

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Modeling atypicality inferences in pragmatic reasoning

Permalink

<https://escholarship.org/uc/item/7630p08b>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 44(44)

Authors

Kravtchenko, Ekaterina

Demberg, Vera

Publication Date

2022

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Modeling atypicality inferences in pragmatic reasoning

Ekaterina Kravtchenko (eskrav@lst.uni-saarland.de)

Department of Language Science and Technology
Saarland University
66123 Saarbrücken, Germany

Vera Demberg (demberg@lst.uni-saarland.de)

Department of Computer Science
Saarland University
66123 Saarbrücken, Germany

Abstract

Empirical studies have demonstrated that when comprehenders are faced with informationally redundant utterances, they may make pragmatic inferences (Kravtchenko & Demberg, 2015). Previous work has also shown that the strength of these inferences depends on prominence of the redundant utterance – if it is stressed prosodically, marked with an exclamation mark, or introduced with a discourse marker such as “Oh yeah”, atypicality inferences are stronger (Kravtchenko & Demberg, 2015, 2022; Ryzhova & Demberg, 2020). The goal of the present paper is to demonstrate how both the atypicality inference and the effect of prominence can be modelled using the rational speech act (RSA) framework. We show that atypicality inferences can be captured by introducing joint reasoning about the habituality of events, following Degen, Tessler, and Goodman (2015); Goodman and Frank (2016). However, we find that joint reasoning models principally cannot account for the effect of differences in utterance prominence. This is because prominence markers do not contribute to the truth-conditional meaning. We then proceed to demonstrate that leveraging a noisy channel model, which has previously been used to model low-level acoustic perception (Bergen & Goodman, 2015), can successfully account for the empirically observed patterns of utterance prominence.

Keywords: world knowledge; experimental pragmatics; Bayesian modeling; noisy channel

When comprehenders encounter utterances that are pragmatically unexpected in the light of world knowledge, they may accommodate them by revising their beliefs about the common ground. Research on pragmatic inferences has to date paid relatively little attention to such common ground inferences, and formal models of pragmatic reasoning, with the notable exception of Degen et al. (2015), similarly do not typically account for the effects unexpected utterances may have on background beliefs about the world. As Degen et al. show, these inferences may substantially alter utterance interpretation. Examples of the types of utterances we concern ourselves with are taken from the experiment by Kravtchenko and Demberg (2015) and include the following conditions (actual materials contain longer stories and are abbreviated here to highlight the critical manipulation):

1. baseline (no informationally redundant utterance)
e.g., “John went shopping.”
2. informationally redundant utterance (no marking)
e.g., “John went shopping. *He paid the cashier.*”
3. informationally redundant utterance (exclamation mark)
e.g., “John went shopping. *He paid the cashier!*”

4. informationally redundant utterance (discourse marker)
e.g., “John went shopping. *Oh yeah, and he paid the cashier.*”

In each case, the speaker first establishes that a stereotypical series of actions, such as the *shopping* event, occurred. In the case of utterance (1), the speaker then stops. However, given world knowledge about the structure and typical activity components of such events, most listeners conclude that habitual events associated with that activity such as *paying the cashier* must have taken place, even if not mentioned explicitly (Bower, Black, & Turner, 1979). According to the cooperative principles (Grice, 1975), listeners expect for rational speakers to not be unnecessarily verbose and to hence omit information that does not need to be explicitly stated to be inferred accurately.

Kravtchenko and Demberg (2015); Ryzhova, Mayn, and Demberg (2022) have established that informationally redundant utterances of the form (2-4) above lead to pragmatic inferences, involving a revision of the beliefs about the habituality of the mentioned activity. In the case of our example, this means that the listener infers that *John* must be a less habitual payer than initially assumed (for instance, he might be a shop lifter), as this is a way to justify the activity’s explicit mention. Ryzhova et al. (2022) have recently demonstrated that such inferences are indeed drawn by comprehenders and that they are expressed in lower ratings regarding the question whether John typically pays the cashier.

Figure 1 is based on data from a large-scale replication of the original study (reported in Kravtchenko & Demberg, 2022). It shows listener distribution of *habituality* estimates after reading utterances of the type (1-4). These estimates were obtained by asking participants to provide ratings of how habitual a given activity was, in the context of a particular activity sequence, see the section on the habituality prior for more details.

In Figure 1, it can be observed that *habituality* estimates – i.e., how often *John* is expected to *pay the cashier* – are overall rather high in the *null utterance* condition. This is expected for a predictable activity. However, they noticeably (and significantly) decrease when encountering a redundant utterance. This decrease in habituality ratings is what we refer to as the atypicality inference.

The ratings of the two marked conditions ((3) exclama-

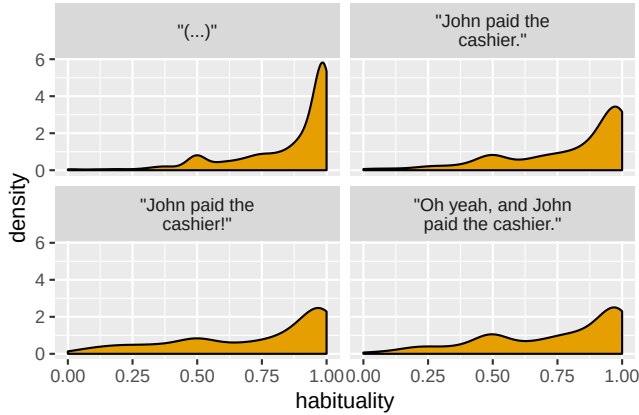


Figure 1: Smoothed distributions (kernel density estimates) showing the comprehender’s estimate of how likely they think the mentioned character (here: John) is to habitually pay the cashier. The estimates are based on 24 items and responses from a total of 2100 participants, reported in Kravtchenko and Demberg (2022). The *atypicality* inference is visible in the shift of estimates away from the right end in the null utterance baseline “(...)” condition to lower estimates in the other conditions.

tion mark and (4) discourse marker) are furthermore significantly lower than the ratings for the full stop condition (2). Kravtchenko and Demberg (2015) suggest that longer or more prominent utterances like (3) and (4) can be thought of as being perceived and attended to better than less prominent but meaning-equivalent utterances like (2). More prominent or attention-drawing utterances are both empirically (see Figure 1) and theoretically (cf. Wilson & Sperber, 2004) predicted to strengthen pragmatic inferences, as they can signal greater speaker intent.

In this paper, we thus aim to

- **Goal A:** model the atypicality inference (difference between conditions (1) and (2-4))
- **Goal B:** model the interpretation of truth-conditionally equivalent utterances that differ in their perceptual prominence (difference between (2) and (3, 4))

We propose to do this using the rational speech act (RSA) (Frank & Goodman, 2012) framework. The rational speech act model combines ideas from Gricean iterative reasoning with probabilistic approaches – probabilistic speakers and listeners recursively reason about each others mental states. The basic “literal listener” simply infers the current world state s based on whether it is compatible with the utterance u and the world state’s prior probability:

$$P_{L_0}(s|u) \propto \llbracket u \rrbracket(s) \cdot P(s)$$

The speaker chooses their utterance based on the utility of the utterance for successfully communicating their intended

meaning to the literal listener, and an utterance cost C which reflects speaker effort, where α and λ are parameters that determine the rationality of the speaker, and the extent to which the speaker weighs the utterance cost:

$$P_{S_1}(u|s; \alpha, \lambda, C) \propto P(u; \lambda, C) \exp(\alpha \log P_{L_0}(s|u))$$

Finally, the pragmatic listener interprets the perceived utterance while reasoning about the speaker’s choice of utterance:

$$P_{L_1}(s|u) \propto P_{S_1}(u|s; \alpha, \lambda, C) \cdot P(s)$$

We will demonstrate that two previously proposed model extensions can be used in order to account for the overall effects: to achieve goal (A), a joint reasoning model (Degen et al., 2015) can be adapted to account for the atypicality inference. It turns out, however, that the resulting RSA model principally cannot predict inferences of different strengths for the utterance prominence conditions, as required for goal B. We therefore propose to re-interpret the noisy channel RSA model proposed by Bergen and Goodman (2015) in terms of attentional processes. While Bergen and Goodman (2015) used the noisy channel to model mishearing at the acoustic level, we here propose to use the same noisy channel construct to account for effects of attentional prominence.

Modeling a shift in background beliefs

The literal meaning of *paid the cashier* communicates nothing about activity *habituality* directly. Standard RSA models, where the listener only infers a world state given an utterance, can therefore accurately predict only that the *cashier* was definitely paid in the case of utterances (2-4), and that they may or may not have been paid, modulated by prior beliefs about the habituality of *paying*, in the case of utterance (1). Activity habituality by itself cannot be modeled within this framework, since all utterances are at face value equally consistent with all possible *habituities*.

To model the *atypicality* inference, it is necessary to minimally incorporate joint reasoning about background knowledge, following the proposal of Degen et al. (2015). Here, the listener reasons jointly about the current world state (s) (i.e., did the activity in question occur, or not), as well as the true *habituality* of the activity (h), given the speaker’s utterance (u).

Habituality RSA (hRSA) model

A RSA model which incorporates joint reasoning (e.g., Degen et al., 2015; Goodman & Frank, 2016) can model both changes in beliefs about the world, and changes in beliefs about the current activity state. Here, we feed empirical priors about event habituality (see next section) directly into the model, where the likelihood of the activity occurring is conditional on the activity *habituality*. Whether a given activity occurred, or not (s), then, is simply a Bernoulli trial with $p = h$. In the habituality RSA (hRSA) model, the literal listener arrives at the most likely current world state (s) (whether the

activity took place, or not), given the utterance (u), and prior beliefs about activity habituality (h):

$$P_{L_0}(s|u, h) \propto \llbracket u \rrbracket(s) \cdot P(s|h)$$

L_0 does not reason about *habituality*, as this is not a part of the literal interpretation.

The pragmatic speaker, S_1 , considers the likelihood that a given utterance will communicate the current activity state to the listener, given common-ground beliefs about *habituality*, while balancing the cost C^1 of uttering the potential utterances relative to one another:

$$P_{S_1}(u|s, h; \alpha, \lambda, C) \propto P(u; \lambda, C) \exp(\alpha \log P_{L_0}(s|u, h))$$

In our model, we set the costs for null utterances to 0, the cost for a plain utterance to 3, the cost for the exclamation mark utterance to 4 and the cost for the *oh yeah* utterance to 4.5. These cost settings were not estimated empirically; they are quite robust to changes in the numerical values. As we will see below, the utterance costs have little effect on the pragmatic listener posterior for which we can compare to empirical data; they mostly affect the speaker choice, which still needs empirical testing.

The relative weights that speakers give to the cost and utility functions are represented by the parameters λ and α . λ expresses the speaker’s prioritization of reducing utterance cost; in our model, we set it to 1. α expresses the speaker’s rationality, i.e., the degree to which the speaker maximizes utterance utility. In our model, α is set at 7. Only one level of recursion is used, as is standard, given limited empirical evidence for deeper levels of recursion in online pragmatic reasoning (Goodman & Stuhlmüller, 2013; Goodman & Frank, 2016).

The pragmatic listener, L_1 , considers the likelihood that a given utterance would be chosen by the speaker, given the probabilities of particular world states and activity habitualities, and arrives at the most likely interpretation of the utterance on this basis:

$$P_{L_1}(s, h|u) \propto P_{S_1}(u|s, h; \alpha, \lambda, C) \cdot P(s|h) \cdot P(h)$$

Note that in the present hRSA model, the habituality is a type of belief about the world, like in other joint reasoning models. However, it is not a belief about the habituality of the activity in general. Instead, the atypicality inference describes whether the general habituality of an activity generalizes to the agent of the story (Generally, people pay when going shopping, hence John pays when going shopping), or whether the habituality of the activity does not extend to John, i.e., *John* doesn’t usually pay.

Estimating the habituality prior

In order to calculate the model predictions, we need to estimate a prior for the habituality of the activities. We decided

¹See <https://michael-franke.github.io/probLang/chapters/app-03-costs.html>, for a discussion of how utterance costs C should be formalized within the RSA framework.

to use a beta distribution for this, and estimate its parameters empirically from the data. Kravtchenko and Demberg (2022) collected ratings for the habituality of various activities, from 2100 participants. These participants were asked to indicate, on a sliding scale from never to always, how often they thought someone engaged in a particular activity (such as paying the cashier) as part of a certain event sequence (such as going shopping), see Figure 2 for an example. The slider was discretized into numbers 0 to 100 for analysis of the data.

Q: How often do you think John usually pays the cashier, when grocery shopping?

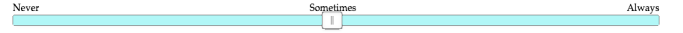


Figure 2: Slider used for collecting habituality estimates from participants.

Importantly, the question was asked in the context of a story which introduced the overall activity (e.g., shopping), but did not contain the redundant target utterance (like condition (1) in the example in the introduction). These ratings are used in our study to estimate the habituality prior. Figure 1, top left panel, shows the distribution of *habituality* estimates collected from participants. We can see that most participants believe that the event is highly likely to happen, as visible from the high number of ratings above 0.9.

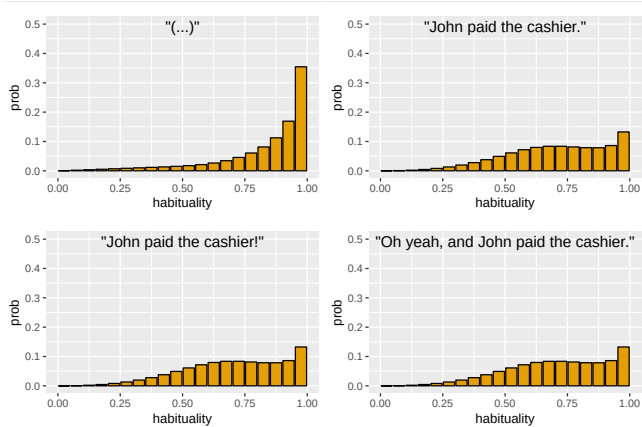
Note that the distributions in Figure 1 are bimodal: there is also a small peak around the 0.5 mark. We believe that this peak around 0.5 is an artifact related to the method of data collection, reflecting the well-known midpoint bias. The midpoint bias says that people tend to select the mid point on a scale, to express uncertainty; even if their “real” estimate might be 0.45 or 0.55, they are more likely to put the slider in the middle, instead of at these values close to the middle of the scale. This midpoint bias affects all of the conditions. However, for our RSA models, we do not aim to replicate this midpoint bias, as the RSA model is a model of human inference, and not a model of how an inference gets mapped onto a slider. Rather, we aim to model the overall pattern of the distribution, i.e. the proportion of responses qualifying the event as highly likely vs. the heavier tail.

In order to set the habituality prior for the computational model, we fit a beta distribution to the empirical response data, using the `fitdistrplus` R package (Delignette-Muller & Dutang, 2015; R Core Team, 2018). The resulting estimate for the prior was fed directly into the model.

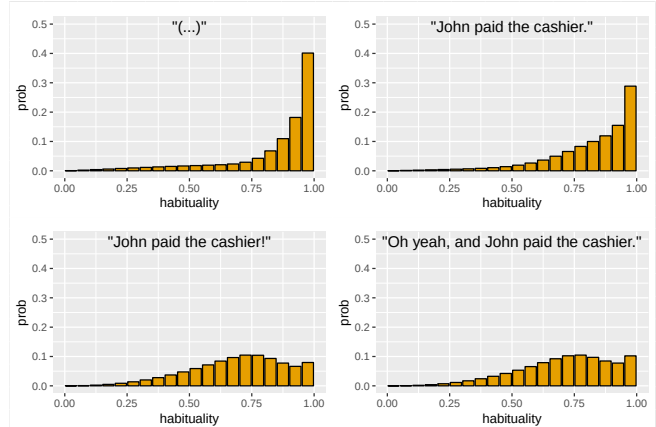
Results

The hRSA model correctly captures the predicted effect of goal (A), as seen in Figure 3a: if an activity is described explicitly, the *habituality* is likely to be low. What it does not, however, capture is the effect of utterance prominence (goal (B)): there is virtually no effect of utterance prominence on interpretation by the pragmatic listener.

The crucial point for understanding this failure to show the



(a) hRSA model.



(b) Noisy channel hRSA model.

Figure 3: Posterior (pragmatic listener) probabilities for the noisy channel hRSA model.

desired effects is as follows: There are three possible ways of articulating the redundant utterance: with a full stop at the end (2); with an exclamation mark (3); or with an attention-drawing and relevance-establishing discourse marker (4), in order of increasing utterance cost. The more attentionally prominent utterances (3-4) will never be of any advantage to the literal listener, in terms of whether they effectively communicate the current world state. They are likewise of no advantage to the speaker, either in terms of likelihood of accurate message transmission to the literal listener, or the speaker’s presumed goal to conserve articulatory effort. As a consequence, the pragmatic listener will not infer that the more effortful utterance is more likely to communicate an atypical meaning.

We note that the model also makes predictions about speaker behaviour, for which we have currently no empirical evidence, but which could be tested in the future. Figure 4a illustrates the model predictions for the speaker: An example of an activity with 95% habituality could be *paying the cashier*; as we can see, the speaker would be predicted not to mention this explicitly. An activity with 50% habituality could for instance be *buying apples*. Here, we can see that the speaker would be predicted to prefer a plain utterance, with some probability also distributed among the other choices. For a surprising event with just 5% habituality, which could for instance correspond to accidentally dropping something, the speaker’s utterance choices are almost identical to the ones for 50% habituality. The only difference is that the empty utterance is not predicted. In particular, the model predicts that speakers are very reluctant to use exclamation marks or other markers even in this condition.

The failure of standard RSA models to derive pragmatic inferences of different strengths, given semantically meaning-equivalent utterances, is directly analogous to their failure to derive M-implicatures or inferences due to prosodic stress, as detailed and mathematically proven in Bergen, Levy, and Goodman (2016). As a result, this model fails to capture any

Table 1: Confusion matrix showing the likelihood of any given utterance being perceived as any other.

	(...)	He paid.	He paid!	Oh yeah...
(...)	0.99	0.01	0.0001	0.0001
He paid.	0.01	0.95	0.02	0.02
He paid!	0.0001	0.02	0.97	0.01
Oh yeah...	0.0001	0.02	0.01	0.97

of the empirically demonstrated effects that increased utterance salience has on utterance choice or comprehension, also predicted by psycholinguistic theories of language comprehension (e.g., Levy, 2008).

Attentional prominence and inference strength

In order to achieve goal (B), it is necessary to assign some attentional benefit to the more costly redundant utterance, to be already active at the L_0 level. Empirically, there is evidence that readers often cannot recall whether elements in a stereotyped activity sequence were explicitly mentioned, or not (Bower et al., 1979), and that informational redundancy, even at the multi-word level, in part serves the purpose of ensuring that listeners attend to and accurately recall relevant information (Walker, 1993; Baker, Gill, & Cassell, 2008). The noisy-channel RSA model proposed by Bergen and Goodman (2015) successfully captures this intuition, although in our case we consider the probability that an utterance is attended to and stored in memory, rather than simply misheard, as represented in Table 1.

The exact values in the table are set somewhat arbitrarily, as we do not have the empirical data for fitting them more exactly. In our experiments, we found that modelling results are quite robust to modifications of these exact values. The most important aspect to capture the empirical effects is that the utterances that are prominent and attract more attention should have a very small confusion probability with the null

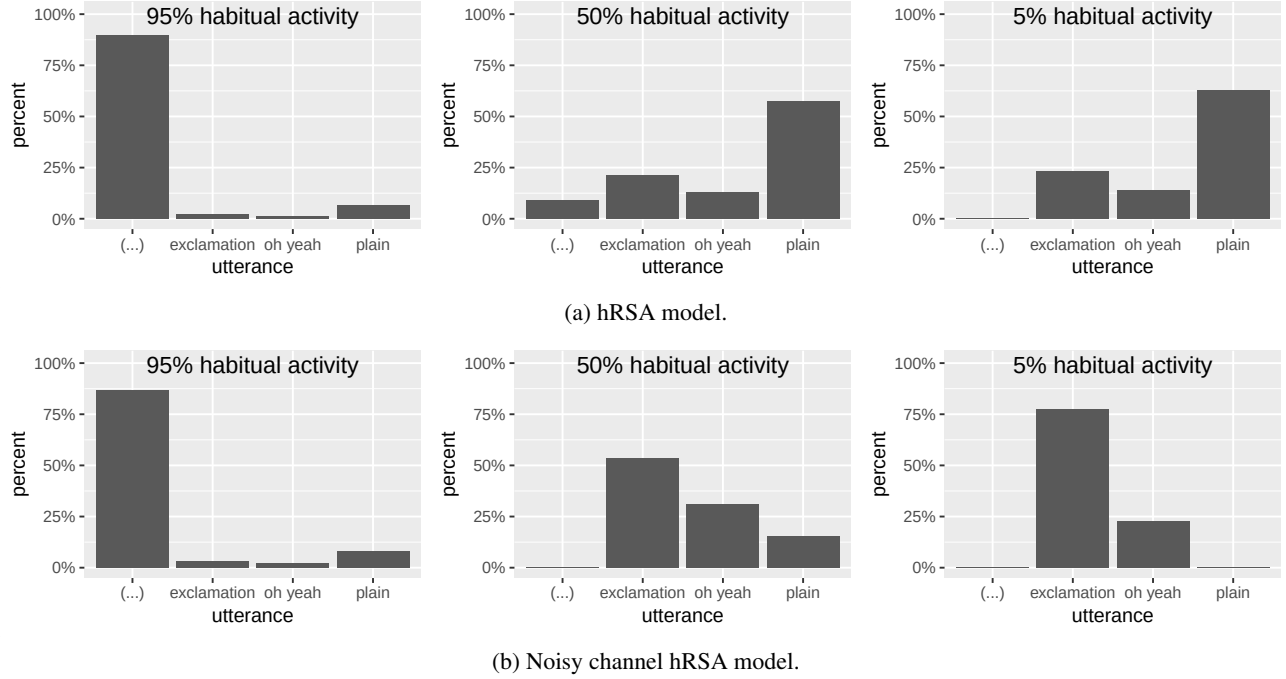


Figure 4: Speaker’s utterance preferences in the hRSA and noisy channel hRSA model for events with different habitualities (e.g., *paying the cashier* vs. *buying apples* vs. *dropping the eggs* in a grocery shopping scenario).

utterance, in particular much smaller than the probability of the full stop condition.

We chose the values shown in the table as we found them to be intuitively plausible. On the diagonal, that each utterance is most likely to be recalled and remembered as itself. Each utterance also has a small likelihood of being misperceived as a perceptually ‘neighboring’ utterance: “*he paid!*” and “*oh yeah...*” both have a small likelihood of being mistakenly recalled as the other, and a higher likelihood of being recalled as the plain utterance: “*he paid.*”. The plain utterance (“*he paid.*”), which does not draw any particular attention, has a small likelihood of not being remembered (“(...”), and the ‘null’ utterance (“(...”) may mistakenly be recalled as the plain utterance. Although this last confusion may appear counterintuitive, Bower et al. (1979) shows that in script contexts, participants frequently recall reading about habitual activities that were not, in fact, mentioned explicitly. To sum up the general intuition, the listener is more likely to notice and accurately recall more perceptually prominent utterances, and is correspondingly more likely to wonder why the speaker went to the extra effort in producing these utterances.

Noisy hRSA model

In the noisy channel hRSA model, it’s assumed that every utterance has a non-trivial likelihood of not being actively attended to, and being mistaken for or mis-recalled as a ‘neighboring’ utterance. Here, u_i represents the utterance intended by the speaker, and u_r represents the utterance actually recalled by the listener. At every level, the listener or speaker

reason about the likelihood that the utterance they actually perceived is not the utterance that was uttered, or, conversely, that the utterance they intend may not be the utterance that is in fact perceived. Again, only one level of recursion is used, and only one is necessary to capture these results. To note, given the mathematical properties of RSA models (Bergen et al., 2016), deeper levels of recursion would not in themselves alter the ability of the model to capture the utterance-dependent variations in inference strength.

The literal listener reasons about the likely world state given the utterance they in fact recall, and the habituality of the activity in question. However, they weigh this by the likelihood that the utterance recalled is not in fact the utterance that was intended.

$$P_{L_0}(s|u_r, h) \propto \llbracket u_r \rrbracket(s) \cdot P(s|h) \cdot \sum_{u_i: \llbracket u_i \rrbracket(s)=1} P(u_r|u_i)P(u_i)$$

The pragmatic speaker chooses an utterance u_i given a world state and activity *habituality*, taking into consideration the likelihood that the listener may misremember or mis-recall the utterance they intend. Intuitively, the confusability is more likely to play a role when the meaning that the speaker intends to transmit is unexpected by the listener, compared to when it is a highly expected meaning.

$$P_{S_1}(u_i|s, h; \alpha, \lambda, C) \propto P(u_i; \lambda, C) \exp(\alpha \sum_{u_r} P(u_r|u_i) \log P_{L_0}(s|u_r, h))$$

The pragmatic listener, as in the hRSA model, infers the current world state and activity habituality, taking into ac-

count the conditions under which the speaker made their utterance choice, their habituality prior, and the likelihood of the state given the habituality. The pragmatic listener again takes into account the possibility that they may mis-recall the speaker’s intended utterance.

$$P_{L_1}(s, h|u_r) \propto P(s|h) \cdot P(h) \cdot \sum_{u_i} P_{S_1}(u_i|s, h; \alpha, \lambda, C) P(u_r|u_i) P(u_i)$$

In sum, a redundant event description that does not somehow draw the listener’s attention is less likely to be attended to, and more likely to be misperceived or misremembered by the literal listener as a ‘*null*’ utterance. The pragmatic speaker must take into account that their utterance may not be attended to or remembered by the listener, and the pragmatic listener likewise considers the possibility that they may fail to attend to or remember what the speaker uttered.

Results

This model qualitatively captures both goals A and B. As can be seen in Figure 3b, pragmatic listeners adjust the common ground such that a typical activity which is uttered overtly is inferred to be less habitual: the peak at the high ratings is greatly reduced in favour of a heavier tail. This is thus reflective of goal (A). The figure also shows that more effortful utterances (3) and (4) lead to stronger *atypicality* inferences (for high-habituality activities): the tails for the exclamation mark condition and the “oh yeah” condition are substantially heavier than for the full stop condition. This is in line with the empirical data in Figure 1 and hence with goal (B) above. The heaviness of the tail depends on the settings of the alpha parameter: heavier tails without bimodality can be achieved by choosing lower values for alpha, thinner tails with a second peak between 0.6 and 0.7 can be achieved by choosing higher values for alpha. Comparing the empirical posterior habituality estimates from Figure 1 to the habitualities inferred by the pragmatic listener from Figure 3b, we can see that the overall pattern is qualitatively similar.

The noisy channel RSA model also makes specific predictions for the speaker, see Figure 4b. For high-habituality activities, speakers are very unlikely to describe the activity explicitly, just like in the plain hRSA model. And if they do, they tend to choose less effortful utterances, as shown in Figure 4b. Moderately habitual activities will always be mentioned, according to the model, and speakers may even choose the marked utterances. The strength of this effect can be modulated by changing alpha and/or by changing the confusion matrix. Non-habitual activities are virtually always described explicitly, and as can be seen, speakers strongly prefer a higher-effort utterance that is more likely to be attended to, and less likely to be mis-recalled as a “*null*” utterance. We also note a difference in speaker preferences between the two marked utterances, with exclamation mark utterances being preferred over the *oh yeah* utterances. This can be attributed to the higher production cost that we postulated in the model for the *oh yeah* utterance compared to the exclamation mark

– the asymmetry between the conditions would disappear if costs were assigned equal values (as we set the confusability values to be equivalent for the two conditions). In summary, the model predicts that speakers are likely to use more effortful utterances to communicate less likely meanings.

Conclusion

In this paper we set out to demonstrate that atypicality inferences can be qualitatively modeled using the RSA framework. We firstly showed that the incorporation of joint reasoning about the utterance and world knowledge can successfully account for common ground atypicality inferences. However, this model can not account for effects of utterance prominence between meaning-equivalent utterances. We addressed this by demonstrating that the noisy channel RSA model by Bergen and Goodman (2015) can be re-interpreted for this case, such that the noisy channel model describes confusability between messages due to higher-level attentional processes, instead of low-level perceptual ones. The noisy hRSA model demonstrates that the empirically observed atypicality inference effects can be captured well qualitatively, and that this can be accomplished using existing and independently motivated mechanisms, and does not require the postulation of any new mechanisms.

Limitations of our work include the fact that we did not empirically estimate the confusion matrix for utterances for the confusion matrix of the noise model, or the cost of the alternative utterances, and consequently did not attempt a quantitative model fit for the noisy hRSA listener model. A further area of future work consists of empirically estimating to what extent speakers choose more prominent utterances as a function of how non-predictable/surprising the target utterance is, to test the model predictions about speaker behaviour.

Finally, we note that our interpretation of the effect of the exclamation mark and the discourse marker as drawing more attention to the utterance, which then causes stronger pragmatic inferences is a stipulation, which should be further tested in future work. The noisy channel mechanism as such is insensitive to our hypotheses regarding what type of process the noise may stem from (acoustic noise, encoding difficulties or noise through memory recall), but trying to experimentally pin down the reason for the effect and modelling it in terms of a process model represent interesting areas of future research.

Acknowledgments

We would like to thank the reviewers for their very constructive comments, which helped us greatly in improving this paper. This work was supported by the Deutsche Forschungsgemeinschaft, Funder Id: <http://dx.doi.org/10.13039/501100001659>, Grant Number: SFB1102: Information Density and Linguistic Encoding. The second author gratefully acknowledges funding from the European Research Council (ERC Starting Grant “Individualized Interaction in Discourse”, 948878) under the European Union’s Horizon 2020 research and innovation programme.

References

- Baker, R. E., Gill, A. J., & Cassell, J. (2008). Reactive redundancy and listener comprehension in direction-giving. In *9th sigdial workshop on discourse and dialogue*.
- Bergen, L., & Goodman, N. D. (2015, apr). The Strategic Use of Noise in Pragmatic Reasoning. *Topics in Cognitive Science*, 7(2), 336–350.
- Bergen, L., Levy, R., & Goodman, N. (2016, may). Pragmatic reasoning through semantic inference. *Semantics and Pragmatics*, 9.
- Bower, G. H., Black, J. B., & Turner, T. J. (1979). Scripts in memory for text. *Cognitive Psychology*, 11(2), 177–220.
- Degen, J., Tessler, M. H., & Goodman, N. D. (2015). Wonky worlds: Listeners revise world knowledge when utterances are odd. *Proceedings of the 37th Annual Conference of the Cognitive Science Society*, 548–553.
- Delignette-Muller, M. L., & Dutang, C. (2015). fitdistrplus: An R package for fitting distributions. *Journal of Statistical Software*, 64(4), 1–34.
- Frank, M. C., & Goodman, N. D. (2012). Predicting pragmatic reasoning in language games. *Science*, 336(6084), 998.
- Goodman, N. D., & Frank, M. C. (2016). *Pragmatic Language Interpretation as Probabilistic Inference* (Vol. 20) (No. 11).
- Goodman, N. D., & Stuhlmüller, A. (2013). Knowledge and Implicature: Modeling Language Understanding as Social Cognition. *Topics in Cognitive Science*, 5(1), 173–184.
- Grice, H. P. (1975). Logic and conversation. In *Speech acts* (pp. 41–58). Brill.
- Kravtchenko, E., & Demberg, V. (2015). Semantically underinformative utterances trigger pragmatic inferences. In *Cogsci* (pp. 1207–1212).
- Kravtchenko, E., & Demberg, V. (2022). Informationally redundant utterances elicit pragmatic inferences. *Cognition, to appear*.
- Levy, R. (2008). A noisy-channel model of rational human sentence comprehension under uncertain input. In *Proceedings of the 13th conference on empirical methods in natural language processing* (p. 234–243).
- R Core Team. (2018). R: A language and environment for statistical computing [Computer software manual].
- Ryzhova, M., & Demberg, V. (2020). Processing particularized pragmatic inferences under load. In *Proc. of the annual conference of the cognitive science society*.
- Ryzhova, M., Mayn, A., & Demberg, V. (2022). What inferences do people actually make on encountering a redundant utterance? an individual differences study. In *arxiv*.
- Walker, M. A. (1993). *Informational redundancy and resource bounds in dialogue*. Doctoral Dissertation, University of Pennsylvania, Philadelphia, PA.
- Wilson, D., & Sperber, D. (2004). Relevance Theory. In L. R. Horn & G. Ward (Eds.), *The handbook of pragmatics* (Vol. 1, pp. 606–632). Oxford, UK: Blackwell Publishing.