

5HC : Chinese boxes or tea chests?

The Urth list — 23 messages, 8 voices, 01 Feb 2005

<https://urth.darkrealm.vip/thread/urth/4348>

Exported from an unofficial mirror of a public mailing list. Copyright in each message remains with the person who wrote it. Email addresses are obscured as `user at host`, the form the original archive published. This file was rebuilt from parsed fields rather than retained headers, so it is faithful in content but not byte-exact.

Jim Raylor — 01 Feb 2005, 07:47

Searle's argument was that the guy in the room who manipulated the symbols had no real knowledge of the Chinese language, just the rules for manipulating the symbols.

It could be argued that all language users 'merely' know how to manipulate the symbols without a full awareness of their meanings.

I think this when I use words like 'Notwithstanding' and 'Nevertheless'.

I sometimes wonder if interpreters at the UN (for example) pay any heed to the meaning of the messages they translate, beyond what is required to translate the words.

For example an eccentric translator could perversely translate a book by cutting it into individual sentences, translating each one in random order and then reassembling them in the original order.

I seem to recall a program which regularly fooled people into thinking it was a client-centred psychotherapist.

How do you feel about a program which regularly fooled people into thinking it was a client-centred psychotherapist?

=====

<http://prisoner143.50megs.com/>

I am not a number I am a free person!

David Lebling — 01 Feb 2005, 08:06 · in reply to Jim Raylor

From: "Jim Raylor" <rjaylor at yahoo.co.uk>

How do you feel about a program which regularly fooled people into thinking it was a client-centred psychotherapist?

Joe Weizenbaum wept.

Dan'l Danehy-Oakes — 01 Feb 2005, 11:43 · in reply to Jim Raylor

On Tue, 1 Feb 2005 15:47:12 +0000 (GMT), Jim Raylor <rjaylor at yahoo.co.uk> wrote:
Searle's argument was that the guy in the room who

manipulated the symbols had no real knowledge of the Chinese language, just the rules for manipulating the symbols.

This is a significant point, because the neurons in my brain have no real knowledge of the English language, just the rules for manipulating the symbols. Thus, the neurons do not "know English." But _I_ appear to.

Similarly, the man in the Chinese Room does not Chinese, but the system-as-a-whole appears to. To say "The Room does not know Chinese," Searle must, _at a minimum_,

1. Define what it means to "know Chinese,"
2. Propose an empirical test for the definition proposed in #1, and
3. Demonstrate (from hypothesis, since no such Room exists) that the Room does not fit the definition proposed in #1.

Searle has not, to the best of my knowledge, proposed either a definition of "knowing Chinese" or an empirical test for determining whether the Room actually knows Chinese. He has instead relied on his readers' paralogical emotional response to assure us that it does not. But that does not mean that it does not.

I seem to recall a program which regularly fooled people into thinking it was a client-centred psychotherapist.

That would be "Eliza." It didn't really fool very many people or for very long.

--Dan'l

"We're going to sit on Scorsese's head"
-- The Goodfeathers

James Wynn — 01 Feb 2005, 12:43 · in reply to Dan'l Danehy-Oakes

Searle has not, to the best of my knowledge, proposed either a definition of "knowing Chinese" or an empirical test for determining whether the Room actually knows Chinese. He has instead relied on his readers' paralogical emotional response to assure us that it does not. But that does not mean that it does not.

Yes, but to be fair Turing's criteria is just as subjective. Roughly, "can a computer succeed in impersonating a human person well enough to fool humans." Where is the "control" here? Mr. Raylor has brought up an excellent point: What if a computer can succeed in fooling someone that it is something we KNOW it is not? That would be a useful control, I think. Is successful mimicry really the final criteria for personhood for a machine?

From bird's point of view, many bugs succeed in fooling them that they are thorns, leaves and sticks. To the mother bird, the coo-coo is her child.

Another possible control that someone has pointed out is that a computer conceivably might be conscious and might think and NOT seem at all human in its responses. Such a computer would not test positive for the Turing test. I used to subscribe to AI Magazine but I'm not much of an authority on this. How does the Turing criteria eliminate for an especially good mimic (to bring this discussion back to V.R.T)?

Wolfe, cleverly, has upped the ante on the Turing test by making the "simulation" much more elaborate:

1) Mr. Million (referred to a "simulation" as a deliberate foil to V.R.T's "simulation") and Rose and the Whorl gods who are much the same thing in another form

2) The Lemur clan and Potto (are the chem bodies whose bodies they have commandeered walkie-talkies as they apparently intended or are they only simulating the psyches of the Ayuntamiento?)

3) What does it mean that the souls of the inhumani look very similar to the physical bodies of their hosts?

~ Crush

Maru Dubshinki — 01 Feb 2005, 12:53 · in reply to Dan'l Danehy-Oakes

The two of you are probably thinking of a) Parry; an improved Eliza which was a lot more effective and b) some experiments where unsuspecting students talked to a program similar to Eliza via a terminal and never suspected anything (that is, they were not told to apply all their critical spirit to it, like they would in a real Turing test.).

~Maru

Microsoft delenda est.

On Tue, 1 Feb 2005 11:43:20 -0800, Dan'l Danehy-Oakes <danldo at gmail.com> wrote:
I seem to recall a program which regularly fooled
people into thinking it was a client-centred
psychotherapist.

That would be "Eliza." It didn't really fool very many
people or for very long.

--Dan'l

--

"We're going to sit on Scorsese's head"

-- The Goodfeathers

Maru Dubshinki — 01 Feb 2005, 12:55 · in reply to Jim Raylor

Of course they do: context matters. When translating it's not a matter of just swapping syntactic structures, grammar, and vocabulary- if it was, machine translation woulda been licked a long time ago. Those are low level rules- there are rules that govern how the story acts, and how the meaning units interact. Contextual rules, not just syntactic rules, if that makes sense.

~Maru

Microsoft delenda est.

On Tue, 1 Feb 2005 15:47:12 +0000 (GMT), Jim Raylor
<rjraylor at yahoo.co.uk> wrote:
I sometimes wonder if interpreters at the UN (for
example) pay any heed to the meaning of the messages
they translate, beyond what is required to translate

the words.

For example an eccentric translator could perversely translate a book by cutting it into individual sentences, translating each one in random order and then reassembling them in the original order.

I seem to recall a program which regularly fooled people into thinking it was a client-centred psychotherapist.

How do you feel about a program which regularly fooled people into thinking it was a client-centred psychotherapist?

=====

<http://prisoner143.50megs.com/>

I am not a number I am a free person!

ALL-NEW Yahoo! Messenger - all new features - even more fun! <http://uk.messenger.yahoo.com>

David Lebling — 01 Feb 2005, 14:05 · in reply to Dan'I Danehy-Oakes

From: "Dan'I Danehy-Oakes" <danldo at gmail.com>

That would be "Eliza." It didn't really fool very many people or for very long.

Indeed it was Eliza, written by Joseph Weizenbaum. For a long discussion of the reaction to the program, see his book "Computer Power and Human Reason," a passionate critique of "strong AI."

Eliza was widely distributed as "Doctor," which imitated a Rogerian psychotherapist. Parry imitated (in my opinion in an even cruder way than Eliza's effort) a paranoid schizophrenic. There is a classic dialog between the two which is easy to Google for (try "parry" and "doctor").

It does not take long to see that either is not human, or able to understand English in even a very crude way, but many people did not, for various reasons, most generally from the social surround of how the program was presented. If the source of the answers is described as a person, many people will believe that it is in fact a person, until the evidence to the contrary is overwhelming.

What lessons about 5HC we may draw from this experience I hesitate to guess.

-- Dave Lebling, aka vizcacha

James Wynn — 01 Feb 2005, 14:44 · in reply to James Wynn

Yes, but to be fair Turing's criteria is just as subjective.

Ummm, no. It is subjective - or, rather, it is based on subjective judgments - but it provides a definition (admittedly, an ostensive rather than a descriptive definition), and a way of testing for whether that definition has been met. Searle (to the best of my memory) does neither.

But since a camouflage device can always get better and better (particularly since our A.I.'s creators can experiment to discover the chinks in its disguise) it is a standard that can likely be met without ever addressing whether the definition is valid (just as the mother bird's criteria for whether the coo-coo is her child is that it was hatched in her nest. Searle's analogy begs all sorts of questions, but it is a set up so that the definitions are implicit: the person in the box represents the expert system A.I. and it "acts" without "thinking". Remember, the original question that Turing was addressing was "can a computer 'think'?".

Roughly, "can a computer succeed in impersonating a human person well enough to fool humans." Where is the "control" here?

Read Turing's article. The "imitation game" was to be played by two players, a human and a computer, with a second human judging between them - "which is the human and which is the computer?" The test is, of course, blind - the judge doesn't know which player is which. (In a really rigorous test, I would suggest that runs be done with all possible combinations (not only human vs computer but human vs human and computer vs computer.) Both players are trying to convince the judge that they are the human. Thus, the responses of the human player provides the control to which the responses of the computer player are compared.

That comes out to little more than an opinion poll if one assume that humans are generally equal in their ability to be duped. Not much of a "control". On the other hand, even if a computer were conscious, a computer programmer who designed such systems would probably always be able to identify certain algorithms to recognize it.

~ Crush

Chris — 01 Feb 2005, 16:42 · in reply to Jim Raylor

It could be argued that all language users 'merely' know how to manipulate the symbols without a full awareness of their meanings.

Well, kinda. If you did make an argument that all language users 'merely' know how to manipulate symbols, then you would need a matching conception of "meaning". (Such theories exist). In such an argument, there's simply isn't any "deeper meaning" for language users to lack awareness of.

I sometimes wonder if interpreters at the UN (for example) pay any heed to the meaning of the messages they translate, beyond what is required to translate the words.

If they didn't, they wouldn't be useful translators. Ambiguities don't map one-to-one in different language; in order to know which French word to use in a translation of an ambiguous English word, you have to work out the probable meaning of the expression it occurs in first.

Also the vast majority of typical communication is also not quite literal. Most of the time you intend to express something more than, or something different from, the literal meanings of the words you use. Preserving this in a word-for-word translation, even without ambiguities, would be unlikely.

Translation is hard.

I seem to recall a program which regularly fooled people into thinking it was a client-centred psychotherapist.

How do you feel about a program which regularly fooled people into thinking it was a client-centred psychotherapist?

I think this was Dr. Sbaitso, although Dr. Sbaitso might well just be a variant of Eliza (Dan'll's suggestion).

Civet

Chris — 01 Feb 2005, 17:33 · in reply to James Wynn

Dan'll said:

Searle has not, to the best of my knowledge, proposed either a definition of "knowing Chinese" or an empirical test for determining whether the Room actually knows Chinese. He has instead relied on his readers' paralogical emotional response to assure us that it does not. But that does not mean that it does not.

Well, hold up for a second. Empirical tests aren't necessary, because this is a thought problem, and in thought problems you can actually **stipulate** such things; of course, this means that what you end up with is a conclusion that a hypothesis is either logically consistent or inconsistent, it doesn't mean that an actual instance of what you describe in the thought problem can or will happen.

I think it was Searle's intention to stipulate the layout of the Room such that not only does it not "know Chinese", but that it **does appear** to "know Chinese". Unfortunately the layout he describes doesn't really fit the bill, for one thing. But really there are a multitude of problems here. Granting that he does have a definition of what it is to understand a language (presumably laid out somewhere in the large body of work he's done in the philosophy of language), for example, doesn't necessarily compel you to agree with that definition. And so on and so on and so on.

Crush said:

Yes, but to be fair Turing's criteria is just as subjective. Roughly, "can a computer succeed in impersonating a human person well enough to fool humans." Where is the "control" here? Mr. Raylor has brought up an excellent point: What if a computer can succeed in fooling someone that it is something we KNOW it is not? That would be a useful control, I think.

And I think this was Searle's intention in trying to lay out the Chinese Room the way he did: to raise this as a logical possibility. Personally I don't think he succeeded with his example, but even if he did succeed, what are the implications? Well, to start with it would mean that obviously you can't **know** that something is intelligent just because it behaves "intelligently". And if this is the case then a strict behaviorist would have a problem. (It is unclear to me but I get the general impression that Searle thinks that Turing's position is necessarily behaviorist, and I don't think it really is.)

Is successful mimicry really the final criteria for personhood for a machine?

This is something I have been trying to get at here. There are two questions to be answered, an "is" question and an "ought" question. If Searle is correct the "is" question (is the machine intelligent?) can't be answered by examination of behavior alone (or possibly can't be answered

at all). However, "personhood" is really an "ought" question. How ought a machine that appears to behave intelligently be treated, what rights should it have, etc?

I would say that Turing is really answering this "ought" question: a machine that appears to be conscious should be granted personhood even if you're unsure of the truth of whether it "is" intelligent. And I think it is even more clear that this is the position that Wolfe holds with the example of Rose.

Maru Dubshinki — 01 Feb 2005, 17:40 · in reply to Chris

I think what he said is even more fundamental than that-(stop me if I've said this before). I think what Turing is saying, is that if you can't tell the difference in output, can't make statistically meaningful differentiation, then you have no choice but to grant it personhood, because you will **never** find a better method, a better test.

~Maru

Now here's an interesting thought problem- postulate a vast group of robots who can pass the Turing Test, and mimic a bright, inventive young engineer so well nobody can tell the difference through a screen, and who are also reasonably physically handy. Suppose every real 'human' died off and all that is left are robots w/o "real consciousness". What would happen if they discovered an asteroid coming, and the only possible way of survival was to blow it up with a nuke? Suppose further that they made like a human, solving all technical problems, and an advanced alien civilization came by impressed by what they had done. Would the aliens accord the robots 'personhood' and would they be right or wrong?

On Wed, 02 Feb 2005 01:32:45 +0000, Chris <rasputin_ at hotmail.com> wrote:

I would say that Turing is really answering this "ought" question: a machine that appears to be conscious should be granted personhood even if you're unsure of the truth of whether it "is" intelligent. And I think it is even more clear that this is the position that Wolfe holds with the example of Rose.

Chris — 01 Feb 2005, 23:09 · in reply to Maru Dubshinki

Maru said:

Suppose further that they made like a human, solving all technical problems, and an advanced alien civilization came by impressed by what they had done. Would the aliens accord the robots 'personhood' and would they be right or wrong?

You've stipulated they don't have "real consciousness", but the important unknown is: do the aliens have the ability to tell the difference, or if they are uncertain, do they have an estimate of probability for the robots having "real consciousness"?

For fun, a few possibilities:

Kantian deontological Martians - treat the robots as persons if there is even the tiniest finite chance of them being conscious, because any attempt to universally apply a maxim which treats a being with a finite chance of being a person as a non-person will eventually use at least one person as a means and thus violates the Categorical Imperative.

Utilitarian Neptunians - attempt to calculate the relative utilities of treating the robots as persons vs. taking advantage of them, weighted for relative probability that each is the case. If unable to accurately estimate the utilities, they spin around in a circle three times chanting "Do

the done thing" and then ask the Martians what to do.

Hobbesian Plutonians - don't care whether the robots are persons or not. They treat the robots as persons because cooperation is to their advantage.

Venusian virtue-ethicists - treat the robots as persons even if they know they're not conscious, because it reinforces their virtue of humane compassion.

Mercurial relativists - Well, do the robots consider *themselves* persons? You have to respect the robots' position on the matter. Unless, of course, the current culture on Mercury has developed anti-robot values, in which case who could blame them for destroying or enslaving the robots? Or, alternately, the more class-conscious Marxists among them elevate them as the ultimate expression of the proletariat and elect the robots to be their leaders. But they might change their minds later. Or do all of the above, because after all there are no wrong answers.

Nathan Spears — 02 Feb 2005, 06:47 · in reply to Chris

I'm not sure I get the "have to nuke the asteroid" bit. If the computers ran off solar energy than it would be a legitimate problem, I suppose, although not for the reasons they were aware of. But don't you think that an entity which has consciousness needs to be "self-aware" enough to make real decisions concerning its own future? Or were you saying that the computers destroyed the asteroid out of a legitimate sense of self-preservation?

Chris, I would have gone with Platonic Plutonians: do the computers aspire to the Form of Personhood?

Dan'l Danehy-Oakes — 02 Feb 2005, 09:39 · in reply to Nathan Spears

Just as a side note:

Differentiate between behaviorism[1] and behaviorism[2]. (Sorry, my email doesn't do subscripts.)

Behaviorism[1] refers to the epistemological position that, since we do not have access to the content of the consciousness of a subject of study (even if the subject describes the content of her consciousness, we have to allow for inaccuracy, lies, etc.), the only aspect of the subject we can actually study is her behavior (including statements about the content of her consciousness, etc.).

Behaviorism[2] refers to the ontological position that no such thing as "consciousness" exists.

Behaviorism[1] seems to-me a legitimate and scientific position.

Behaviorism[2] seems to-me to have been concocted as a parody of behaviorism[1].

However, some extreme behaviorists[1] often act and speak as if they actually believed in behaviorism[2].

The point of all this? Turing's test works in the scientific realm of behaviorism[1], not the essentially religious (because incapable of empirical test) realm of behaviorism[2].

Searle's Chinese Room gedankenexperiment works in the space between behaviorism[1] and behaviorism[2].

What I mean by this:

Turing proposes an open test and says, "If we obtained a certain type of results, the most reasonable conclusion would be that the machine in question was conscious." He leaves the question open of whether the conclusion would actually be true, or whether such a result might actually occur.

Searle proposes a much more limited test and says, "No matter what the results, we must conclude that the Room does not 'know Chinese,' because we know *a priori* that it does not."

--Dan'l

"We're going to sit on Scorsese's head"
-- The Goodfeathers

Maru Dubshinki — 02 Feb 2005, 10:29 · in reply to Nathan Spears

What I meant was something along these lines: ' Suppose further that the robots ran into some difficult, advanced, hard-to-solve life-threatening problem. Let's call it X. Then the robots solve X just like the humans would have, if they were still around (because the robots are so very good at imitating what humans would do). Then some advanced aliens come...' etc.

~Maru

Microsoft delenda est.

Nathan Spears <spearofsolomon at yahoo.com> wrote:
I'm not sure I get the "have to nuke the asteroid" bit. If the computers ran off solar energy then it would be a legitimate problem, I suppose, although not for the reasons they were aware of. But don't you think that an entity which has consciousness needs to be "self-aware" enough to make real decisions concerning its own future? Or were you saying that the computers destroyed the asteroid out of a legitimate sense of self-preservation?

Chris, I would have gone with Platonic Plutonians: do the computers aspire to the Form of Personhood?

Maru Dubshinki — 02 Feb 2005, 10:35 · in reply to Chris

Ahh, now that depends on your philosophical position. If you go with Penrose and suchlike, then all you have to do is see if they pass the Turing Test and also run on the right substrate. But if you take the Strong AI position, and argue that intelligence and consciousness are independent of the physical instantiation, that what is relevant is the pattern, then they have to fall back on observations, to what are essentially variations on the Turing Test. I'll leave out how every thought about reality is but a probability, since it isn't really relevant. But you left out the Bayesian Beta Regulisians: Examine carefully their actions, material culture, and the results of various tests you carry out. Weight them all into the formula, obtain your best probability, then evaluate expected utility and go from there.

~Maru

Microsoft delenda est.

Chris <rasputin_ at hotmail.com> wrote:
Maru said:

Suppose further that they made like a human, solving all technical problems, and an advanced alien civilization came by

impressed by what they had done. Would the aliens accord the robots 'personhood' and would they be right or wrong?

You've stipulated they don't have "real consciousness", but the important unknown is: do the aliens have the ability to tell the difference, or if they are uncertain, do they have an estimate of probability for the robots having "real consciousness"?

Nathan Spears — 02 Feb 2005, 10:38 · in reply to Maru Dubshinki

But are the similarities superficial? That is, are they solving problems for themselves using human methods, or for the same reasons humans would have? Do they believe that they are organic creatures dependent on oxygen and sunlight for survival, or are they "aware" (even if that only means, "they factor into their problem-solving the fact that they are machines") that they are machines mimicking human thought-processes?

--- Maru Dubshinki <marudubshinki at gmail.com> wrote:
What I meant was something along these lines: ' Suppose further that the robots ran into some difficult, advanced, hard-to-solve life-threatening problem. Let's call it X. Then the robots solve X just like the humans would have, if they were still around (because the robots are so very good at imitating what humans would do). Then some advanced aliens come...' etc.

~Maru
Microsoft delenda est.

Nathan Spears <spearofsolomon at yahoo.com> wrote:

I'm not sure I get the "have to nuke the asteroid" bit. If the computers ran off solar energy then it would be a legitimate problem, I suppose, although not for the reasons they were aware of. But don't you think that an entity which has consciousness needs to be "self-aware" enough to make real decisions concerning its own future? Or were you saying that the computers destroyed the asteroid out of a legitimate sense of self-preservation?

Chris, I would have gone with Platonic Plutonians: do the computers aspire to the Form of Personhood?

James Wynn — 02 Feb 2005, 11:14 · in reply to Dan'l Danehy-Oakes

Behaviorism[2] refers to the ontological position that no such thing as "consciousness" exists.
Behaviorism[1] seems to-me a legitimate and scientific position.
Behaviorism[2] seems to-me to have been concocted as a parody of behaviorism[1].

I don't know the whole history of Behaviorism[2], but wasn't it most popularized in B.F. Skinner's "Beyond Freedom & Dignity"? He is hardly a fringe personality in the science of Behaviorism. Skinner argued that problems of human overpopulation could be handled by the emerging 'technologies of behavior' which in order to be fully developed had to be sundered from the notion of human "free will", that an "automaton" resides within us (as opposed to the

understanding that we are the sum of Nature and Nurture). There are lots of critics of Skinner (count me with them), but none that I know of do so with empirical data.

Searle proposes a much more limited test and says, "No matter what the results, we must conclude that the Room does not 'know Chinese,' because we know _a priori_ that it does not."

I agree with this. The Searle's conclusion (and definitions) are implicit in his thought experiment. But the argument against him always comes down to: "Well, how do you know that *humans* 'think' in the superlative manner Searle's analogy seems to presume they do?" But that sounds like a tilt toward B.F. Skinner to me.

I think Wolfe has posed a much more complex scenario. Rather than building his simulated personalities from scratch, he has them imprinted directly from a human pattern and says "Are they people?" Lemur says "no". He calls them 'ghosts'. Sort of moving, monuments to people long dead. (And then he found out that HE was such a monument as well). Mr. Million seems to agree with Lemur.

Rose say "yes", because she already WAS a person. She's not a fraud, mimicking personality.

V.R.T. is trying to *become* Marsch but does not seem to have succeeded because Number Five and his jailers seem to see right through him....but if he got away with it, would he BE Marsch? Does this question not strike to the heart of the Turing "imitation game"?

~ Crush

Maru Dubshinki — 02 Feb 2005, 11:23 · in reply to James Wynn

There's a relevant idea in computer science which goes 'a perfect emulation *is* the thing being emulated.'

~Maru

On Wed, 2 Feb 2005 14:14:25 -0500, James Wynn <thewynns at earthlink.net> wrote:

V.R.T. is trying to *become* Marsch but does not seem to have succeeded because Number Five and his jailers seem to see right through him....but if he got away with it, would he BE Marsch? Does this question not strike to the heart of the Turing "imitation game"?

~ Crush

Dan'I Danehy-Oakes — 02 Feb 2005, 11:29 · in reply to James Wynn

First, I do believe that we are reaching the (or a) heart of 5HC in a way I've never seen before. I wonder if anyone here can tie "A Story" into this analysis?

Second...

Yes, Skinner is _the_ paradigmatic behaviorist. Behaviorism[2] is very much pointed at the extreme behaviorist[1] position he takes in books like _Beyond Freedom and Dignity_ and, to a lesser extent, _Walden Two_.

The key point is that Skinner, as near as I can make out, does not claim that "consciousness" does not happen - though he _does_ claim, if I'm not mistaken (it's been a looong time since I

read S) that there is no such "thing" - no nounish extant - as "consciousness." Rather, he is claiming that "consciousness" is irrelevant to pragmatic matters (such as controlling the "population explosion"), which _can_ (he says) be handled by "technologies of behavior" that don't take "consciousness" into account.

In this, he seems to me to fall into the same overall category as Hitler, Stalin, Mao, Pol Pot, etc. - an anti-humanist. I also suggest that this is a possible theme of 5HC, and particularly of "VRT".

--Dan'l

"We're going to sit on Scorsese's head"
-- The Goodfeathers

Iorwerth Thomas — 03 Feb 2005, 04:29 · in reply to Dan'l Danehy-Oakes

From: Dan'l Danehy-Oakes <danlido at gmail.com>

Behaviorism[1] seems to-me a legitimate and scientific position.

Behaviorism[2] seems to-me to have been concocted as a parody of behaviorism[1].

However, some extreme behaviorists[1] often act and speak as if they actually believed in behaviorism[2].

Hmm, I think (going by conversations with a psychologist flatmate) that [1] is historically prior to [2]... That may be because one gets so used to a stance taken for the sake of scientific method that one starts to take it as reality (add in logical postivism, and [1] pretty much becomes equivalent to [2], since the only things we can talk about are scientifically provable things, if you're a logical positivist). Which do you class Skinner as?

I keep coming across articles in the New Scientist that seem to be attempting to argue from [1] to [2]. That depresses me.

Iorwerth

Iorwerth Thomas — 03 Feb 2005, 04:39 · in reply to James Wynn

From: "James Wynn" <thewynns at earthlink.net>

I don't know the whole history of Behaviorism[2], but wasn't it most popularized in B.F. Skinner's "Beyond Freedom & Dignity"? He is hardly a fringe personality in the science of Behaviorialism. Skinner argued that problems of human overpopulation could be handled by the emerging 'technologies of behavior' which in order to be fully developed had to be sundered from the notion of human "free will", that an "automaton" resides within us (as opposed to the understanding that we are the sum of Nature and

Nurture). There are lots of critics of Skinner (count me with them), but none that I know of do so with empirical data.

From conversations with a psychologist friend, I get the impression that he vacillated between the two (he was an epiphenomenonologist (sp?), so consciousness - when he believed in it - didn't matter anyway). He apparently subjected his daughter to operant conditioning, but that may be exaggeration - my friend doesn't like him overmuch.

Iorwerth

Chris — 03 Feb 2005, 11:21 · in reply to Dan'l Danehy-Oakes

At the risk of terrible boredom, this is a reasonable reference:

<http://plato.stanford.edu/entries/behaviorism/>

Analytical behaviorism seems to go just as far back as the other kind. But what I would point out here is that the mere epistemic position - that behavior provides the only data we have to go by - is not enough to constitute any of these types of behaviorism, and I don't think that what is being referred to here as "behaviorism[1]" is really behaviorism.

To draw a parallel with physics: when studying the motion of a ball in space, the only data we have to go by is its position in time. Two scientists could entirely agree on this much. But let's say that one of these scientists concludes that the proper way to study the motion of the ball is by mathematically analyzing its motion, whereas the other insists that the object of inquiry should not be the mere motion of the ball itself but rather the action of the forces which act on the ball, and the laws of motion in general.

This is not meant to be a perfect analogy. But my point is that being a behaviorist of **any** stripe requires more than just a particular view of behavior as data. It requires a certain outlook as to the limits, method, and object of study. And to accept Turing's position I don't think you have to accept any variety of behaviorism.

As a side note I think that the term has tended to refer to the more radical versions as time has gone on at least in part because people took the general principles of the most extreme version and imported it into other arenas under the name "behaviorism". For example "linguistic behaviorism", something Searle was concerned with, is the notion that there is no deeper "meaning" than simply the set of external stimuli that would lead one to affirm/assert or deny a sentence. [I can only say that his arguments against this are much sharper, in my opinion, than the whole Chinese Room thing.] In any event anyone familiar with these other arenas rather than psychology is going to get a skewed impression of what "behaviorism" means.