

Bayesian Semantics and Japanese Evidential Particle *Ya*

Daiki Matsumoto

Kanazawa Seiryō University

matsumoto@seiryō-u.ac.jp

Observation 1: We present a formal semantics for an understudied discourse particle *ya* in Japanese. *Ya* (neither a dialectal variant of *yo* nor a connective *ya*) has received scarce attention, but recent descriptive investigations have revealed some of its aspects. [1] observes that *ya* is typically used for the speaker (S)’s self-talk (inner speech), as in (1). If (1) is *ya*-less, it just asserts the prejacent proposition *p*. By contrast, if *ya* is used, it conveys that S newly obtains *p* upon utterance (cf. [1]). This is supported by the fact that *ya* is illicit in (1) if S knew that it’s cold in the room beforehand. But “newness of *p*” *per se* doesn’t pertain to *ya*: (2) is licit even if S has the prior knowledge that it’s cold in Canada. Rather, according to [2, p.62], *ya* is used in a sentence which expresses the renewal of S’s doxastic state with *p* (cf. [3, p.118]). [2] argues that this explains (1) and (2), as well as data such as (3) and (4) (both from [2]).

- (1) Uwa (wow), kono (this) heya (room) samui (cold) (*ya*). ‘Wow, it’s cold in this room.’
- (2) (Uttered as S arrives in Canada:) Kanada-wa (C.-TOP) samui *ya*. ‘It’s cold in Canada.’
- (3) Huraito (flight) kekkou (be.canceled) siteru (do) *ya*. ‘The flight has been canceled.’
- (4) Uwa, mata (again) bagut-teru (glitch-PROG) *ya*. ‘Ugh, the app is glitching again.’

Problems: What it means is unclear that *ya* is used in a “sentence expressing the renewal of S’s doxastic state”. In (1), S comes to believe *p* as they utter the sentence (i.e. S renews their doxastic state). But no doxastic change wrt *p* alone can be witnessed in (2). Rather, S utters (2) as S comes to acquire new evidence for *p*: they make it public that they have acquired perceptual evidence for *p* upon arrival in Canada, and *ya* highlights this. However, this is still not enough. (5) is uttered as S (newly) holds the belief that the addressee (A) doesn’t like sea urchin sushi from A’s facial expression as (new) evidence. But *ya* is infelicitous. Interestingly, (6), in which S asserts that they don’t like it, is perfectly fine with *ya*.

- (5) (Seeing A’s face twist after eating sea urchin sushi:)
Kimi-wa (you-TOP) sore (it) suki-zyanai (like-NEG) (*#ya*). ‘You don’t like it.’
- (6) (Uttered as S eats sea urchin sushi for the first time:)
Uwe (ew), watasi (I) kore suki-zyanai (*ya*). ‘Ew, I don’t like this.’

The difference between (5) and (6) lies at the fact that while (5) asserts that a certain property (not liking sea urchin sushi) holds true of A, (6) expresses that the same property holds true of S, the utterer of the sentence. The contrast is particularly telling. Generally speaking, each agent should be more certain of which property holds true of themselves than of any other individual. In (5), A should be more aware of their own taste than S, but S intends to assert something about A’s taste. The assertion *per se* is not illicit, as the felicity of the *ya*-less sentence in (5) shows (adding some other particle, such as *ne* and *ka*, makes it sound more natural, which we don’t discuss in this talk). However, when *ya* is used, the utterance becomes awkward. By contrast, (6), in which *ya* can be used, asserts that the property of liking sushi holds true of S, whose taste is the most well-known by S themselves. The comparison between (5) and (6) indicates that for the use of *ya* to be licit, the new evidence S holds must be strong. More specifically, the use of *ya* suggests that S holds new evidence for *p* which is the strongest among all types of evidence

available to all discourse agents (DAs). This is violated in (5), as A has stronger evidence for p (e.g., their own perceptual evidence about their liking or disliking should be stronger than S's inference about it based on A's facial expression). The same condition is satisfied in (6).

That a piece of evidence is the “strongest” means that there are no other pieces of evidence that support p more strongly than it does. This captures the fact that ya is usable in both (1) and (2): In (1), S obtains the most reliable evidence (direct, perceptual evidence) as they utter the sentence. Before the utterance, S didn't hold any credible evidence concerning the temperature. In (2), by contrast, S presumably had several types of evidence prior to utterance that it's cold in Canada (e.g., hearsay evidence, inferential evidence based on the geographic location of Canada and the time of year, etc.). But none of them supports p more strongly than the evidence to the possession of which S explicitly commits themselves by uttering (2) (again, sensory evidence that it's cold in Canada). Based on this, we claim that ya expresses two things: **(i)** S commits to the possession of a new piece of evidence e with respect to p , and **(ii)** e is stronger than any other piece of evidence possessed by all of the discourse agents (**NB:** e can be shared by S and other agents; it just has to be the strongest evidence in a strict sense among all the pieces accessible to some agents). We model this based on the Table model [4][5].

Formal Model: The idea that e is the strongest support for p can be modeled based on Bayesian probability semantics [6][7]. That is, e supports a proposition p in c more strongly than e' when the probability of p given e is higher than that of p given e' . That e is the strongest support for p means this holds for all pieces of evidence except e itself possessed by all DAs. This is one of the use conditions for ya . The other condition is e 's newness: The presence of e in S's “evidential base” for p is new. The Table model is suitable to formalize ya 's meaning, as in (7). (We assume each proposition is of type $\langle st, t \rangle$, a downward closed set of sets of possible worlds [4]. Under this assumption, $\bigcup p$ is the set of possible worlds in which p is true; but unless this causes confusion, we keep using p as a conventional shorthand for $\bigcup p$ in body texts).

$$(7) \quad \llbracket ya \rrbracket := \lambda L_{st,tk} . \lambda p_{st,t} . \lambda c_k . [\exists e : \underbrace{[P(\bigcup p|e) \geq \theta \wedge Ev_{s_c, \bigcup p} = Ev_{s_c, \bigcup p}^c \cup \{e\}]}_{(i)} \wedge \underbrace{[\forall e' \in \bigcup_{DA} Ev_{\bigcup p} - \{e\} : P(\bigcup p|e) > P(\bigcup p|e')]}_{(ii)}] \wedge c' = L(p)(c) \text{ in all other respects}]^{c'}$$

(7) dictates that the evidential base for p of S in c (s_c) is updated with e . Thus, s_c is publicly committed to e 's being *new* evidence for p in c' . This models the intuition in **(i)**. The evidence must be newly acquired (thus ya is incompatible with reportatives). The second part formally represents **(ii)**: it creates the grand union of all of the DAs' Evs , and subtracts e from it ($\bigcup_{DA} Ev_p - \{e\}$). It further says that for all pieces of evidence in this set, they are weaker than e as evidence for p . In other words, e counts as the strongest evidence for p . Finally, $P(\bigcup p|e) \geq \theta$ means the truth probability of p conditionalized on the strongest evidence must be above the contextually determined threshold: this precludes cases where e is the strongest evidence but not strong enough for S to assert p . A $\llbracket ya \rrbracket$ utterance ($\llbracket UTT \rrbracket$) yields (cf. [5] for $\llbracket UTT \rrbracket$):

$$(8) \quad \llbracket ya \rrbracket(\llbracket UTT \rrbracket) = \llbracket ya \rrbracket(\lambda p . \lambda c . [\text{MAX}(T) = \{p\} \wedge DC_{s_c} = DC_{s_c}^c \cup \{\bigcup p\}]^{c'}) = \lambda p . \lambda c . [\exists e : [Ev_{s_c, \bigcup p} = Ev_{s_c, \bigcup p}^c \cup \{e\} \wedge \forall e' \in \bigcup_{x \in DA} Ev_{\bigcup p} - \{e\} : P(\bigcup p|e) > P(\bigcup p|e')] \wedge P(\bigcup p|e) \geq \theta] \wedge \text{MAX}(T) = \{p\} \wedge DC_{s_c} = DC_{s_c}^c \cup \{\bigcup p\}]^{c'}$$

The infelicity of ya in (5) follows: Since A has stronger evidence for their disliking sushi (direct knowledge), (7) doesn't hold. It further says that e should trump all other pieces of evidence held by S themselves: this explains why (6) is fine with ya : S has new direct, sensory evidence for p , and this should be stronger than S's any other types of evidence (e.g., hearsay evidence that

sea urchin sushi tastes bad). Note that the condition only says S holds the strongest evidence e : it does not preclude the possibility that e is shared between S and other agents: This explains examples such as (2), when p - ya is uttered in the presence of A , who holds the same factive, perceptual evidence (e). Hence **NB** explained. Furthermore, (7) provides an explanation for the fact that ya is often used in self-talk: If there are no agents other than S , only Ev_{s_c} pertains to the use condition in (7). As long as e is stronger than S 's any other evidence, ya can be used. Thus, ya 's use condition is loose in such cases. This explains why it is often used in self-talk [1][2]. Note though that [2, pp. 45–46] shows that ya can be used to answer a question asked by the addressee, which means “self-talk” is not ya 's function, as in (9). This is compatible with our proposal: S doesn't feel like having sushi now, which is the new strongest evidence for p . **Advantages:** This presents a fine-grained formal semantics of ya with greater empirical adequacy than previous accounts. One further virtue is that (7) gives us an explanation for the fact that ya is unusable in interrogatives: ya is illicit in (9), as the grand union of the polar interrogative $p?$ is the union of p -worlds and $\neg p$ -worlds, which is the entire logical possibility $\mathcal{W}$. Since $P(\bigcup p|e) = P(\mathcal{W}|e) = P(\bigcup p|e') = P(\mathcal{W}|e') = 1$, the condition (ii) in (7) cannot be satisfied, rendering the infelicity of ya in interrogatives.

(9) A: Let's have some sushi! – S: Ee (well), ii (fine) ya . ‘Well, no.’

(10) Kanada-wa samui ($\#ya$)? ‘Is it cold in Canada?’

References: [1] Nakazaki, T. 2014. *Syuuzitu Hyoogen Bunka* 8. / [2] Oki, K. 2017. Tokyo, Hitsuji Shobo. / [3] Shirakawa, R. 2021. PhD diss., the U. of Tsukuba. / [4] Farkas, D. F. and F. Roelofsen. 2017. *J. of Sem.* 34. / [5] Rudin, D. 2022. *J. of Sem.* 39. / [6] Davis, C., C. Potts and M. Speas. 2007. In *Proc. of SALT XVII*. / [7] Yalcin, S. 2012. *Proc. Arist. Soc.* 112.