

Attitude reports: From modal logic to event semantics

Day 5: Bad attitudes

Tom Roberts

Utrecht University

OVA @ Debrecen, 31 July 2026

The final question

Languages clearly have a variety of possible attitude meanings.

The final question

Languages clearly have a variety of possible attitude meanings.

Are there plausible attitude meanings that nevertheless never get put in a single lexical package?

The final question

Languages clearly have a variety of possible attitude meanings.

Are there plausible attitude meanings that nevertheless never get put in a single lexical package?

Today: A case study on this question.

Case study: the contrafactive puzzle (Joint work with Deniz Özyıldız)

Factives and contrafactive

Actual **Factives**

x Vs that $p \rightsquigarrow_{psp} p$

- (1) Mary *knows* that Bonn is in West Germany.
- a. Assertion
Mary believes that Bonn is in West Germany.
 - b. Presupposition
Bonn is in West Germany.

Factives and contrafactualives

Actual **Factives**

$x \text{ Vs that } p \rightsquigarrow_{\text{psp}} p$

- (1) Mary *knows* that Bonn is in West Germany.
- a. Assertion
Mary believes that Bonn is in West Germany.
 - b. Presupposition
Bonn is in West Germany.

Hypothetical **Contrafactualives**

$x \text{ Vs that } p \rightsquigarrow_{\text{psp}} \neg p$

- (2) Mary *contras* that Bonn is in West Germany.
- a. Assertion
Mary believes that Bonn is in West Germany.
 - b. Presupposition
Bonn is not in West Germany.

The contrafactive gap

There are many factive predicates* across the world's languages (*know, be surprised, forget*, etc.), but...

Commitment (to be motivated):

Contrafactive predicates do not exist—or, if they do, they are much rarer than factive ones.

The contrafactive gap

	inference of truth	inference of falsity
via entailment	<i>be right</i>	<i>be wrong</i>
via presupposition	<i>know</i>	<i>['contra']</i>

The contrafactive gap

	inference of truth	inference of falsity
via entailment	<i>be right</i>	<i>be wrong</i>
via presupposition	<i>know</i>	<i>['contra']</i>

Adding to the puzzle is that predicates that *entail* the falsity of their complement do exist, namely, at least *be wrong*.

⇒ The gap might have something to do with presupposing falsity.

Central queries

- ✧ (Why) should we believe in the No Contra-factives Conjecture?
- ✧ If contra-factives don't exist (or are rare), is there a principled reason for their absence (or rarity)?

Coexisting stances

✧ There are no contra-factives

- because there are no contra-facts for their complements to refer to. Holton (2017)
- because they are difficult to learn, maybe. Strohmaier & Wimmer (2022, cf. 2025)
- because we prefer to attribute true beliefs. Sander (2024)

Coexisting stances

- ✧ There are no contra-factives
 - because there are no contra-facts for their complements to refer to. Holton (2017)
 - because they are difficult to learn, maybe. Strohmaier & Wimmer (2022, cf. 2025)
 - because we prefer to attribute true beliefs. Sander (2024)

- ✧ There are contra-factives, e.g., *creerse* or *wähnen*, and many different ways that languages express falsity of belief (no 'gap')
 - Rosenberg (1975)
 - Sander (2024); Anvari et al. (2019), cf. Maldonado & Percus (2024)
 - Glass (2022); Bossi (2022); McGregor (2023); a.o.

Coexisting stances

- ✧ There are no contra-factives
 - because there are no contra-facts for their complements to refer to. Holton (2017)
 - because they are difficult to learn, maybe. Strohmaier & Wimmer (2022, cf. 2025)
 - because we prefer to attribute true beliefs. Sander (2024)

- ✧ There are contra-factives, e.g., *creerse* or *wähnen*, and many different ways that languages express falsity of belief (no 'gap')
 - Rosenberg (1975)
 - Sander (2024); Anvari et al. (2019), cf. Maldonado & Percus (2024)
 - Glass (2022); Bossi (2022); McGregor (2023); a.o.

- ✧ What even is a contra-factive? Glass (2024); Sander (2024)

Coexisting stances

We all agree that natural languages have ways or dedicated morphemes to attribute false beliefs to people.

But whence the unease regarding contra-factives?

At least three sources:

- ✧ Deciding on and interpreting the criteria for contra-factivity.
- ✧ Categorizing predicates with variable behavior.
- ✧ Lack of complete descriptions.

(+The search can only end if we find one, exhaust all predicates in all languages, or agree on why *contra* can't exist.)

We don't have enough empirical evidence to reject the “No Contra-factives Conjecture,” and we propose to explain it.

The proposal

There is a pressure against lexicalizing *contra* from **ontological incoherence** of the state it describes.*

- ★ Lexical presuppositions must be **causally upstream** of at-issue content

*À la Roberts & Simons' (2024) presuppositions as ontological preconditions.

Outline

1. Empirical discussion:

Putative counterexamples are not *contra*

Here we will echo some of what we know already and show our work.

2. Proposal

I. Causation and lexicalization

II. Explaining ✓ *know* vs **contra*

Empirical discussion

What are we looking for?

Holton's (2017) criteria

For a predicate *contra* to count as a contra-factive, it must...

- ✧ be doxastic
 - “express a mental attitude”
- ✧ take that-clauses
- ✧ be stative
- ✧ be monomorphemic
- ✧ presuppose that its complement is false
 - “stand [in the same relation] to the falsity of its complement as factives stand to the truth of theirs”

These are at least good in that we can collectively refer to them.

Commentary

✧ Are these criteria too stringent?

- Are we looking for *a* contra-factive, or the simplest one?

Is there reason to believe that none can exist if there isn't a simplest one? (Can 'come to know' be lexicalized in a language that doesn't have 'know'?)

Commentary

✧ Are these criteria too stringent?

- Are we looking for *a* contra-factive, or the simplest one?

Is there reason to believe that none can exist if there isn't a simplest one? (Can 'come to know' be lexicalized in a language that doesn't have 'know'?)

- Factives are exempt from some of the criteria: *find out*, *be surprised*.

Commentary

✧ Are these criteria too stringent?

- Are we looking for *a* contra-factive, or the simplest one?

Is there reason to believe that none can exist if there isn't a simplest one? (Can 'come to know' be lexicalized in a language that doesn't have 'know'?)

- Factives are exempt from some of the criteria: *find out*, *be surprised*.

✧ Regarding the presupposition of falsity, our standards shouldn't be different for the presupposition of truth.

- False belief inferences disappear (as we will see shortly). But so do true belief inferences.

Commentary

✧ Are these criteria too stringent?

- Are we looking for *a* contra-factive, or the simplest one?

Is there reason to believe that none can exist if there isn't a simplest one? (Can 'come to know' be lexicalized in a language that doesn't have 'know'?)

- Factives are exempt from some of the criteria: *find out*, *be surprised*.

✧ Regarding the presupposition of falsity, our standards shouldn't be different for the presupposition of truth.

- False belief inferences disappear (as we will see shortly). But so do true belief inferences.
- Should we then scrutinize our belief in factives? (Yes.)

Commentary

✧ Are these criteria too stringent?

- Are we looking for *a* contra-factive, or the simplest one?

Is there reason to believe that none can exist if there isn't a simplest one? (Can 'come to know' be lexicalized in a language that doesn't have 'know'?)

- Factives are exempt from some of the criteria: *find out*, *be surprised*.

✧ Regarding the presupposition of falsity, our standards shouldn't be different for the presupposition of truth.

- False belief inferences disappear (as we will see shortly). But so do true belief inferences.
- Should we then scrutinize our belief in factives? (Yes.)
- Is there still an asymmetry between factives and contra-factives? (Also yes.)

(White & Rawlins 2018; Glass 2024; Sander 2024)

What do we find?

As we see it, candidate contra-factives fall into three bins:

- ✧ They ‘obviously’ fail one or more of these criteria.
- ✧ It is difficult to decide if they satisfy them.
- ✧ The data describing their inferential patterns are conflicting or incomplete.

For each candidate here, we don’t have all the judgments that we would want to be able to conclude one way or the other.

Let’s go through some representative examples.

Defeasible falsity inferences

Mandarin *yǐwéi* (Glass 2022), Kipsigis *par* (Bossi 2022)

- (3) a. rénmen *yǐwéi* tā shì yìwànfùwēng...
 person.PL YǐWÉI 3S be billionaire
 b. ér tā díquè shì.
 and 3S indeed be
 ‘People are under the impression that she’s a
 billionaire... and she actually is.’
 (adapted from Glass 2022: ex. 9)

The continuation in (3b) is “non-contradictory (though marked).”

Non-defeasible falsity inferences...

Spanish *creerse* (Anvari et al. 2019, Maldonado & Percus 2024)

- (4) a. Juan *se cree* que Ana tiene 30 años...
Juan REFL believe that Ana has 30 years
- b. #¡y tiene razón!
and has right
Intended: Juan believes that Ana is 30 years old... and
he's right! (Anvari et al. 2019: ex. 10)

In the positive, the FBI associated with *creerse* can't be canceled
(while it can be with plain *creer*).

...that aren't always triggered

The inferential profile of *creerse* is nuanced, and is conditioned (at least) by negation and mood.

Ambiguity with negation + indicative

- (5) Hilda **no se cree** que Jirafales **es** honesto.
Hilda not REFL believe that Jirafales is.IND honest.

a. Contra-factive understanding

↪ Jirafales is not honest and Hilda doesn't believe he is.

b. Factive understanding

↪ Jirafales is honest and Hilda believes he isn't.

(Maldonado & Percus 2024: ex. 5)

...that aren't always triggered

Spanish *creerse* (Anvari et al. 2019, Maldonado & Percus 2024)

The inferential profile of *creerse* is nuanced, and is conditioned (at least) by negation and mood.

No FBI with negation + subjunctive

(6) Hilda **no se cree** que Jirafales **sea** honesto.

Hilda not REFL believe that Jirafales is.SUBJ honest.

↪ Hilda fails to believe that Jirafales is honest

Maldonado & Percus (2024: ex. 7)

Non-defeasible falsity inferences that aren't always triggered

How to categorize a predicate like *creerse*?

Obvious and non-obvious violations of the criteria

Dutch *wijsmaken* (Hoeksema 2021)

- (7) Ze maakte me wijs dat ze rijk was...
she made me wise that she rich was

#en ze was wel rijk.
and she was PRT rich

‘She fooled me into thinking that she was rich #and she *was* rich.’

(modified from Hoeksma 2021: ex. 3a to include the continuation)

- ✧ A clear doxastic component, but this is a *change of state* verb.
- ✧ The ‘attitude holder’ is a direct object.
- ✧ (Not monomorphemic.)

Conflicting data: Strong inference?

Tagalog *akala* (Kierstead 2013, Jed Pizarro-Guevara, p.c.)

Kierstead: *akala* comes with a strong falsity inference.

(8) *Akala* ni Jim na may party.

AKALA IND Jim COMP exist party

‘Jim falsely believes that there’s a party.’

a. ✓ if I know there’s no party.

b. #if I know there’s a party.

c. #if I don’t know if there’s a party or not.

(adapted from Kierstead 2013: ex. 9–11)

Conflicting data: Or not?

But a speaker different from Kierstead's allows for suspending the predicate's FBI in an explicit ignorance context.

- (9) Di ko alam kung nakascore siya,
NEG I know Q was.able.score 3SG
peru **akala** niya siguro nakascore siya.
but AKALA 3SG perhaps was.able.score 3SG
'I don't know whether he was able to score, but he thinks
that he might have scored.' (Jed Pizarro-Guevara, p.c.)

Incomplete data

Yanyuwa (Pama-Nyungan) *yidirri* (Kirton and Charlie 1996: 205; via McGregor 2023)

- (10) Katharra-*yudirri* *katha* babalu
we.DU.EXCL-YUDIRRI KATHA buffalo

ngala kurdardi ngala wakardawakarda.
but no but bull

‘We mistakenly thought a buffalo (was there) but (it was) not, (there was a) bull.’

(Kirton and Charlie 1996: 205; via McGregor 2023)

- ✧ *Katha*, with which the verb *yudirri* occurs, is described as a ‘mistaken belief’ particle.
- ✧ Does *yudirri* ever occur alone? does it imply falsity? (We don’t know.)

Some notable candidatess we didn't talk about

✧ Taiwanese Southern Min *lih-tsun*

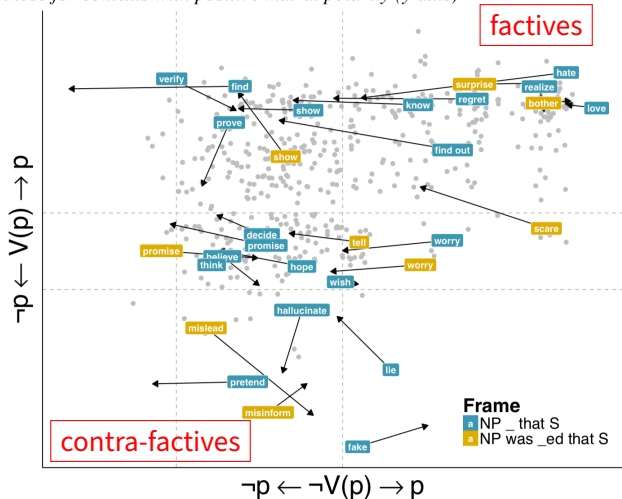
Hsiao (2017)

✧ German *wähnen*

Sander (2024)

No contra-factives in the English lexicon

- (14) *Normalized responses for contexts with negative matrix polarity (x-axis) against those for contexts with positive matrix polarity (y-axis)*



Central queries

- ✧ (Why) should we believe in the No Contra-factives Conjecture?

No clear example yet?

- ✧ If contra-factives don't exist (or are rare), is there a principled reason for their absence?

Causation and lexicalization

Assumptions

Lexical items encode (at least) two ‘tracks’ of information:
presuppositions and at-issue content

Assumptions

Lexical items encode (at least) two ‘tracks’ of information:
presuppositions and at-issue content

Core intuition: Presuppositions are conceptually prior to at-issue content (Schlenker 2021, Roberts & Simons 2024)

Assumptions

Lexical items encode (at least) two ‘tracks’ of information:
presuppositions and at-issue content

Core intuition: Presuppositions are conceptually prior to at-issue content (Schlenker 2021, Roberts & Simons 2024)

★ Our main point is compatible with a view in which
presuppositions are pragmatically enriched entailments (à la Stalnaker
1970, 1973, 1974; Rosenberg 1975; Karttunen & Peters 1979; Kadmon 2001; Abusch 2002,
2010; Simons et al. 2010, 2016; Abrusán 2011, 2016; ...)

Constraints on lexicalization

Main idea: We cannot lexicalize configurations of meaning in which the presupposition does not 'lead to' at-issue content

Caveat: only in attitudes?

Constraints on lexicalization

Main idea: We cannot lexicalize configurations of meaning in which the presupposition does not 'lead to' at-issue content

Caveat: only in attitudes?

$\rightsquigarrow$ Since *know* presupposes p and asserts *believe p* (Bp), we expect p 'leads to' belief in p

Constraints on lexicalization

Main idea: We cannot lexicalize configurations of meaning in which the presupposition does not ‘lead to’ at-issue content

Caveat: only in attitudes?

$\rightsquigarrow$ Since *know* presupposes *p* and asserts *believe p* (*Bp*), we expect *p* ‘leads to’ belief in *p*

Goal: Cash out this intuition using causal models and show that *contra* does not plausibly obey the constraint

Constraints on lexicalization

Main idea: We cannot lexicalize configurations of meaning in which the presupposition does not ‘lead to’ at-issue content

Caveat: only in attitudes?

$\rightsquigarrow$ Since *know* presupposes p and asserts *believe p* (Bp), we expect p ‘leads to’ belief in p

Goal: Cash out this intuition using causal models and show that *contra* does not plausibly obey the constraint

We focus here on the link between p and Bp , but certainly knowledge involves more than that!

Constraints on lexicalization

Main idea: We cannot lexicalize configurations of meaning in which the presupposition does not ‘lead to’ at-issue content

Caveat: only in attitudes?

$\rightsquigarrow$ Since *know* presupposes p and asserts *believe p* (Bp), we expect p ‘leads to’ belief in p

Goal: Cash out this intuition using causal models and show that *contra* does not plausibly obey the constraint

We focus here on the link between p and Bp , but certainly knowledge involves more than that!

- (11) **Lexicalization constraint (preliminary, to be revised)**
An eventuality-denoting predicate with at-issue content α can have a presupposition π iff we can construct a **causal model** with a chain from π to α .

Causal models (simplified from Pearl 2000)

A CM is a triple $\langle U, V, E \rangle$:

U : set of exogenous variables

V : set of endogenous variables

E : set of structural equations f_i such that $\forall v_i \in V, v_i = f_i(pa_i)$,
where $pa_i \subset U \cup V$

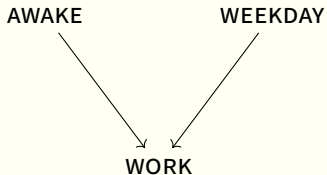
Profitably used in many linguistic domains Nadathur & Lauer 2020; Baglini & Bar-Asher Siegal 2020; Copley & Mari 2022; Glass 2023; McHugh 2023; Nadathur 2023

Causal model example

You go to WORK iff you are AWAKE on a WEEKDAY.

Causal model example

You go to WORK iff you are AWAKE on a WEEKDAY.



U: { AWAKE, WEEKDAY }

V: { WORK }

E: { $WORK = AWAKE \wedge WEEKDAY$ }

Deriving *know* vs. *contra*

Why *know* obeys the constraint

Why *know* obeys the constraint

What connects the presupposition p with at-issue Bp ?

Why *know* obeys the constraint

What connects the presupposition p with at-issue Bp ?

A sketch:

- 1) The fact of p yields some in-principle observable evidence for p ('indicator')

Why *know* obeys the constraint

What connects the presupposition p with at-issue Bp ?

A sketch:

- 1) The fact of p yields some in-principle observable evidence for p ('indicator')
- 2) Some agent becomes aware of this indicator

Why *know* obeys the constraint

What connects the presupposition p with at-issue Bp ?

A sketch:

- 1) The fact of p yields some in-principle observable evidence for p ('indicator')
- 2) Some agent becomes aware of this indicator
- 3) The acquaintance with an indicator is a precondition for the belief

A three-step process for (part of) knowing

Step 1: The fact of p is sufficient to generate some **indicator** i_p for p

A three-step process for (part of) knowing

Step 1: The fact of p is sufficient to generate some **indicator** i_p for p

$$p \longrightarrow \text{indic}(p)$$

Dependency: $p \rightarrow \text{indic}(p)$

($\text{indic}(p)$ = 'there is an indicator i for p ')
(i_p = 'the indicator i for p ')

Step 2: The existence of i_p is necessary for an agent to become acquainted with it

$$\text{indic}(p) \longrightarrow \text{acq}(a)(i_p)$$

Dependency: $\neg \text{indic}(p) \rightarrow \neg \text{acq}(a)(i_p)$

Step 3: a 's acquaintance with i_p is necessary for them to believe that p ($\approx$ people form beliefs iff they have reason to)

$$acq(a)(i_p) \longrightarrow B(a)(p)$$

Dependency: $\neg acq(a)(i_p) \rightarrow \neg B(a)(p)$

A causal chain of *knowing*

- (12) *Context: Marianne is learning about geography by looking at an accurate contemporary map.*
Marianne knows that Bonn is in Germany.

A causal chain of *knowing*

- (12) *Context: Marianne is learning about geography by looking at an accurate contemporary map.*
Marianne knows that Bonn is in Germany.

B I N G

A causal chain of *knowing*

- (12) *Context: Marianne is learning about geography by looking at an accurate contemporary map.*
Marianne knows that Bonn is in Germany.

B IN G $\longrightarrow$ INDIC(B IN G)

A causal chain of *knowing*

- (12) *Context: Marianne is learning about geography by looking at an accurate contemporary map.*
Marianne knows that Bonn is in Germany.

$$B \text{ IN } G \longrightarrow \text{INDIC}(B \text{ IN } G) \longrightarrow \text{ACQ}(i_{B \text{ IN } G})$$

A causal chain of *knowing*

- (12) *Context: Marianne is learning about geography by looking at an accurate contemporary map.*
Marianne knows that Bonn is in Germany.

$$B \text{ IN } G \longrightarrow \text{INDIC}(B \text{ IN } G) \longrightarrow \text{ACQ}(i_{B \text{ IN } G}) \longrightarrow B_M(B \text{ IN } G)$$

***Contra* is badly behaved**

Contra is badly behaved

What if Marianne sees this map and concludes that Bonn is in West Germany?



Does she *contra* that Bonn is in West Germany?

***Contra* is badly behaved**

$$\llbracket \text{M kontras that B in WG} \rrbracket = \neg \underline{\text{B in WG}}.B_M(\text{B in WG})$$

Contra is badly behaved

$$\llbracket M \text{ contra that } B \text{ in WG} \rrbracket = \neg B \text{ in WG} . B_M(B \text{ in WG})$$

$$\neg B \text{ IN WG} \longrightarrow \text{INDIC}(\neg B \text{ IN WG})$$

??

$$\neg \text{INDIC}(B \text{ IN WG}) \longrightarrow \text{ACQ}(i_{B \text{ in WG}}) \longrightarrow B_M(B \text{ IN WG})$$

Contra is badly behaved

$$\llbracket M \text{ contra that } B \text{ in WG} \rrbracket = \neg B \text{ in WG} . B_M(B \text{ in WG})$$

$$\neg B \text{ IN WG} \longrightarrow \text{INDIC}(\neg B \text{ IN WG})$$

??

$$\neg \text{INDIC}(B \text{ IN WG}) \longrightarrow \text{ACQ}(i_{B \text{ in WG}}) \longrightarrow B_M(B \text{ IN WG})$$

★The fact of $\neg p$ does not itself lead to p !

Retooling the lexical constraint

Retooling the lexical constraint

Problem: Sometimes the fact of $\neg p$ leads to the belief p .

Retooling the lexical constraint

Problem: Sometimes the fact of $\neg p$ leads to the belief p .

Spy case: Marianne thinks that Sasha is not a spy, because she is not behaving suspiciously. In reality, Sasha is a spy who is very good at concealing her spy-ness by acting normal.



Spy problem

(13) Marianne **contras** that Sasha is not a spy.

Spy problem

(13) Marianne **contras** that Sasha is not a spy.

$$spy(s) \longrightarrow indic(\neg spy(s)) \longrightarrow acq(m)(i_{\neg spy}) \longrightarrow B_m(\neg spy(s))$$

Solution: generalizing the constraint

We don't want bizarre cases to prevent *contra* from being ruled out
⇒ need to **generalize our constraint**

Solution: generalizing the constraint

We don't want bizarre cases to prevent *contra* from being ruled out
⇒ need to **generalize our constraint**

- (14) **Lexicalization constraint, final:** A verbal predicate V with at-issue content α can have a presupposition π iff in the **normative world** w_n a causal chain from $\pi(w_n)$ to $\alpha(w_n)$ can be constructed for **any assignment** of the variables of the arguments of V .

Solution: generalizing the constraint

We don't want bizarro cases to prevent *contra* from being ruled out
⇒ need to **generalize our constraint**

- (14) **Lexicalization constraint, final:** A verbal predicate V with at-issue content α can have a presupposition π iff in the **normative world** w_n a causal chain from $\pi(w_n)$ to $\alpha(w_n)$ can be constructed for **any assignment** of the variables of the arguments of V .

⇒ *contra* unlexicalizable because under normal circumstances, there is some $\neg p$ that isn't causally upstream of Bp

The punchline

We have made a lot of progress in formalizing the semantics of attitudes.

The punchline

We have made a lot of progress in formalizing the semantics of attitudes.

Much of this progress was centered on solving compositional problems.

Understanding these two main accounts will equip you to make sense of much of contemporary semantics literature

The punchline

We have made a lot of progress in formalizing the semantics of attitudes.

Much of this progress was centered on solving compositional problems.

Understanding these two main accounts will equip you to make sense of much of contemporary semantics literature

These problems have not all been solved yet, but other issues surrounding the lexicalization of attitudes have not been attended to all that much.

The punchline

We have made a lot of progress in formalizing the semantics of attitudes.

Much of this progress was centered on solving compositional problems.

Understanding these two main accounts will equip you to make sense of much of contemporary semantics literature

These problems have not all been solved yet, but other issues surrounding the lexicalization of attitudes have not been attended to all that much.

We need a **much** wider range of linguistic data to make serious traction here. Your turn!