

Logic I – or – Moving in on the Monkey & Bananas Problem

We said that an agent receives percepts from its environment, and performs actions on that environment; and that the action sequence can be based on a plan that is formed by searching (blindly or in an informed way) among all possible such sequences. However, something much richer can go on in this plan-formation stage. The agent can learn (thereby becoming aware of new actions or consequences of actions), and it can reason with what it knows (possibly seeing quickly how to achieve something that might take much longer if done by ordinary searching).

Thus, an agent benefits a lot from having a knowledge base (KB) that it can add to and draw on. [Note: calling this a *knowledge* base should not be taken to mean that everything in it is unquestionably true. But it is to be thought of consisting of items that the agent *takes* to be true, at least for the moment.] And to *draw on* the KB means to access certain items in the KB and use them to guide its decisions, by some sort of reasoning process. That is, the agent will also have an *inference engine* at its disposal.

All this suggests that the agent will be using some form of logic, where the items in the KB are taken as axioms, and the inference process is governed by the logic's rules of inference. This was the view long championed by John McCarthy, in the so-called *logician* approach to AI. And it still is one vigorous approach being pursued today, among others (some of which we will examine later).

We will need to take a look at Propositional Logic and then First-Order Logic, before we can assemble the machinery to solve the Monkey & Bananas problem. Roughly speaking, the monkey will know (have in its KB) basic information about its environment (positions of box, bananas, and itself; actions it can take; etc) and what its goal is (have the bananas); and it will have access to inference rules to allow it to draw conclusions about what to do. The details are rather involved, but the basic idea seems fairly simple.

However, a worry might arise: there are in general *many* ways to start proofs, and most of them won't lead to what we might want to prove. Aren't we then right back where we started, with a horrible search problem? Won't we need a special set of search techniques particular to what we want to prove?

Well, it turns out that there is a clever way to approach this quite generally, and one that we already know: proof by contradiction. Instead of trying to prove a result *R* directly, we assume $\neg R$ and try to prove a contradiction; if we succeed, then we have shown $\neg R$ cannot be correct, hence we have shown *R*. And so all we need is one (powerful) technique for proving contradictions. We will return to this later.

PROPOSITIONAL LOGIC

A logic in general consists of

1. a formal language
2. inference rules
3. axioms
4. semantics (what the symbols mean)

Propositional logic (PL, for short) has a language consisting of letters (A, B, C , etc; so-called sentential variables), connectives ($\rightarrow, \wedge, \vee, \sim$), and parentheses; and the usual ways of combining these into so-called well-formed formulas (or wffs). At times we will want to refer to arbitrary wffs, and we can use Greek or boldface letters ($\mathbf{A}, \mathbf{B}$, etc) to indicate unspecified wffs that may be simple letters or more complex wffs such as $A \wedge \sim C, C \rightarrow (B \vee C)$, etc.

The inference rules for PL can be of various sorts. A very common choice is simply that of *modus ponens* (MP): from $\mathbf{A} \rightarrow \mathbf{B}$ and $\mathbf{A}$, infer $\mathbf{B}$. It is often enough all by itself; what we mean by *enough* will become clear later.

As for axioms, a simple choice is to allow all tautologies: wffs that are always true – but to say what that means will require us to look at semantics, so we turn to that now.

Semantics for PL is that of the familiar *truth tables*. We simply list all the letters of interest, and under them we create rows with all possible arrangements of Ts and Fs. Each such row is called an *interpretation* of PL. It does not really go so far as to state meanings for the letters; it simply marks each as true or false. And then we extend the listing to include whatever wffs we are interested in (that employ those letters), and we also mark down the proper Ts and Fs for them in each row, based on the usual understandings. For instance, the truth table for $\mathbf{A} \rightarrow \mathbf{B}$ is

$\mathbf{A}$	$\mathbf{B}$	$\mathbf{A} \rightarrow \mathbf{B}$
T	T	T
T	F	F
F	T	T
F	F	T

I find it easy to remember this as follows: consider $\rightarrow$ to indicate a promise: if you do what is on the left-hand side, then I promise to do what is on the right-hand side. Now, how can it happen that the promise turns out to be false (that is, how can I break the promise)? You will have to do your part ($\mathbf{A}$) and I will have to refrain from doing my part ($\mathbf{B}$). That is, $\mathbf{A} \rightarrow \mathbf{B}$ is false iff $\mathbf{A}$ is true and $\mathbf{B}$ is false, which happens only in row 3.

Now we can say precisely what a tautology is: any wff whose truth-table assigns it the value True in *every* row. And a *contradiction* is a wff that is assigned False in every row. Finally, a *contingent* wff is one that is neither a tautology or a contradiction: it is assigned both True and False (in different rows, of course).

So, the tautologies are pretty special. But they can take a great many forms, and there is an infinite number of them. So it is a bit awkward taking all of them as axioms, when we don't even know what they all look like. A more streamlined set was proposed by Lukasiewicz:

1. $(\sim A \rightarrow A) \rightarrow A$
2. $A \rightarrow (\sim A \rightarrow B)$
3. $(A \rightarrow B) \rightarrow ((B \rightarrow C) \rightarrow (A \rightarrow C))$

The three forms above are *schemata*: they each represent an infinite number of possible wffs, gotten by replacing the boldface letters by any actual wffs we like. Each of the three forms however always will be a tautology (this is easily verified by truth-tables), so the Lukasiewicz schemata (let's call them L) are a subset of all tautologies.

It turns out that the rule MP together with axioms L lead to exactly the same set of consequences (things that can be proven) as does MP with all tautologies as axioms. However, it is still cumbersome that L represents an infinite number of axioms. We will see a way around that a bit later on.

What does it mean to *prove* something, say the wff **W**, in PL? It means: to start with whatever one is using as axioms (which need not consist of tautologies, by the way – just whatever one wants to take to be true for whatever reason), and to apply one's rule(s) of inference over and over until one arrives at the result that is wanted. We write

$Ax \vdash W$

where Ax is our chosen set of axioms (we perhaps could call it KB instead); the inference rules are understood to be whatever we have specified. This notion of proof is formal, or syntactic: it does not make any mention of truth or meaning. We sometimes read the above as “ Ax proves **W**”, which really means: there is a proof of **W** from Ax (using the understood rules of inference, eg MP).

There is another totally different way to try to “conclude” a wff: write out the truth-table for all the axiom wffs as well as the wff **W** that one wants to prove. Then check all those rows in which the axioms are all true, to see if **W** holds in them. If so, we say **W** is *entailed* by the axiom set (Ax , say), and we write

$Ax \models W$

This means, then, that **W** is true in all interpretations in which all wffs of Ax are true. It is a semantic notion, making use of truth-values, but does not use inference rules at all.

So we have two extremely different notions of “arriving at” a wff as a conclusion, starting from initial (axiom) wffs. One way is much like the usual mathematical notion of proof, where we reason step-by-step, starting from axioms and using inference rules as we go along, trying to be clever enough to choose the right axioms and the right rule(s) to get to the result we want; and then it looks much like a search problem, in a huge tree of possible proofs-sequences.

And the other method is much more routine, possibly tedious, but very methodical: just fill out the (possibly enormous) truth-table and check to see if **W** is true in every row where all the axioms are true. There is no thinking needed here, no cleverness.

How do they match up? Well, with the right combination of axioms and inference rules, they match up perfectly. For instance, we have this famous result:

Theorem (Soundness and Completeness of PL): Consider any version of PL that includes the Lukasiewicz schemata among its axiom set Ax , and MP as one of its inference rules. Then for every wff **W**, $Ax \models W$ if and only if $Ax \vdash W$.

The left-to-right direction is *completeness* (the proof method is able to prove everything it *should*), and the right-to-left is *soundness* (everything provable is true).

A digression on speed/efficiency: Even with the Lukasiewicz axioms – indeed even with the much more streamlined Robinson resolution method that we will see shortly – it still usually takes a long time to find proofs in PL. There is a famous problem, called *SAT*: the satisfiability problem for PL. It is this: given a wff **W**, is it satisfiable (is it true in at least one interpretation) or not? We can settle the matter by drawing the entire truth table for **W**, of course, but that is very time-consuming, requiring in general the filling out of as many as 2^n rows, where n is the number of letters in **W**, and where each row has at least $n+1$ entries (one for each letter plus one for **W**). We say that this solves SAT, but it is an exponential time and space solution.

SAT is in the class of problems known as *NP*-problems. This is the class of problems that, when a supposed solution is suggested, it is easy (i.e., can be done in time proportional to a polynomial in the “size” n of the problem) to check if it really is a solution. In fact, SAT was the first problem shown to be *NP-hard*: any solution to SAT can be transformed into one that solves any other NP problem in essentially the same amount of time. So SAT is maximally hard among NP problems. No one has

ever found a fast (polynomial-time) solution to SAT; and if anyone does, then that will mean *all* NP problems can be solved in polynomial time.

What makes SAT so special like this? Perhaps because, being based on a basic logic framework, it can be used to encode all the other NP problems. This idea was exploited by Stephen Cook in 1971, when he first proved that SAT is both NP-hard and a member of NP (together these say SAT is NP-complete).

One final note on PL proofs: to prove **W** from a given axioms set, KB, we can suppose $\sim\mathbf{W}$ (enlarge our axiom set to $\mathbf{KB}' = \mathbf{KB} + \mathbf{W}$) and try to show that $\mathbf{KB}'$ is inconsistent (contradictory). This we mentioned before: proof by contradiction, also called indirect proof. In particular it is typically used along with Robinson's resolution method, and also is basic to the PROLOG programming language. We will study both of these soon.