

Derive the biased reading of *A-not-A* questions in Mandarin

Yaxuan Wang*

Abstract. Among many forms of *A-not-A* questions in Mandarin Chinese, the *shi-not-shi* question is found to be unique in that it obligatorily gives rise to a biased reading toward its prejacent, so-called positive bias. The previous pragmatic approach by Ye (2020) establishes a link between focus in polar questions and question bias to explain this biased reading. However, the current study finds that two other *A-not-A* questions formed by epistemic modals, *hui-not-hui* and *keneng-not-keneng* which are not focus markers, obligatorily produce positively biased readings as well. I propose that biased *A-not-A* questions are a type of high-negation questions with *A-not-A* residing outside of TP. *shi*, *hui* and *keneng* should all be analyzed as epistemic modals which are the overt realization of Goodhue’s (2019) epistemic operator scoped by the negator. The positively biased reading is derived from the resulting unbalanced partition based on general pragmatic principles. This analysis from the semantic aspect provides new evidence for the argument that the first *A* has reality only in PF. Furthermore, the Mandarin Chinese data lends evidence to Goodhue’s (2019) argument that there exists a doxastic speech operator between NegP and TP in high negation questions. The paper also provides explanations for previously remaining questions on bias cancelation by the stress marker *daodi* and factive predicates like *zhidao* (“know”).

Keywords. biased questions; polar questions; high negation questions

1. Introduction. In Mandarin Chinese, multiple items can form *A-not-A* polar questions. For example, *A* can be replaced by verbs (1a), prepositions (1b), adverbs (1c), epistemic modals (1d), deontic modals (1e), and the word *shi* (1f).

- (1) a. Lisi xihuan-bu-xihuan huahua?
Lisi like-not-like painting
Does Lisi like painting?
- b. Lisi gei-mei-gei ni dadianhua?
Lisi to-not-to you call
Did Lisi call you?
- c. Lisi chang-bu-chang chi pingguo?
Lisi often-not-often eat apple
Does Lisi often eat apples?
- d. Lisi hui-bu-hui shengbing le?
Lisi possible-not-possible be.ill LE
Is it possible that Lisi is ill?

* I would like to thank my advisor Brian Buccola, my classmate Jingying Xu in the semantics course at MSU, and the audiences at IACL 2023 and LSA 2024 for their valuable suggestions.

Author: Yaxuan Wang, Michigan State University (wangyax3@msu.edu).

- e. Lisi yinggai-bu-yinggai xie zuoye?
Lisi should-not-should write homework
Should Lisi do homework?
- f. Lisi shi-bu-shi xihuan huahua?
Lisi SHI-not-SHI like painting
Does Lisi like painting?

It was observed that among various forms of *A-not-A* questions, the *shi-not-shi* question is rather unique in that it entails the questioners bias toward the answer entailed by the prejacent of *shi-not-shi*, while other *A-not-A* questions do not. As shown by (2), Zhangsan as the questioner obtained the clue that Lisi likes painting. To find out if his guess is correct, Zhangsan can use either the *shi-not-shi* or the *V-not-V* polar question to ask Wangwu. But it is only under the *shi-not-shi* question can Wangwu felicitously responds with a confirmative expression. This example indicates that the *shi-not-shi* but not *V-not-V* question entails the questioners bias toward the truth of its prejacent.

(2) Context: Zhangsan saw Lisi bought many paintbrushes yesterday. Wangwu knows Lisi very well. Today, Zhangsan asks Wangwu:

- a. Lisi shi-bu-shi xihuan huahua?
Lisi SHI-not-SHI like painting
Is it painting that Lisi likes?
Wangwu replies: How do you know that!
- b. Lisi xihuan-bu-xihuan huahua?
Lisi like-not-like painting
Does Lisi like painting?
Wangwu replies: #How do you know that!

However, I find that the *shi-not-shi* is not the only type of biased *A-not-A* question. Under the same context as described in (2), the confirmative response is also felicitous under both (3a) and (3b), which suggests that the *A-not-A* questions formed by Mandarin epistemic modals, *hui* ('probably') and *keneng* ('probably')¹, are positively biased as well.

(3) Zhangsan asks Wangwu:

- a. Lisi hui-bu-hui xihuan huahua?
Lisi probably-not-probably like painting
Is it probable that Lisi likes painting?
Wangwu replies: How do you know that!
- b. Lisi keneng-bu-keneng xihuan huahua?
Lisi probably-not-probably like painting
Is it probable that Lisi likes painting?
Wangwu replies: How do you know that!

¹ Both *hui* and *keneng* mean "probably". The future work needs to figure out whether one is semantically stronger than the other.

The current paper provides a unified analysis that explains the facts mentioned above by treating Mandarin *A-not-A* questions as a type of high-negation question and adopting Goodhue’s (2019) epistemic operator and deriving the positively biased reading from the resulted unbalanced partitions.

In addition, there are two situations where the positive bias is canceled. First, as Ye (2020) points out, when biased *A-not-A* has an intonational stress on the first *A* under the scope of the optional stress marker *daodi*, the bias is canceled and the question becomes neutral. (4a) shows that the stressed question is felicitous to use under the context where the questioner has equal amounts of prior expectations for the positive and the negative answer. But if the questioner lacks prior evidence for either answer, stressed questions are never licensed, as illustrated by (4b). To explain this, I propose that *daodi* is an exhaustive partition operator with the semantics of partitioning a set of worlds into sets of balanced cells.

- (4) a. Zhangsan thought Lisi likes painting. But Lisi dropped the painting class yesterday. Now Zhangsan is not sure if Lisi likes painting or not. Zhangsan asks Wangwu:
 Lisi (daodi) {shi-bu-shi/hui-bu-hui/keneng-bu-keneng}_F xihuan huahua?
 Lisi at.all {SHI-not-SHI/probably-not-probably} like painting
 Does Lisi like painting at all?
 Wangwu replies: Lisi likes painting. /#You are right.
- b. Zhangsan has no idea if Lisi likes painting or not. Zhangsan asks Wangwu:
 #Lisi (daodi) {shi-bu-shi/hui-bu-hui/keneng-bu-keneng}_F xihuan huahua?
 Lisi at.all {SHI-not-SHI/probably-not-probably} like painting
 Does Lisi like painting at all?’

Secondly, Ye (2020) leaves a puzzle that the positive bias is concealed when the *A-not-A* question is embedded under a responsive predicate but not when under a rogative predicate, as shown by (5). I argue *know* is a factive verb that selects DPs as its complement. To turn the question into a DP, there is an implicit THE ANSWER TO operator which partitions a set of worlds into sets of balanced cells.

- (5) a. Zhangsan zhidao Lisi shi-bu-shi xihuan huahua
 Zhangsan know Lisi SHI-not-SHI like painting
 Zhangsan knows if Lisi likes painting.
 ⇨ Zhangsan is neutral about whether Lisi likes painting or not.
- b. Zhangsan xiangzhidao Lisi shi-bu-shi xihuan huahua
 Zhangsan wonder Lisi SHI-not-SHI like painting
 Zhangsan wonders if Lisi likes painting.
 ⇨ Zhangsan has a prior bias to believe that Lisi likes painting.

In the remainder of the paper, Section 2 argues against Ye’s (2020) focus account. Section 3 displays the syntactic and pragmatic properties of *A-not-A* questions and high negation questions in English. I argue that the Mandarin biased *A-not-A* question is a type of high negation question. Section 4 reviews previous accounts for high negations and discusses their application to Mandarin. Section 5 shows how Goodhue’s (2019) operator and pragmatic principles can successfully

apply to the Mandarin data. This section also provides explanations for bias cancelations. Section 6 discusses the implications of the analysis and concludes the paper.

2. Against the focus account. Ye (2020) proposes that in the *shi-not-shi* question, *shi* is a focus marker which presupposes its prejacent *p* is a possible complete answer to the current Question Under Discussion (QUD), that is, no other true answer entails *p*. *shi-not-shi* questions are treated as the same focus marker, so the same presupposition projects out of *shi-not-shi* questions. *shi-not-shi* questions are used to check the truth of possible complete answers to the current QUD, so-called F(ocus)-strategy of inquiry, while other *A-not-A* questions are used to check the truth of possible partial answers, so-called C(ontrastive)T(opic)-strategy of inquiry. The difference between the two strategies is that when using F-strategy, the *yes* answer to the polar question closes the QUD; while using CT-strategy, the QUD will not be closed until the questioner has gone through all the subquestions. This can be illustrated by the following example in (6).

(6) QUD: Which branches of linguistics does Lisi like?

a. Is it pragmatics_F that Lisi likes?

Answer: Yes (close the QUD)

b. CT-strategy of inquiry:

Subquestion: Does Lisi like pragmatics?

Answer: Yes (move to the next subquestion, such as *How about syntax?*)

Ye further assumes that the F-strategy of inquiry indicates the questioners intention to close the current QUD, and to achieve this goal, they expect a positive answer. As *shi-not-shi* question is part of the F-strategy, it is expected to bias toward a positive answer. In short, Ye's approach establishes a link between focus in polar questions and question bias and predicts that *shi-not-shi* questions are forced to have a positively biased reading.

However, based on a simple survey of native Mandarin speakers that I conducted, *shi-not-shi* questions without salient focus are still felicitous to induce a positively biased reading while followed by other subquestions, as shown by (7). This suggests that *shi-not-shi* does not necessarily presuppose a possible complete answer set. Thus, to assume *shi-not-shi* questions as part of the F-strategy of inquiry seems dubious.

(7) QUD: Which branches of linguistics does Lisi like?

a. Zhangsan:

Lisi shi-bu-shi xihuan yuyongxue?

Lisi SHI-not-SHI like pragmatics

Is it pragmatics that Lisi likes?

Wangwu replies: How do you know that!

b. Cont. Zhangsan:

Subquestion: How about syntax?

Wangwu: Yes, that's also his interest.

Furthermore, it is unclear how Ye's proposal works for the current new observation of *hui-not-hui* and *keneng-not-keneng* biased questions as they do not function as focus markers. As (8a) illustrates, the focused constituent is required to appear within the scope of *shi*, otherwise, the

sentence is unacceptable. Yet, the focus is felicitous when appears above the epistemic modals² in (8b).

- (8) a. #Lisi_F shi xihuan huahua.
 Lisi FOCUS like painting
 Int.: It is Lisi who likes painting.
- b. Lisi_F keneng/hui xihuan huahua.
 Lisi probably like painting
 It is Lisi who probably likes painting.

In fact, Mandarin *shi* can be used in more than one way. Besides the focus marker function mentioned above, it is also widely accepted to be a copula. However, the copula meaning also fails to apply to the analysis of *shi* in *shi-not-shi* questions. (9) shows that in a situation where Zhangsan lacks prior bias toward the prejacent, the COPULA-*not*-COPULA question is felicitous to use.

- (9) Context: Zhangsan doesn't know Lisi's identity, but Wangwu knows Lisi very well. Today, Wangwu asks Zhangsan to guess Lisi's occupation. Zhangsan randomly guesses:
 Lisi shi-bu-shi yi wei huajia?
 Lisi SHI-not-SHI one CL painter
 Is Lisi a painter?

I propose to treat *shi* in the positively biased *shi-not-shi* question as an epistemic modal. It is widely agreed that epistemic modals vary in the degree of certainty toward the proposition in its scope, based on the amount of evidence accumulated within the epistemic state in support of that proposition. For example, in English, *Q believes For a fact p* if and only if (iff) *Q* has direct evidence for *p*; *Q believes Must p* iff *Q* has at least indirect evidence for *p* and no evidence against *p*; *Q believes Probably p* iff *Q* has much more evidence for *p* than against *p*; etc., $\langle \text{for a fact, must, probably, ...} \rangle$. Similarly, Mandarin Chinese has a set of epistemic modals scaled on the same basis, in which *shi* is the strongest one for a speaker to use if they have direct evidence: $\langle \text{shi, hui/keneng, ...} \rangle$. As (10) shows, when Zhangsan has direct evidence, it is natural to use *shi* instead of a less probable *hui* or *keneng*. The following section displays more evidence for this argument by showing that *shi* both syntactically and pragmatically resembles canonical epistemic modals when forming biased *A-not-A* questions.

- (10) Zhangsan heard Lisi say : "I like painting!" Then Zhangsan met Wangwu and chatted about Lisi:
- a. Wo juede Lisi shi xihuan huahua.
 I believe Lisi SHI like painting
 I believe it is a fact that Lisi likes painting.

² *hui* as an epistemic modal seems to be less natural when used in a positive declarative sentence than in its negative counterpart and polar interrogatives. This is independent of the current discussion, so I leave the question for the next step of research.

- b. #Wo juede Lisi keneng/hui xihuan huahua.
 I believe Lisi probably like painting
 I believe Lisi probably likes painting.

3. Compare properties of *A-not-A* question to high negation question. In the literature, biased questions are observed in English and many other languages such as Greek, Spanish, German, and Korean as well (e.g., Romero & Han 2004; Sudo 2013; AnderBois 2019). In this section, I take the English high- and low-negation polar question pair as an example to compare with Mandarin *A-not-A* questions in terms of their syntactic and pragmatic properties.

3.1. SYNTACTIC PROPERTIES. Previous studies argue that the biased *A-not-A* is triggered by an outer *A-not-A* above TP, while the neutral counterpart is licensed by an inner *A-not-A* morpheme situated below TP (e.g., Law 2006; Liu 2004, 2010; Tsai & Yang 2015). Law (2006) proposes an adverbial test to evidence this argument. (11a) shows the neutral *A-not-A*, here formed by *V-not-V*, is typically blocked by sentential frequency adverbials like *often*, whereas *shi-not-shi* is not (11b-c). Observe that *hui-not-hui* and *keneng-not-keneng* pass the tests as well.

- (11) a. *Lisi kan-bu-kan jingchang dianshi?
 Lisi watch-not-watch often TV
 Int.: Does Lisi often watch TV?
- b. Lisi {shi-bu-shi/hui-bu-hui/ke-bu-keneng} jingchang kan dianshi?
 Lisi {SHI-not-SHI/probably-not-probably} often watch TV
 Is it the case that Lisi often watches TV?
- c. *Lisi jingchang {shi-bu-shi/hui-bu-hui/ke-bu-keneng} kan dianshi?
 Lisi often {SHI-not-SHI/probably-not-probably} watch TV
 Int.: Is it the case that Lisi often watches TV?

Goodhue (2019) proposes several tests for English to support the argument that the negation operator in high negation questions is above the matrix TP as well. For example, the word *again* triggers the presupposition that the proposition denoted by its complement has happened before. When negation is within the prejacent of *again*, negation is part of the presupposition (12a). Applying *again* to high negation questions as in (12b), we find the question does not presuppose the negative proposition. *again* here is in the scope of the negator, and its presupposition projects out of it. Another example, *until*-adverbial only combines with clauses that have durative rather than punctual aspect (13a), and negating a verb with punctual aspect creates durative aspect (13b-c). However, high negation questions do not (13d). Both tests show that *again* and *until*-adverbial cannot scope above negation in high negation questions, indicating the negator is outside of TP.

- (12) a. Did John not come to class again?
 ~↗ John did not come to class before.
- b. Didn't John come to class again?
 ~↗ John came to class before.

- (13) a. #Lily discovered the thief until 9.
 b. Lily didn't discover the thief until 9.
 c. Did Lily not discover the thief until 9?
 d. #Didn't Lily discover the thief until 9?

3.2. PRAGMATIC PROPERTIES. As discussed above, Mandarin outer *A-not-A* questions are biased in that they necessarily require the speaker to have a prior expectation that the proposition embedded is true. By contrast, inner *A-not-A*s do not have such a requirement. English negation questions parallel the same property: English high negation questions necessarily convey that the speaker is epistemically biased, but low negation questions do not have biased readings. As shown by (14) taken from Goodhue (2019), in the context of (14a) where the questioner Lucy lacks expectations toward Kate's being home or not, the low negation question is perfect but the high negation question is unacceptable; while in the context of (14b) where Lucy has a prior expectation that the proposition embedded under negation is true, the high negation question becomes felicitous.

- (14) a. Lucy is looking for her roommate Kate. She has no idea whether Kate is home or not. She looks around and can't find her. But Lucy sees Grace, so she asks Grace:
 (i) Is Kate not home?
 (ii) #Isn't Kate home?
 b. Lucy is looking for her roommate Kate. She expects Kate to be home. But Lucy can't find her anywhere. Lucy asks Grace:
 (i) Is Kate not home?
 (ii) Isn't Kate home?

Given the similarity between *A-not-A* questions and negation questions, I propose to analyze Mandarin *A-not-A* questions as a type of negation question that have been cross-linguistically observed, outer *A-not-A* questions resembling high negation questions, and inner *A-not-A* questions resembling low negation questions.

4. Previous accounts for high negation questions. To analyze high negation questions, previous theories propose an operator scoped by the negator in high negation questions. The meaning defined for such an operator differs among various proposals.

Romero & Han (2004) observe that both high negation questions and questions with polarity focus convey the speaker's epistemic bias, as shown by (15). They thus relate the two questions and argue that high negation scopes above an epistemic, conversational operator in the form of a VERUM focus, which can be overtly spelled out with the English adverb *really* and relates to the strength with which a proposition should be added to the common ground. VERUM(*p*) reads as *it is for sure that p should be added to the common ground*. Following Hamblin's set denotation for questions, with the application of VERUM to high negation questions, we get an unbalanced partition, $\{\neg\text{VERUM}(p), \text{VERUM}(p)\}$. $\neg\text{VERUM}(p)$ represents a larger set of worlds in which *it is not for sure that p should be added to the CG* or *it is for sure that p shouldn't be added to the CG*. The questioner using high negation questions to asks whether *p* can be assumed with a high degree of certainty, thus, a pragmatic bias toward *p*. It is also argued that VERUM can apply above

the negator, thus creating an ambiguous high negation question. However, this theory does not apply to Mandarin *A-not-A* questions. As (4) in Section 1 illustrates, when *A-not-A* is focused, the question does not obtain the biased reading anymore. Relating outer *A-not-A* to polarity focus thus seems dubious. Moreover, Goodhue (2019) points out that polarity focus but not high negation requires the sort of antecedent needed to license prosodic prominence shifts more generally. A unified account of high negation and polarity may not be appropriate even for other languages.

- (15) a. Isn't Kate home?
 $\rightsquigarrow$ The speaker believes that Kate is home.
 b. Is_F Kate home?
 $\rightsquigarrow$ The speaker believes that Kate is not home.

Repp (2013) is against Romero & Han (2004) and argues that high negator is a conversational epistemic operator itself, supported by the evidence for the parallel behavior of high negation and negation in denials concerning their syntax and scope behavior in negative gapping sentences. She assumes high negation expresses an additional operator *FALSUM* that states the degree of strength with which the proposition should be added to the common ground is zero. This operator allows answers with a reduced degree of certainty such as *Yes, I think there is a station, but I'm not sure*. While applying this account to Mandarin can account for the biased reading of outer *A-not-A*s, since it is the negator that plays an important role in deriving the bias, it remains unclear why only epistemic modals form outer *A-not-A*s with the negator.

By contrast, Goodhue (2019) proposes high negation scopes over an epistemic operator as defined in (16). He further develops a pragmatic analysis based on the general pragmatic principles in (17-18).

- (16) $\llbracket \textit{Epis} \rrbracket = \lambda p. \lambda w. \forall w' \in \textit{Dox}_x(w) [p(w') = 1]$
 $\textit{Dox}_x(w)$ is the set of worlds compatible with x 's beliefs in w . x is a free variable for individuals whose value is contextually determined. In high-negation questions, x is the addressee. In plain English, (15) simply states that the addressee believes p .

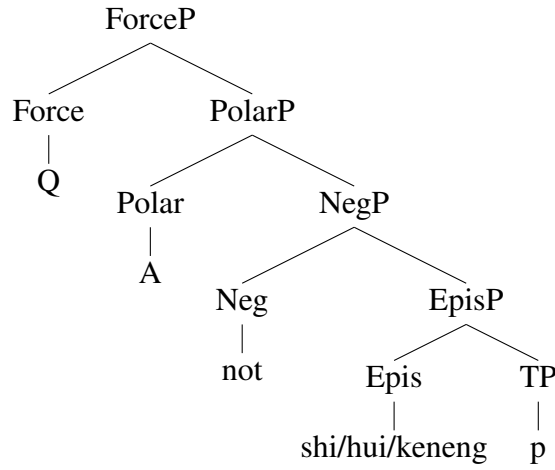
- (17) a. Felicity Condition for the use of questions:
 A question Q is felicitous only if Q is at least as useful as other questions Q' .
 b. The lemma:
 Outer *A-not-A* questions are at least as useful as other Q s only if the speaker is biased for p .
 c. Conclusion:
 Outer *A-not-A* questions are felicitous only if the speaker is biased for p .

- (18) Strategies for comparing the utility of questions
 a. Gain information strategy: Q_1 is more useful than Q_2 iff Q_1 partitions cells produce epistemic states that are more informed relative to p than the cells of Q_2 do.
 b. Determine agreement strategy: Q_1 is more useful than Q_2 iff Q_1 partitions cells make it easier to determine whether the addressee agrees with the speaker about p than the cells of Q_1 do.

Since Mandarin only allows for epistemic modals to form an outer *A-not-A*, plus the parallel between *A-not-A* and high negation questions, I propose that the epistemic operator in (16) is spelled out with the Mandarin epistemic modals, and the biased reading of *A-not-A* questions should be derived in the same way as for negation questions. In the next section, I show how the positive bias of *A-not-A* questions is obtained based on Goodhue’s theory.

5. Analysis. The syntactic structure for outer *A-not-A* questions is displayed below. I assume that the first *A* in the PolarP head has only phonological reality, thus, I treat the denotation of it as an identity function. Note that it is plausible that the second *A* in the EpisP head lacks realization in LF, which I will discuss later. Moreover, I assume the denotation for *Q* is a Hamblin’s set (19). The semantics of the outer *A-not-A* question is computed in (20).

$$(19) \quad \llbracket Q \rrbracket = \lambda p. \lambda q. [q = p \vee q = \neg p]$$



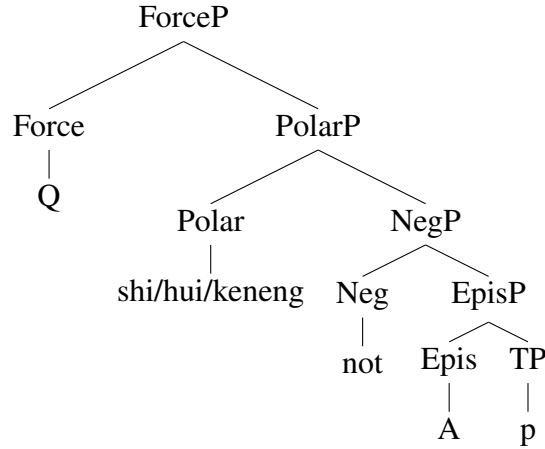
- (20) a. $\llbracket EpisP \rrbracket = \llbracket Epis \rrbracket(\llbracket p \rrbracket) = \lambda w. \forall w' \in Dox_x(w) [p(w') = 1]$
b. $\llbracket NegP \rrbracket = \llbracket not \rrbracket(\llbracket EpisP \rrbracket) = \lambda w. \neg \forall w' \in Dox_x(w) [p(w') = 1]$
c. $\llbracket PolarP \rrbracket = \llbracket A \rrbracket(\llbracket NegP \rrbracket) = \lambda w. \neg \forall w' \in Dox_x(w) [p(w') = 1]$
d. $\llbracket ForceP \rrbracket = \llbracket Q \rrbracket(\llbracket PolarP \rrbracket) = \lambda q. [q = \lambda w. \neg \forall w' \in Dox_x(w) [p(w') = 1] \vee q = \lambda w. \forall w' \in Dox_x(w) [p(w') = 1]]$

In plain English, the resulting answer set in (20d) is $\{the\ addresssee\ does\ not\ believe\ p, the\ addresssee\ believes\ p\}$, schematically $\{\neg \Box p, \Box p\}$. This is an unbalanced partition in that *not believing p* includes a wider range of situations where the addressee lacks belief either way or the addressee believes not *p*.

If the questioner lacks belief about *p* either way, their goal will be to gain information about *p* from the addressee, thus adopting the gain information strategy (18a). Either *p* or $\neg p$ would be a perfect answer to increase the questioners information. Therefore, under such situations, the questioner would prefer to ask a question with balanced partition $\{p, \neg p\}$, such as the neutral polar question with an inner *A-not-A* structure, *Lisi like-not-like painting?* On the other hand, if the questioner has a prior expectation toward one answer, their goal will be to determine whether or not they and the addressee agree on *p*, thus adopting the determine agreement strategy (18b). If the questioner biases toward *p*, that *the addressee believes p* will agree with them, while that *the*

addressee does not believe p will indicate disagreement. In this case, the questioner would rather ask a question with an unbalanced partition $\{\Box p, \neg\Box p\}$. By contrast, if the questioner biases toward $\neg p$, that *the addressee believes* $\neg p$ will agree with them, while that *the addressee does not believe* $\neg p$ will indicate disagreement. Again, the questioner would prefer to ask a question with an unbalanced partition $\{\Box\neg p, \neg\Box\neg p\}$. This explains why outer *A-not-A* questions always have a biased reading toward its prejacent proposition.

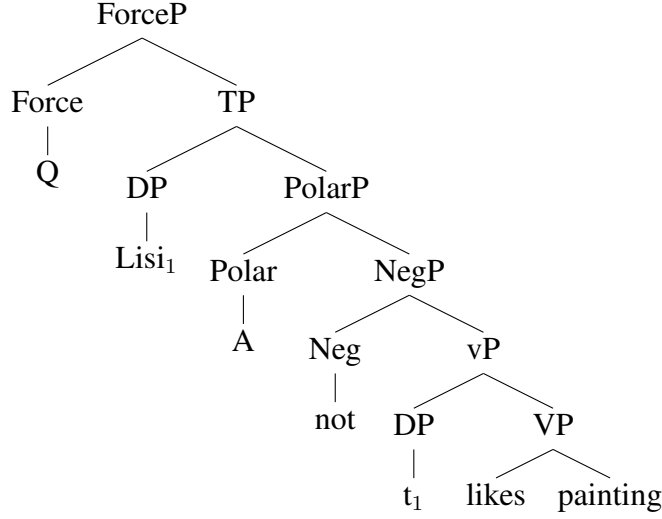
As mentioned earlier, it is plausible that the first *A* has reality in LF while the second *A* is realized only in PF, as shown by the structure below. In such cases, we obtain the resulting answer set $\{the\ addressee\ believes\ not\ p, the\ addressee\ does\ not\ believe\ not\ p\}$ (21), schematically, $\{\Box\neg p, \neg\Box\neg p\}$.



- (21) a. $\llbracket EpisP \rrbracket = \llbracket A \rrbracket(\llbracket p \rrbracket) = \lambda w.p(w) = 1$
b. $\llbracket NegP \rrbracket = \llbracket not \rrbracket(\llbracket EpisP \rrbracket) = \lambda w.\neg p(w) = 1$
c. $\llbracket PolarP \rrbracket = \llbracket shi/hui/keneng \rrbracket(\llbracket NegP \rrbracket) = \lambda w.\forall w' \in Dox_x(w)[\neg p(w') = 1]$
d. $\llbracket ForceP \rrbracket = \llbracket Q \rrbracket(\llbracket PolarP \rrbracket) = \lambda q.[q = \lambda w.\forall w' \in Dox_x(w)[\neg p(w') = 1] \vee q = \lambda w.\neg\forall w' \in Dox_x(w)[\neg p(w') = 1]]$

As discussed above, according to the determine agreement strategy, the questioner only uses questions with such a partition when they are biased toward the negative answer. However, a negative bias toward *p* is not attested in Mandarin outer *A-not-A* questions with the form of *A-not-A p*. It is only attested when the prejacent is negated, *A-not-A* $\neg p$. Therefore, I conclude that the high position *A* only has reality in PF while the lower position *A* is realized in LF. This analysis from the semantic aspect thus provides new evidence for the argument that the first *A* is realized by phonological processes such as reduplication or syllable copy and deletion. Furthermore, the Mandarin Chinese data lends additional language evidence to Goodhue's 2019 argument that there exists an epistemic operator between NegP and TP in high negation questions and this structure is not ambiguous.

Since no epistemic operator is available for inner *A-not-A* questions, whose structure is presented below, the resulting partition of the question is a balanced one by default (22). The inner *A-not-A* questions are predicted to have the same answer set as that of canonical polar questions. Pragmatically, since the questioner lacks belief either way, the gain information strategy requires the most informative answer set $\{p, \neg p\}$.



- (22) a. $\llbracket vP \rrbracket = \lambda w.t_1 \text{ likes painting in } w$
 b. $\llbracket NegP \rrbracket = \llbracket not \rrbracket(\llbracket vP \rrbracket) = \lambda w.t_1 \text{ doesn't like painting in } w$
 c. $\llbracket PolarP \rrbracket = \llbracket A \rrbracket(\llbracket NegP \rrbracket) = \lambda w.t_1 \text{ doesn't like painting in } w$
 d. $\llbracket TP \rrbracket = \llbracket PolarP \rrbracket(\llbracket Lisi_1 \rrbracket) = \lambda w.Lisi \text{ doesn't like painting in } w$
 e. $\llbracket ForceP \rrbracket = \llbracket Q \rrbracket(\llbracket TP \rrbracket) = \lambda q.[q = \lambda w.Lisi \text{ doesn't like painting in } w \vee q = \lambda w.Lisi \text{ likes painting in } w]$

Now let's turn to the two bias cancelation cases. First, when the outer *A-not-A* is stressed with an optional stress marker *daodi*, the speaker's prior positive bias toward the prejacent is canceled, and they become neutral toward the answer. The example is repeated in (23). Notice that the prior positive bias is a presupposition for using the stressed *A-not-A* question, as it is infelicitous to ask such a question when the context does not entail the initial positive bias (23b).

- (23) a. Zhangsan thought Lisi likes painting. But Lisi dropped the painting class yesterday. Now Zhangsan is not sure if Lisi likes painting or not. Zhangsan asks Wangwu:
 Lisi (daodi) {shi-bu-shi/hui-bu-hui/keneng-bu-keneng}_F xihuan huahua?
 Lisi at.all {SHI-not-SHI/probably-not-probably} like painting
 Does Lisi like painting at all?
 Wangwu replies: Lisi likes painting. /#You are right.
 b. Zhangsan has no idea if Lisi likes painting or not. Zhangsan asks Wangwu:
 #Lisi (daodi) {shi-bu-shi/hui-bu-hui/keneng-bu-keneng}_F xihuan huahua?
 Lisi at.all {SHI-not-SHI/probably-not-probably} like painting
 Does Lisi like painting at all?'

Thus, I treat *daodi* as an exhaustive partition operator that takes a partition as input and produces a balanced partition. It presupposes the input partition is unbalanced. The definition is given in (23).

- (24) $\llbracket daodi \rrbracket = \lambda C_{QUD} : \exists c_1, c_2 \in C[|c_1| \neq |c_2|]. C_{QUD}$ such that $\forall c_1, c_2 \in C[|c_1| = |c_2|]$
 C_{QUD} is a function that partitions a question under discussion into non-overlapping sets of worlds, so-called cells.

This proposal is endorsed by independent evidence from *wh*-questions. For instance, in (25), Zhangsan's first response is blurred by hiding the other two detailed answers, that *Zhangsan had hotpot* and that *Zhangsan ate ice cream*. The doctor believed that Zhangsan didn't tell the whole story because her question "what did you eat?" can denote an unbalanced answer set. To gain more information, the doctor needs a complete and detailed answer, thus, she continued with a *daodi*-question that requires an obligatorily balanced partition.

- (25) Context: Zhangsan had hotpot, cakes, and ice cream. After a few minutes, the pain started in his stomach. He then went to the hospital. The doctor asked: What did you eat? Zhangsan replied: Cakes and something else. The doctor was not satisfied with this answer, so she continued:
 Ni daodi chi le shenme?
 you at.all eat LE what
 What on earth did you eat?
 Zhangsan replied: OK... I had cakes, hotpot, and ice cream.

Finally, Ye (2020) leaves a puzzle that the bias of outer *A-not-A* is survived when embedded under rogative verbs like *xiangzhidao* ('wonder') but concealed when embedded under responsive verbs such as *zhidao* ('know'), repeated in (26).

- (26) a. Zhangsan zhidao Lisi shi-bu-shi xihuan huahua
 Zhangsan know Lisi SHI-not-SHI like painting
 Zhangsan knows if Lisi likes painting.
 $\rightsquigarrow$ Zhangsan is neutral about whether Lisi likes painting or not.
 b. Zhangsan xiangzhidao Lisi shi-bu-shi xihuan huahua
 Zhangsan wonder Lisi SHI-not-SHI like painting
 Zhangsan wonders if Lisi likes painting.
 $\rightsquigarrow$ Zhangsan has a prior bias to believe that Lisi likes painting.

Notice there is a parallel pattern in terms of positive bias inferences and presuppositions. (27) shows that *wh*-questions embedded under *zhidao* ('know') have existential presupposition projected to the matrix level, while the no existential presupposition projects when embedded under *xiangzhidao* ('wonder').

- (27) a. Zhangsan zhidao shui qu can hui le.
 Zhangsan know who go join meeting LE
 Zhangsan knows who went to the meeting.
 Presuppose: Someone went to the meeting.
 b. Zhangsan xiangzhidao shui qu can hui le.
 Zhangsan wonder who go join meeting LE
 Zhangsan wonders who went to the meeting.
 Presuppose: Someone went to the meeting.

Based on the comparison above, I propose it is the contrast between being factive or non-factive matters, instead of being responsive or rogative. According to Kastner (2015), the complements of factive attitudes have a definite DP-like behavior, while complements of non-factives do not necessarily. I assume that factives can embed DPs via a silent THE ANSWER TO operator, as shown in (28).

(28) Zhangsan knows [the answer to] who went to the meeting.

I further assume the semantics of THE ANSWER TO operator is similar to that of *daodi* which takes a partition as input and produces a balanced partition, but it does not have any presupposition requirement (29). This proposal is evidenced by examples like (30): when *the answer to* is overtly stated in the response, we expect a complete and exhaustive answer instead of a vague one.

(29) $\llbracket \text{the answer to} \rrbracket = \lambda C_{QUD}. C_{QUD}$ such that $\forall c_1, c_2 \in C [|c_1| = |c_2|]$

(30) Context: Lisi asks Wangwu: “*What did Zhangsan eat last night?*” Wangwu replies:

- a. “The answer to your question is that he had hotpot, cakes, and ice cream.”
- b. ? “The answer to your question is that he had cakes and something else that I’m not sure about.”

6. Conclusion. This paper provides a unified analysis for Mandarin positively biased *A-not-A* questions. It is observed that only those formed by *shi-not-shi*, *hui-not-hui*, and *keneng-not-keneng* can derive positive bias. I argue that the previous focus theory is not able to account for the data and propose to analyze these three *A-not-A*s are epistemic modals. Based on the similarity shared by Mandarin *A-not-A* questions and English negation questions in syntactic and pragmatic properties, I argue that *A-not-A* question is a type of high negation question. I further adopt Goodhue’s epistemic operator and pragmatic principles to derive the positively biased reading based on the resulting unbalanced partition. This analysis from the semantic aspect provides new evidence for the argument that the first *A* has reality only in PF. Furthermore, the Mandarin Chinese data lends evidence to Goodhue’s (2019) argument that there exists an epistemic operator between NegP and TP in high negation questions. In addition, I argue that the positive bias can be canceled by exhaustive partition operators. In stressed *A-not-A*, the operator is surfaced as the stress marker *daodi*, while in factive embeddings, the operator is surfaced as the DP convertor THE ANSWER TO.

References

- AnderBois, Scott. 2019. Negation, alternatives, and negative polar questions in American English. In Klaus von Stechow, V. Edgar Onea Gaspar & Malte Zimmermann (eds.), *Questions in discourse*, 118–171. Leiden: Brill. https://doi.org/10.1163/9789004378308_004.
- Goodhue, Daniel. 2019. High negation questions and epistemic bias. *Proceedings of Sinn Und Bedeutung* 23(1). 469–485. <https://doi.org/10.18148/sub/2019.v23i1.544>.
- Kastner, Itamar. 2015. Factivity mirrors interpretation: The selectional requirements of presuppositional verbs. *Lingua* 164. 156–188. <https://doi.org/10.1016/j.lingua.2015.06.004>.

- Law, Paul. 2006. Adverbs in *A-not-A* questions in Mandarin Chinese. *Journal of East Asian Linguistics* 15(2). 97–136. <https://doi.org/10.1007/s10831-005-4916-5>.
- Liu, Haiyong. 2004. *Complex predicates in Mandarin Chinese: Three types of bu-yu structures*. Los Angeles: UCLA dissertation.
- Liu, Haiyong. 2010. What is *A* in Mandarin *A-not-A* questions? *Proceedings of the 22nd North American Conference on Chinese Linguistics (NACCL- 22) and the 18th International Conference on Chinese Linguistics (IACL-18)* 2. 287–304.
- Repp, Sophie. 2013. Common ground management: Modal particles, illocutionary negation and verum. In Daniel Gutzmann & Hans-Martin Gärtner (eds.), *Beyond expressives: Explorations in use-conditional meaning*, 231–274. Leiden: Brill. https://doi.org/10.1163/9789004183988_008.
- Romero, Maribel & Chung-hye Han. 2004. On negative yes/no questions. *Linguistics and Philosophy* 27(5). 609–658. <https://doi.org/10.1023/b:ling.0000033850.15705.94>.
- Sudo, Yasutada. 2013. Biased polar questions in English and Japanese. In In Daniel Gutzmann & Hans-Martin Gärtner (eds.), *Beyond expressives: Explorations in use-conditional meaning*, 275–295. Leiden: Brill. https://doi.org/10.1163/9789004183988_009.
- Tsai, Wei-tien Dylan & Ching-yu Helen Yang. 2015. Inner vs. outer *A-not-A* questions. Paper presented at International Workshop on Cartography of Syntax. Beijing Language and Culture University.
- Ye, Shumian. 2020. From maximality to bias: Biased *A-not-A* questions in Mandarin Chinese. *Semantics and Linguistic Theory (SALT)* 30. 355–375. <https://doi.org/10.3765/salt.v30i0.4826>.