

# Reporting *Fears*

*Frank P. DeVita*

August 2026

## A Puzzle About Fear and Belief

What does it mean to say that an agent *fears*? What are the conditions that make fear reports true? Is uncertainty necessary? Are brain states sufficient? These questions are partly linguistic, partly philosophical, and partly scientific. In this paper, I get clearer on the semantics of the intensional verb *fears* by reviewing discussions of it in analytic philosophy with particular focus on *Prior's puzzle*—a Frege puzzle-like substitution failure of co-referential that-clauses and noun phrases that occurs under *fears* and some other intensional verbs, but not under *believes*. The main challenge posed by Prior's puzzle is to explain why sentences like "Sally fears that Fido barks" mean something different than sentences like "Sally fears the proposition that Fido barks" when "that Fido barks" and "the proposition that Fido barks" co-refer, and further, why this substitution failure doesn't occur with *believes*.

Extant solutions to Prior's puzzle emphasize linguistic issues pertaining to polysemy and intensional verbs, the meaning of that-clauses, and the predication of content to nouns. Here, I propose a solution to Prior's puzzle that engages with the linguistic issues to shed some light on broader questions about the attitudes and intentionality. My main argument is that *fears* and *believes* relate agents to different kinds of mental representations—*fears* reports an agent's relation to negative *evaluations* of possible outcomes, while *believes* reports an agent's relation to *judgements* of information as true. I argue that Prior's puzzle arises because *fears* and similarly behaving verbs semantically compose with any entity in their complement that is associable with a possible negative outcome, including information, vehicles of that information, or objects. By contrast, I argue that the Prior's puzzle doesn't arise for *believes* and similarly behaving verbs because these verbs semantically compose with information that can be assigned a truth value, but not with vehicles of information, or with objects. Put differently, while *fears* reports an agent's evaluation of possibilities and projections of what *might* happen, *believes* reports an agent's judgements of what is already the case.

### ***Fears and Substitution***

Substitution puzzles are essential to the philosophy of language. One of the most famous is from Frege (1948), and highlights *substitution failure* of co-referential names

under *believes*. Frege's original formulation uses the names Hesperus and Phosphorus, which *co-refer* to the planet Venus, to show that an agent can be reported to have beliefs about Hesperus—e.g. that it is the first star to appear in the evening—and different, seemingly incompatible beliefs about Phosphorus—e.g. that it is the first star to appear in the morning. The Fregean solution to the puzzle is that though the *referent* of both “Hesperus” and “Phosphorus” is the planet Venus, an agent can believe different things about Venus under different *senses* of the referent—some under the sense of Hesperus/The Morning Star and others under Phosphorus/The Evening Star.

Other variants of substitution puzzles bring out similar phenomena. Kripke (2011) tells the story of puzzling Pierre, who, as a native French speaker believes, in French, “Londres est jolie,” but, upon moving to a rough London neighborhood, starts learning English and acquires the belief, in English, “London is not pretty”—beliefs that are seemingly contradictory since both sentences refer to the same proposition. Another modern variant poses that though Superman and Clark Kent are the same individual, it is true that Lois Lane believes that Superman can fly and that Clark Kent cannot fly, and thus holds what seems like contradictory beliefs about the same individual.

Explaining substitution failures of co-referential expressions under attitude verbs is a puzzling and complex affair, leading to intricate theories about language, logic, mind, and reality. Though many substitution puzzles involve the verb *believes*, there are other intensional verbs that generate puzzles too. *Fears* is one such verb, and an interesting case

at that because there are many robust scientific theories about what it means to fear that discuss brain activation in the amygdala (Adolphs 1995) or cognitive factors that result in a mixture of characteristic nervous physiologic responses and defensive behaviors (Ledoux 2013) that generally trace to William James's theory that the emotions are consciousness of bodily activity (James 1890). From a scientific perspective, one might think that observing conjunctions of certain characteristic physiologic responses and behaviors is sufficient to justify whether an agent *fears*, and much modern neuroscience and psychology operates on this strategy.

However, *fears* may not *necessarily* report nervous physiologic responses and defensive behavior in all cases. One might fear layoffs, that it will rain tomorrow, or fear the news that inflation is increasing, in addition to more straightforward fears of things like bears, snakes, and thunderstorms. In the former set of cases, one would look nothing like an agent exhibiting fight, flight, or freeze behavior in response to stimuli like the latter set of cases, but still be accurately reported as fearing. There are also cases of patients with amygdala damage that still report fear (Ledoux 2017), undermining the idea that particular brain states are necessary conditions for fearing. Further, an agent might fear something and yet endure its presence or occurrence, undermining the idea that a specific set of defensive behaviors is a necessary condition for fearing either. Working out the semantics of *fears* is thus less straightforward than the scientific perspective suggests.

## Representing Attitudes

What does it mean to describe and report the mind or mental states of an agent?

*Attitude reports* represent the mental states of agents with language, so they must have accuracy conditions: they are true if they represent an agent's state of mind accurately and false if they don't. The hard question is articulating those accuracy conditions. Do accurate representations of agents' mental states and processes say if they have, or would have, particular occurrent thoughts in particular situations? Respond to questions in certain ways but not others? Perform certain actions in response to certain conditions? Have certain brain states?

One idea is that since an attitude report describes a mental state at a time, attitude reports are true when some set of psychological processes occur or tend to occur in an agent under certain conditions. However, in his *Philosophical Investigations*, Wittgenstein (2009) resists the idea that any "inner process"—e.g., fearing, remembering, being in pain—sheds any light on the meaning of language about mental states:

"What we deny is that the picture of an inner process gives us the correct idea of the use of the word "remember." Indeed, we're saying that this picture, with its ramifications, stands in the way of our seeing the use of the word as it is."

(Wittgenstein 2009: 109<sup>e</sup>, §305)

If the nature of psychological processes obscures our understanding the meaning of psychological language, we must search for determinants of the meanings of attitude reports elsewhere. Considering language related to the concept of belief in remarks on the philosophy of psychology<sup>1</sup>, Wittgenstein suggests focusing on *behavior* to make sense of the meanings of the kind of psychologically descriptive language that is characteristic of attitude reports:

“This is how I’m thinking of it: Believing is a state of mind. It persists; and that independently of the process of expressing it in a sentence, for example. So it’s a kind of disposition of a believing person. This is revealed to me in the case of someone else by his behavior; and by his words.”

(Wittgenstein 2009: 201<sup>e</sup>, x: §102)

On Wittgenstein’s view, we need information about an agent’s *behavior*—e.g., what they do, what they say, how they choose—to report their attitudes, or at least to report their beliefs. While Wittgenstein resists an explanatory approach from psychology to language, there is promise for going the other way — “*by his words*” — that is, examining the language of attitude reports to for clues about their semantics. If Wittgenstein is right, then this strategy should help us get clearer about the semantics of *fears*.

---

<sup>1</sup> In *Philosophy of Psychology—A Fragment* (Wittgenstein 2009), previously published as *Philosophical Investigations II* and *The Blue and Brown Books*. This material was revised several times preceding the publication of *Philosophical Investigations* and was unfinished at the time of Wittgenstein’s death.

Linguistically, the complements of the intensional verb *fears* in the context of attitude reports if the form *S fears X* can come from a range of categories and parts of speech. They can be concrete (*S fears snakes*), abstract (*S fears the possibility of nuclear war*), folkloric (*S fears ghosts*), fictional (*S fears Freddie Kreuger*), events (*S fears plane crashes*), activities (*S fears skiing*), to enumerate a few possibilities. The plurality of possible objects and expressions that may complement *fears* raises questions about what ties these objects together as a class of fearable things, and what captures the conditions under which an agent can be accurately reported to fear them.

There are different varieties of fear that a theory of fear will need to capture, and more nuance to what replaces the “X” in *S fears X*. The most familiar kind of fear is *experiential fear*, in which an agent has the characteristic nervous physiological responses—e.g. fast heartbeat, sweating, freezing—in response to X (Davis 1987, and cf. Adolphs and Ledoux). The *formal object*—that is, the property attributed to a target by an emotion (Sarantino and de Sousa 2021)—of this kind of fear has been analyzed as the dangerous or unsafe (de Vignemont 2024). Experiential fear can be captured in the linguistic form *S fears NP*, or *S is afraid of NP*, where NP is a noun or noun phrase, e.g. *bears, the tall man in the scary mask, loneliness, ghosts*, that an agent is reported to fear. Another variety of fear is *propositional fear*, where X is a proposition denoted by a that-clause. An agent doesn’t need to experience fear to have a propositional fear, though the two may co-occur. One might be reported to fear *that it will rain on the party* or *that the*

*Yankees will have a terrible year*, with no nervous response, fight, flight, freeze, or occurrent experiential fear response, or (a little) one may fear *that monsters come out at night*, and experience fear responses in the evenings. Propositional fear can be captured in the linguistic form *S fears that P*. Though clearly delineated, these forms of fear generate a difficult semantic puzzle about the meanings of fear ascriptions.

### **Prior's Puzzle**

*Fears* generates an instance of a puzzle about intentionality and intensional verbs called *Prior's puzzle*. Prior's puzzle highlights a substitution failure of co-referential that-clauses and noun phrases under *fears* and some other attitude verbs. The contemporary formulation of the puzzle goes like this: if a that-clause of the form *that P* refers to proposition, and a *propositional phrase* of the form *the proposition that P* also refers to that proposition, then these linguistic items *co-refer*, and therefore should be *substitutable salva veritate* in attitude reports without a generating a change in meaning or semantic value. However, consider the phenomenological difference in (1a) versus (2a):

(1a) Sally fears that Fido barks.

(1b) Sally fears the proposition that Fido barks.

(1a) reports a propositional fear that communicates that Sally fears Fido's barking. More formally, there is some individual, Fido, that barks, such that Sally fears the event or possibility of his barking. (1b) converts the that-clause in the complement of (1a) to a co-referential noun phrase. However, (1b) communicates that Sally fears *a proposition*, which is almost certainly false. (Propositions aren't scary.) If (1a) is true and (1b) is false, we need an explanation for why, since that-clauses and their propositional phrase forms seem to co-refer. Further, the phenomenon does not arise in the context of *believes*, as illustrated in (2a-b):

(2a) Sally believes that Fido barks

(2b) Sally believes the proposition that Fido barks.

Both (2a) and (2b) communicate that Sally takes it to be the case that there is some individual, Fido, that performs the action called barking. Why can *believes* report with that-clauses *or* their co-referential noun phrases to communicate what an agent believes, while *fears* cannot? Put technically, why can co-referential expressions be substituted *salva veritate* in the complement of *believes*, but not in the complement of *fears*? A solution to Prior's puzzle should explain why substitution failure occurs in data like (1a-b) and not (2a-b).

Before discussing solutions to the contemporary formulation of Prior's puzzle, let's examine its original motivations. The earliest expression of Prior's puzzle occurs in his (1963) discussion of *oratio obliqua*, or *indirect speech* that reports the content of an utterance without direct quotation. (3b) illustrates *oratio obliqua* of the utterance in (3a):

(3a) James says, "man is mortal."

(3b) James says that man is mortal.

In (3a), the verb "says" *directly quotes* James, while in (3b), "says" *indirectly reports* the content of his utterance. Prior thinks (3b) is fine as an instance of indirect speech, and so resists interpreting (3b) as a *relation* between James and a proposition, and therefore that parsing it as "James says / that man is mortal" is a "fundamental mistake" (1963: 116). Rather, Prior argues that (3b) should be parsed as "James says that / man is mortal," reading the verb as *says that*, and everything that follows as a *sentential connective*.

For Prior, two interpretations of indirect reports like (3b) are always available: one that takes intensional verbs to be *two-place predicates relating objects designated by names*, and another that takes intensional verbs as *two-place relations between names and sentences*. His (1971) heartedly endorses the latter:

“For expressions like ‘— fears that —’ and ‘— thinks that —’ have precisely this function of forming sentences from other expressions of which the first is a name and the second another sentence. They are as it were predicates at one end and connectives at the other.”

(Prior 1971: 19)

Prior also believes *oratio obliqua* that use *thinks, hopes, fears, wishes*, and other intensional verbs should be interpreted in the same way. That is, as two-place relations between *names* and *sentences*. More radically, he denies not only that that propositions are relata in intentional relations, but that they are objects at all:

“...there is no need to take the component phrase “the proposition that” at all seriously as a name of an object; the phrase “the proposition that—” is merely part of the connective”

(Prior 1963: 119).

On Prior’s *non-propositionalism*, the construction “— *Vs* that—” where *V* is an intensional verb, expresses a relation between a name on the left and a sentence on the right. Prior dispenses with reliance propositions in the interpretation of indirect reports entirely because he thinks that the relata of intensional relations are between objects designated by names and sentences, and further, that sentences don’t designate “anything whatever” (*ibid.*). Prior does qualify that “although we do not *think* sentences, we do think *in*

sentences” (1971: 14), noting that we do *not* hope, believe, fear, desire etc. *sentences*, but rather that attitude reports *express our relation to sentences*.

Prior concludes that relations expressed by intensional verbs are “parasitic” on our relation to our own mental states (1971: 15-16). His puzzle is thus a puzzle about *intentionality*—specifically what the objects of the attitudes can be, given the distinction between a mental act and its object—a distinction going back to Brentano (1874: 88-89) (Grazankowski 2016). Prior expresses the puzzle most succinctly in its canonical gloss:

“(a) X’s thinking of Y constitutes a relation between X and Y when Y exists, but (b) not when Y doesn’t; but (c) X’s thinking of Y is the same sort of thing whether Y exists or not. Something plainly has to be given up here; what will it be?”

(Prior 1971: 130)

Prior’s solution was to give up the denotations of sentences and propositions, but other philosophers have devised solutions without the high cost Prior was willing to pay.

### **Solutions to Prior’s Puzzle**

Philosophers have proposed different solutions to Prior’s puzzle that include arguing that *fears* is polysemous or ambiguous and expresses different relations in different syntactic environments (King 2002), that that-clauses refer to propositions while

propositional phrases refer to functions that return true propositions (Nebel 2019), and that that-clauses refer to predicates rather than to propositions (Güngör 2022). These different strategies target features of *fears* and other intentional verbs that exhibit the phenomenon, or features of the verbs' complement to explain the puzzle.

King (2002) argues that the elements Prior wishes to eschew as denoting and designating nothing whatsoever—sentences and propositions—do indeed designate, and further, can co-refer while having different semantic functions (342). To express Prior's puzzle as linguistic data, King considers environments like (4a-b) and (5a-b) that show the semantic contrast between *believes* and *fears*:

(4a) Russell believes that mathematics reduces to logic.

(4b) Russell believes the proposition that mathematics reduces to logic.

(5a) Amy fears that mathematics reduces to logic

(5b) ? Amy fears the proposition that mathematics reduces to logic.

(5a) and (5b) illustrate the substitution failure under *fears* that needs explaining. King's explanation of the data is that *fears* is ambiguous and likely polysemous relative to the *syntactic type* in its complement—that-clauses determine one kind of relation, while noun

phrases (NPs) determine another. Thus, by King's lights, sentences like (6a-b) express *different relations*, though they seem to report the same mental act:

(6a) Sally remembers that flash floods are dangerous

(6b) Sally remembers the proposition that flash floods are dangerous.

Further, King thinks verbs that generate Prior's puzzle demarcate two major classes of propositional attitude verbs that can take both NP and that-clause complements: (1) *ambiguous verbs of propositional attitude* or AVPs—like *fears*—that express different relations based on the syntactic type of their complement, and (2) *univocal verbs of propositional attitude* or UVPs—like *believes*—that express the same relation regardless of the syntactic type in their complements (*ibid.*: 356-358). AVPs, for King, include some psychological, perceptual, inquiry, desire, and speech verbs such as *fear*, *remember*, *feel*, *understand*, *explain*, *expect*, *hear*, *mention*, *indicate*, *suspect*, *demand*, *desire*, *suggest*, *request*, *imagine*, *know*, and *recommend*. UVPs, for King, include some speech, evaluative, and inquiry verbs, such as *believe*, *doubt*, *deny*, *prove*, *accept*, *assert*, *state*, and *assume*.

Nebel (2019) opposes King's polysemy/ambiguity by syntactic type solution to Prior's puzzle because *fears* fails the *zeugma test*, by which polysemous or ambiguous verbs should produce infelicitous constructions when purportedly different senses of the

verb in question are combined into one expression with one instance of the verb elided.

For example, (7c) combines the supposed different *fears* relations in (7a) and (7b):

(7a) Sally fears ghosts.

(7b) Sally fears the proposition that ghosts haunt the attic.

(7c) Sally fears ghosts and the proposition that they haunt the attic.

Since (7c) is felicitous, *fears* is *not* ambiguous/polysemous by the zeugma test. If *fears* were ambiguous/polysemous, (7c) would sound odd. See for example how different relations expressed by the verb “called” are brought out by the zeugma test in (8):

(8) # Sally called Bert’s bluff and her grandmother for her birthday.

*Call* is ambiguous by the zeugma test because it produces an odd sentence when it must do double work due to the elision—(8) brings out that “calls” is ambiguous between *identifying* and *communicating*.

For Nebel, the evidence in (7c) undermines King’s argument that *fears* is polysemous.<sup>2</sup> The zeugma result for *fears* leads Nebel to focus instead on the meaning of

---

<sup>2</sup> It is debated whether the zeugma test definitively rules that a verb is ambiguous, but assuming it does, (7c) is strong evidence against King’s argument.

intensional verb complements to explain Prior's puzzle. Nebel traces the source of substitution failure to a difference in the *referents* of intensional verb complements when they contain that-clauses and propositional phrases. He thinks the forms *S fears that P* and its presumed co-referring propositional phrase, *S fears the proposition that P*, differ in meaning because they differ in designation and denotation.

In line with standard views of propositions, Nebel thinks that clauses refer to propositions, but takes it that propositional phrases refer to what he calls *propositional concepts*—*functions* from worlds to individuals that return a true proposition about an individual in a situation of evaluation (2019: 90). Thus, for Nebel, when the complement of *fears* includes a *content noun* (Moulton 2009) that contains some information about the world such as “the proposition,” “the rumor,” or “the message,” that phrase *doesn't* denote a proposition, rumor, or message—it refers to a function. As such, fear reports report that an agent fears a function (*ibid.*: 91).

Nebel's position is surprising given that its contention that a propositional phrase of the form *the proposition that P* doesn't refer to a proposition. Although Nebel's account explains Prior's puzzle by fixing the reference of content noun phrases (CNPs), and specifically propositional phrases, to propositional concepts, there is a residual oddity in the denial that the intuitive referents of CNPs in ordinary speech are *not* in fact the referents of those phrases. Rather, it seems phenomenologically intuitive that noun phrases such as “the report that global temperatures are rising” or “the rumor that the

president embezzled” refer to reports and rumors, respectively, and similarly, that “the proposition that flamingos are pink” refers to the proposition that flamingoes are pink. If someone tells me that they fear the rumor of upcoming layoffs, they are certainly in a sense communicating that there is a *rumor* that they fear, in addition to communicating a fear about that rumor’s content. This suggests that the referents of content nouns can be objects of fear independently of the information they contain. Put differently, one can not only fear some information or content, but also fear its packaging.

This still leaves us with the case of fearing propositions. In contrast to Nebel, Güngör (2022) thinks it is that-clauses don’t refer to propositions, and further, that Nebel’s accounting of content noun phrases as *abstract* propositional concepts is problematic because content nouns like “the rumor” or “the news” can compose with causal predicates. One can be hurt, surprised or otherwise *causally affected* by the referents of many content nouns. A rumor may hurt, a statement offend, or a decision disappoint, suggesting that content nouns do not always refer to *abstracta*. Wrestling with the oddity of Nebel’s assertion that propositional phrases refer to abstract functions, Güngör instead explains the substitution failure in Prior’s puzzle by classifying that-clauses as *predicates*:

“That-clauses can be defined as predicates that express relations between contentful entities and their contents. They relate a contentful entity like rumor, statement or proposition to its content expressed by the sentence complementing

the that-clause. In effect: they say what the content of the contentful entity is.”  
(Güngör 2022: 2774)

On this view, the form *fears CNP*, e.g. “Sally fears the rumor /news / report that Fido barks” are about the *content* of the sentence that complements the that-clause in virtue of the that-clause’s relation to a contentful entity that can compose with causal predicates. To explain attitude reports like “Sally fears that Fido barks” that *don’t* have a visible contentful entity, Güngör argues that *predicate restriction* (Kratzer 2006) limits the range of objects a verb can have to propositions expressed by the that-clause in its complement. In these cases, Güngör postulates a hidden item and “presupposes that there is a contentful entity whose content is the proposition that Fido barks without explicitly specifying what the contentful entity is” (2776). Güngör thus explains Prior’s puzzle by the differential specification of contentful entities by predicates—the difference in semantic value is traced to the difference in predicates that explicitly specify a contentful entity (*S fears NP*), and those that don’t (*S fears that P*).

### ***Fears and Representation***

An issue at the core of Prior’s puzzle is what *fears* and other intensional verbs, that-clauses, and noun phrases *represent* in the context of attitude reports. I want to offer a solution to Prior’s puzzle that follows Nebel’s in spirit, but distinguishes between the

*referents* of seemingly co-referential that-clauses and propositional phrases and the *mental representations* agents have of them, and the *role* of those representations play in thought, decision, and action.

I think that-clauses refer to pieces of information, e.g. *that grass is green*, and noun phrases refer to either concrete objects, e.g. snakes, or to *informational vehicles*, e.g. messages, rumors, news etc., that contain pieces of information. I offer that the truth of fear reports hinges on how agents *represent* objects, information, and vehicles in terms of possible outcomes or future states of the world associated with them (cf. Stalnaker 1984: 4)—if an agent *S* negatively evaluates possible outcomes associated with some object, information, or vehicle denoted by a noun, that clause or content noun phrase, respectively—some *X*—and uses that representation to think, decide, and act in ways that avoid those outcomes, then the report *S fears X* is true.

I contend that the substitution failure in Prior's puzzle occurs with *fears* and not with *believes* because *fears* and similarly patterning verbs, e.g. *hopes*, *loves*, and *desires*, report an *evaluation*—an agent's subjectively valenced representations of possible outcomes associated with their complements that play a role in thought, decision making and action, while *believes* and similarly patterning verbs, e.g. *accepts*, *denies*, and *asserts*, report a *judgement*—an agent's representing of some information as true that plays its own, different role in thought, decision making, and action. We can also add a probability condition—if an agent *S* represents some object, information, or vehicle *qua* a set of

associated possible outcomes with a negative valence, and assigns a high probability to those outcomes, *S fears X* more than they would if they assigned a lower probability.

One might fear *rumors*, for instance because of what they *are*—suspected or uncertain assertions—or what they *say* about individuals, things, or situations—their content. One can fear the message or the messenger. This differential is what generates the phenomenology of Prior’s puzzle—since one can fear information *or* its packaging, there is a salient reading of the report *S fears the proposition that P* that is about an agent being afraid of a proposition *qua* object. On my view, those reports are only true and accurately report fear if they *also* meet the valenced possible outcome criteria, including the probability constraint, which will vary across individuals and explains different levels of fearing across agents in relation to the same objects, information, or circumstances. Consider the fear reports in (9a-b) that generate Prior’s puzzle:

(9a) Sam fears that smoking kills.

(9b) Sam fears the proposition that smoking kills.

(9a) is true if Sam assigns a negative valence and a probability within or above some individual-dependent threshold to outcomes associated with the information *smoking kills*, i.e. dying from smoking. This *evaluation of prospects* gives Sam reasons to avoid smoking and shapes behavior consistent with fearing, e.g. Sam refusing cigarettes or any

invitations to smoke. If Sam finds dying by smoking probable but *doesn't* tag it with a negative valence, then (9a) would be false. If Sam assigns a *positive* valence to dying by smoking, not only would (9a) be false, but Sam could be reported as *hoping, loving, or desiring* that smoking kills, depending on further contextual detail.

If (9a) is true per the valence/possible outcome/probability conditions I've proposed, then (9b) is false because Sam assigns negative valence and high probabilities to possible outcomes associated with the information *that smoking kills* captured by the *that*-clause, not with the propositional phrase denoting that information. It is uncoincidental that (9a) is synonymous with "Sam fears the information that smoking kills," and Sam is accurately reported as such if she negatively evaluates and takes the possible outcomes associated with that information, i.e. death from smoking, to be probable. If (9b) and similar constructions are intended to express, for instance, that Sam would be afraid of the proposition *that smoking kills* being used in an argument against her, say if she were a lawyer of a tobacco company testifying to Congress, then (9b) could be interpreted as true with that additional context. In any case, Sam's fear would have to be determined by some further fact pertaining to a proposition *itself*, so the report as phrased in (9b) is false without that additional context.

One *fears* the possible outcomes associated with some object, information, or informational vehicle. By contrast, one *believes* information, not objects or informational vehicles themselves. We do not take it to be the case that snakes or rumors—these

expressions are awfully bad—but in believing we do take it to be the case that *something* is the case. My suggestion is that the something is *information*. Further, when natural language makes it seem as if one can believe objects or informational vehicles, there is a Prior-like indirect speech rejoinder to level. While *fears* reports that an agent represents some possible outcomes negatively and with some above-threshold probability, *believes* rather reports that an agent has already judged some information as true. This *judgement* is about information or content, not about the way that its packaged or contained, nor about particular things. Consider the belief reports in (10a-c):

(10a) Sam believes that smoking kills.

(10b) Sam believes the proposition that smoking kills.

(10c) Sam believes the report that smoking kills.

If (10a) is true, then (10b-c) are true because (10b) can be read as “Sam believes the information that smoking kills,” and (10c) is true insofar as Sam believes what the report *says*—it’s content or the information contained within it. Belief reporting is not concerned with Sam’s consideration of reports in general, and thus targets the content or information that an agent judges to be true. Accordingly, constructions that include informational vehicles express beliefs in the information that they *carry* or *transmit*, not in the vehicle itself. Belief, as it were, goes directly to content.

There are interesting edge cases to consider however, especially when vehicles themselves imply something about truth, as in (12a-b):

(12a) Sam believes the lie that vaccines contain microchips.

(12b) Sam believes the rumor that vaccines contain microchips.

These vehicles aren't neutral, like others such as "the proposition" or "the information," or even "the message," or, in increasingly rare cases, "the news." Rather, they express that Sam believes, in the case of the lie, something false, and, in the case of the rumor, something uncertain that may turn out to be true or false. Still, when we say that she *believes* in these cases, we report what information it is that Sam believes, which itself is denoted by the *that*-clause. The bits about lies and rumors in the reports are rather better read as about *our* evaluation of the objects of *her* belief. This makes for a complex form of attitude report and interpretive strategy. By contrast, substituting *fears* in (12c-d) invites a different semantic story:

(12c) Sam fears the lie that vaccines contain microchips.

(12d) Sam fears the rumor that vaccines contain microchips.

In these cases, Sam may well fear lies and rumors, for all of the reasons that come along with the potential nastiness of these vehicles. Whether Sam *fears* hinges not on her judgement of the information contained in the vehicle, but, as argued, whether she attaches a negative valence and some above threshold probability that possible outcomes associated with lies and rumors about vaccines containing microchips. *Fears* reports Sam's subjective evaluations, not ours.

Zeugma testing generates evidence that *fears* is not polysemous on grounds of its ability to compose with "bare" information denoted by that-clauses, contained information denoted by vehicles of information, or objects. Consider the combined construction of a that-clause and vehicle produced by fusing and eliding (9a) and (12d) with (13a-b) into (13c):

(13a) Sam fears snakes.

(13b) Sam fears ghosts.

(13c) Sam fears snakes, ghosts, that smoking kills, and the rumor that vaccines contain microchips.

Since (13c) is felicitous, the zeugma test rules that *fears* is not polysemous or ambiguous on grounds of the object/information/informational vehicle distinction. However, this leaves us with "the proposition" — which, if zeugma-tested in (13c) yields:

(14) Sam fears snakes, ghosts, that smoking kills, and the rumor that vaccines contain microchips, and ?the proposition that the earth is flat.

I'm unsure that (14) is felicitous with the addition of the propositional phrase. Thus, it may be the case that while there isn't vast polysemy and ambiguity of *fears*, that there is a degree of ambiguity, namely between *fears* when complemented by objects, information, and informational vehicles, versus *fears* complemented by propositional phrases. There may be some quirk to propositional phrases that I am missing. Nevertheless, we can still classify the propositional phrase fragment of (14) as bad by the other criteria I have stipulated for the truth of fear reports. These are topics for further research.

Overall, I hope to have given force to the idea that truth conditions of fear reports have to do with the way that an agent *evaluates* possible outcomes, while the truth conditions of belief reports have to do with how an agent *judgets* information.

### **A Remark on Affect**

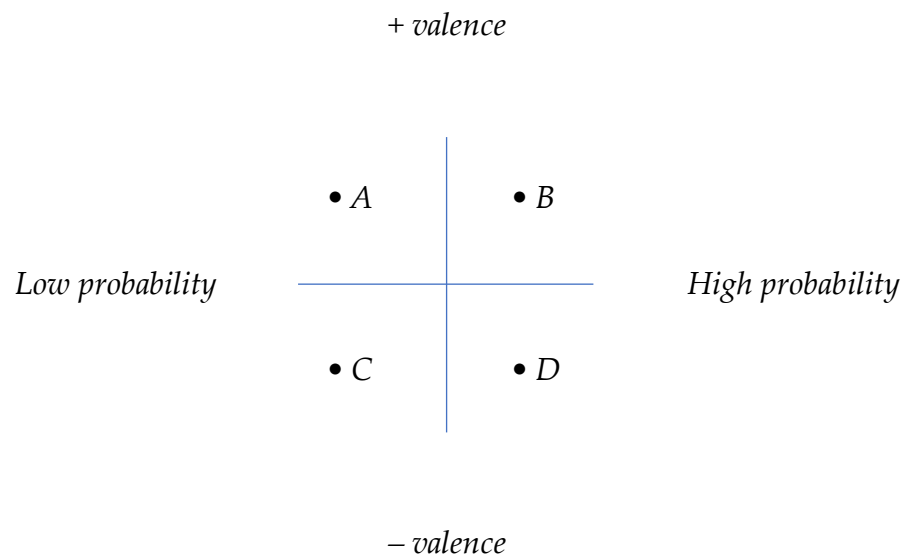
What further conclusions might we draw about fear and belief *the attitudes* given the above account? If I am right and fear reports are true based on whether an agent represents possible outcomes with a negative valence and a threshold probability, then

*fears* represents a particular kind of relation that an agent has to their own representations of the world. The application of a valence and a probability to these representations makes them agent-relative, and accounts for the wide variation in fearing across individuals. Different individuals fear differently—some fear, Fido’s bark while others don’t. On the present theory, this is accounted for by an agent’s relation to negatively valenced representations of possible outcomes fearers have in their minds that play a causal role in behaviors that aim to avoid those outcomes. One might look for these behaviors in observation, then look for neural evidence of these representations in the brain in circuits associated with subjective value, projection and action planning.

The assignment of negative valence is a function of what I’ll call an agent’s *affect*—a base for evaluation and decision that biases an agent toward certain judgements and emotions (see McMartin and Pickavance 2024). Let’s go back to “Sam fears that smoking kills.” Say that several people in Sam’s family have died of lung cancer induced by smoking, and that Sam is a biologist who studies the effects of carcinogens found in cigarettes on cancer development. Sam’s exposure to the consequences of smoking for others and her awareness about the causal link between smoking and cancer contribute to her *affective state* when it comes to processing situations that involve smoking. These subjective factors induce her mind to negatively tag information about smoking that then influences her behavior, say to pressure her friends that smoke to stop. Say, by contrast, that Sam has only had a positive history with smoking—she finds it pleasurable, has

never met anyone who succumbed to its health effects, isn't aware of the relevant biology, and is enamored by its coolness in films. These facts about Sam contribute instead to a *positive* affective state that induces neutral or even positive evaluations of information about smoking—*hopes*, *loves*, and *desires* for instance. She may even *disbelieve* that smoking kills until she gets some evidence to the contrary.

Fear thus can be understood as lying on at least a two dimensional axis that plots outcomes as a function of valence and probability:



On this *affective axis*, an agent *fears* when some possible outcome they represent in association with some situation falls below the x-axis, and the degree to which the fear increases from left to right. If this axis models an agent *S*, then *S fears D* more than *S fears C*, and we might report that *S hopes for A* and that *S anticipates / expects B*. Modeling an

agent on this way can help get a grip on what it is we mean when we report attitudes, especially those that involve projected future states of the world and varying degrees of uncertainty.

On the subject of belief, the present theory holds that to believe is to represent some information as being the case—to represent it as true. An agent's beliefs may play roles in downstream psychological processes that generate heavily valenced emotional or physical responses, but it is important to separate the idea of these representations from the role they play in mental processes. The thought is that beliefs themselves, since they are about what is the case, are neutrally valenced—belief *itself* is the treatment of some information as true, and does not necessarily include some affective appraisal or weighing of possible outcomes as fear, desire, or other attitudes do. In fact, it is part of the nature of beliefs that they can persist *despite* such evaluations.

## References

- Adolphs, Ralph, Daniel Tranel, Hanna Damasio, and Antonio R. Damasio. Fear and the Human Amygdala. *The Journal of Neuroscience* 15(9): 5879-5891).
- Davis, Wayne A. (1987). The Varieties of Fear. *Philosophical Studies* 51 (3): 287-310.

- de Vignemont, Frédérique (2024). Fear Beyond Danger. *Mind and Language* 39 (5): 647-663.
- Frege, Gottlob (1948). Sense and Reference. *Philosophical Review* 57 (3): 209-230.
- Gordon, Robert M. (1980). Fear. *The Philosophical Review* 89 (4): 560-578.
- Grzankowski, Alex (2016). Attitudes Towards Objects. *Noûs* 50 (2):314-328.
- Güngör, Hüseyin (2022). That solution to Prior's puzzle. *Philosophical Studies* 179 (9):2765-2785.
- James, William (1890). The Emotions. In *The Principles of Psychology*, Chapter 25. Henry Holt and Company.
- King, Jeffrey C. (2002). Designating Propositions. *Philosophical Review* 111 (3), 341-371.
- Kratzer, A. (2006). Decomposing attitude verbs. Talk in honor of A. Mittwoch. The Hebrew University, Jerusalem.
- Kripke, Saul A. (2011). A Puzzle about Belief. In *Philosophical Troubles: Collected Papers, Volume 1*. New York: Oxford University Press. pp. 125-160.
- Ledoux, Joseph E. (2013). The Slippery Slope of Fear. *Trends in Cognitive Sciences* 17 (4): 155-156.
- Ledoux, Joseph E. (2017). Semantics, Surplus Meaning, and the Science of Fear. *Trends in Cognitive Sciences* 21(5): 303-306.

- McMartin, Jason and Pickavance, Timothy (2024). Affective Reason. *Episteme* 21: 819-836.
- Moulton, K. (2009). *Natural Selection And the Syntax of Clausal Complementation*. PhD dissertation, UMass, Amherst.
- Nebel, Jacob M. (2019). *Hopes, Fears, and Other Grammatical Scarecrows*. *Philosophical Review* 128 (1): 63-105.
- Prior, Arthur Norman (1971). *Objects of Thought*. Oxford: Clarendon Press. P. T. Geach & Anthony Kenny, Eds.
- Prior, Arthur N. (1963). Symposium: Oratio Obliqua. *Proceedings of the Aristotelian Society*, Supplementary Volumes 37: 115-46.
- Scarantino, Andrea and Ronald de Sousa, "Emotion", *The Stanford Encyclopedia of Philosophy* (Summer 2021 Edition), Edward N. Zalta (ed.)
- Stalnaker, Robert C. (1999). Belief Attribution and Context. In *Context and Content: Essays on Intentionality in Speech and Thought*. Oxford University Press.
- Stalnaker, Robert (1984). *Inquiry*. Cambridge University Press.
- Wittgenstein, Ludwig (1922) *Tractatus Logico-Philosophicus*. D.F. Pears and B.F McGuinness (trans.). New York and London: Routledge.
- Wittgenstein, Ludwig (2009). *Philosophical Investigations* (4th ed.). G. E. M. Anscombe, P. M. S. Hacker and J. Schulte, Trans. Wiley-Blackwell.