Washington* and *The Great Salt Lake*)
6. Describe the truth conditions of the following sentence, and give its truth value in at least two different well-known possible worlds, telling what the worlds are and why you chose the truth values you did:

Dorothy threw water onto the Wicked Witch of the West.

7. Consider the following simple sentences, and discuss the challenges they pose for the mentalist and/or referential points of view:
 - a. *Fred is six feet tall.*
 - b. *A unicorn has a single horn.*
 - c. *My brother is an only child.*
8. Having read through the philosophical background of pragmatics, what are your initial thoughts concerning the relationship between truth, reference, and meaning? Do any of these philosophers seem to have it "right"? Note that most of them will be discussed in much greater detail later on, so your opinions might change.