

Cost-Based Pragmatic Inference about Referential Expressions

Judith Degen (jdeggen@bcs.rochester.edu)

Dept. of Brain and Cognitive Sciences
University of Rochester

Michael Franke (m.franke@uva.nl)

Institute for Logic, Language and Computation
Universiteit van Amsterdam

Gerhard Jäger (gerhard.jaeger@uni-tuebingen.de)

Institute of Linguistics
University of Tübingen

Abstract

We present data from three experiments addressing how much theory of mind reasoning is involved in production and interpretation of ambiguous referential expressions in an artificial language task, and how this interacts with the cost and availability of alternative utterances. When an unambiguous alternative is not available, listeners tend to draw simple Quantity inferences reminiscent of scalar implicatures (Grice, 1975). When an unambiguous alternative *is* available, fewer inferences are observed, but gradually more as the cost of unambiguous alternatives increase. We outline a novel game theoretic model of pragmatic reasoning based on probabilistic back-and-forth reasoning about interlocutors' rational choices and beliefs. The model provides a good fit to the data and raises interesting issues for future research.

Keywords: Pragmatics; Game theory; Referential Expressions; Language production; Language comprehension.

Introduction

People are lazy: when they speak, they like to save effort. But if speakers are too lazy and say too little, their listeners will not understand them. A good example is the choice and interpretation of referential expressions. A rational speaker who wants to establish reference should choose the most economic (shortest, easiest, least effortful) description that, according to his beliefs about the listener's dispositions to interpret utterances, will allow for the listener to safely infer the correct referent. A rational listener should take the speaker's production costs¹ into account and so a rational speaker should in turn take that into account, etc. But this is an idealized picture. From an empirical point of view two related questions arise: **1) How much do speakers and listeners take into account each other's perspective? 2) How much influence do economy considerations have; do listeners weigh the speaker's production costs?**

When it comes to referential language use, the latter question has not been investigated thoroughly, but the former question has been addressed in a variety of ways, both theoretically (Clark & Marshall, 1981) as well as experimentally (Hanna, Tanenhaus, & Trueswell, 2003; Keysar, Barr, & Brauner, 2000). This paper aims to address both questions and to bring theoretical and empirical approaches closer together. The paper's empirical contribution is to report on

experimental data from referential language games. Its theoretical contribution is a novel probabilistic model of back-and-forth reasoning that synthesizes recent Bayesian models (Frank & Goodman, 2012) and game theoretic approaches (Camerer, Ho, & Chong, 2004; Rogers, Palfrey, & Camerer, 2009; Franke, 2011; Jäger, 2013).

Our experiments probed interlocutors' perspective-taking ability in a task where an artificial language left some relevant meaning features inexpressible or made some expressions costly. When critical meaning features are inexpressible, the situation is reminiscent of scalar implicature calculation (Grice, 1975). We manipulated how many steps of such reasoning were needed for communicative success and tested both comprehension (Exp. 1) and production (Exp. 2) to investigate question 1. Our design was chosen so as to improve on previous related studies where non-linguistic pictorial messages were used (Degen & Franke, 2012) and where the availability of alternative expressions was not explicitly controlled (Stiller, Goodman, & Frank, 2011; Frank & Goodman, 2012). In addition, rather than making some messages entirely unavailable, we investigated the effect of variable production costs on interpretation behavior to address question 2 (Exp. 3). Rohde, Seyfarth, Clark, Jäger, and Kaufmann (2012) showed that listeners take into account message costs when messages are assigned an explicit dollar value. Here we investigate whether these results replicate when costs, as in real language use, are implicit.

The data from our experiments is explained well by a probabilistic model of back-and-forth reasoning. The model parameterizes how deeply interlocutors reason about each other's perspective and how close they are to being rational. Parameter values that best explain our data suggest that participants were reasonably rational and applied a small but non-negligible amount of theory of mind reasoning and that they took production costs into account in comprehension.

Referential Language Games

Referential communication can be conceived of as a *signaling game* (Lewis, 1969): a sender (speaker) *S* knows which referent he wants to talk about, but a receiver (listener) *R* does not; *S* chooses a referential description; if *R* can identify the intended referent, communication is successful, otherwise a failure. Different games ensue for different sets of potential referents and referential expressions. In the critical trials of our experiments the referential games were isomorphic to the

¹Many factors have been identified as contributing to production preferences (see e.g. Jäger & Tily, 2011). Here we take participants' empirically estimated (Exp. 3) relative preference for shorter over longer messages as a measure of subjective production cost.

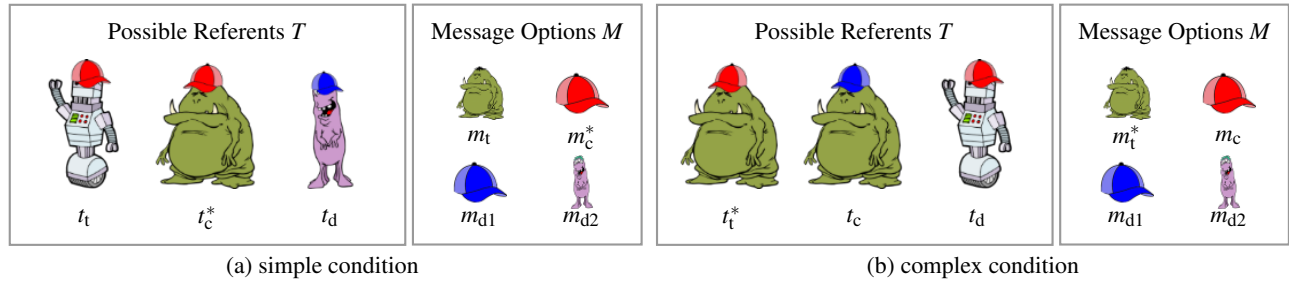


Figure 1: Target implicature conditions. Hearers choose one of the POSSIBLE REFERENTS $T = \{t_t, t_c, t_d\}$. Speakers have MESSAGE OPTIONS $M = \{m_t, m_c, m_{d1}, m_{d2}\}$, shown here for ease of interpretation visually (the experiment used artificial words). Trigger items are indicated with asterisks: e.g., t_t^* is the referent to be communicated on complex production trials.

situations in Fig. 1. There were three possible referents in the form of monsters and robots wearing hats or scarves (not depicted in the example) as accessories. Additionally, there is a fixed set of possible descriptions that are available to the sender. Messages for monsters and hats were always available and were equally costly. Messages for robots and scarves were either not available at all (Exp. 1 and 2) or more costly (Exp. 3).

Experiments 1 and 3 tested participants’ choice of referent for a given *trigger message* (comprehension). Experiment 2 tested their choice of message for a *trigger referent* (production). Trigger items for the critical conditions are marked with an asterisk in Fig. 1. Indices t, c, d stand for *target*, *competitor* and *distractor* respectively. We refer to a game as in Fig. 1a as the *simple condition*, because it involves one step of Quantity reasoning similar to scalar implicature calculation (Grice, 1975). Hearing *trigger message* m_c^* , R should reason that S must have meant *target state* t_t , and not *competitor state* t_c , because if S had wanted to refer to the latter he could have used an unambiguous message. Conversely, when S wants to refer to *trigger state* t_c^* , he should not use the true but semantically ambiguous message m_c , because he has an unambiguous message m_t . Similarly, we refer to a game in Fig. 1b as the *complex condition*, because it requires performing similar reasoning twice in sequence (see also Degen and Franke (2012) for details).

Experiment 1 - comprehension

Exp. 1 tested participants’ behavior in a *comprehension* task using instantiations of the signaling games just described.²

Methods

Participants Using Amazon’s Mechanical Turk, 48 self-reported native English speakers were paid \$1.00 to participate. All were naïve as to the purpose of the experiment.

Procedure and Materials Participants engaged in an artificial language referential comprehension task. The experi-

ment proceeded in two stages: a *language learning* stage and an *inference* stage. Only the inference stage was of theoretical interest. In the language learning stage, participants learned four 3-character words (*RAV*, *ZUB*, *XEK*, *KOR*) of the alien language Zorx. The words referred to visual features: ontological kinds (one of two monster species) and accessories (red or blue hat). Each unique word-to-feature mapping (24 total) occurred twice between participants to ensure effects were not artifacts of the particular mapping.

Language learning occurred in three steps. First, participants saw each word twice with a visual representation of its meaning. They were then presented with each word alongside two choices for the meaning of the word and had to click on the correct meaning. Finally, they were presented with each meaning in succession and had to produce the correct word by clicking on characters in a two-row character array. They repeated this process until achieving 100% accuracy on the production task. They then proceeded to the inference stage.

On each trial in the inference stage, participants saw three possible referents on a display (as in Fig. 1). Each referent differed systematically along two dimensions: its ontological kind (robot or one of two monster species) and accessory (scarf or either blue or red hat). In addition to these three referents, participants saw a Zorx word that they were told was sent to them by a previous participant whose task it was to get them to pick out one of these three referents. They were told that the previous participant was allowed to send a message expressing only one feature of a given referent, and that the other participant had learned the same words of Zorx they did (i.e., they could send monster/hat messages, but not robot/scarf messages).

Participants initially completed four production trials. They saw three referents, one of which was highlighted with a yellow rectangle, and were asked to send one of the Zorx words to another Mechanical Turk worker to get them to pick out the highlighted object. They were told that the other worker did not know which object was highlighted but knew the same language they did. The four production trials contained three unambiguous and one ambiguous trial which functioned as fillers in the main experiment.

²Exps. 1 and 2 here were identical to Exps. 1 and 2 in Degen and Franke (2012) with the difference that we use linguistic instead of pictorial messages.

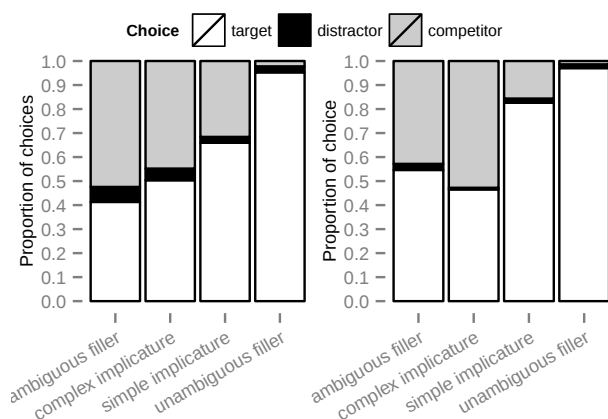


Figure 2: Proportions of target, competitor, and distractor choices in critical and filler conditions for Exp. 1 (comprehension, left) and Exp. 2 (production, right).

Participants saw 36 experimental trials, with a 2:1 ratio of fillers to critical trials. Of the 12 critical trials, 6 occurred in the simple (one iterated reasoning step) and 6 in the complex (two steps) condition as described above (see Fig. 1).

Target position was counterbalanced (each critical trial occurred equally often in each of the 6 possible orders of target, competitor, and distractor), as were the target's features and the number of times each message was sent. Of the 24 filler trials, half used the displays from the critical conditions but the target was either t_c or t_d (as identified unambiguously by the trigger message). This was intended to prevent learning associations of display type with the target. On the other 12 filler trials, the target was either entirely unambiguous or entirely ambiguous given the message. That is, there was either only one object with the feature denoted by the trigger message, or there were two identical objects that were equally viable target candidates. Trial order was pseudo-randomized such that there were two lists (reverse order) of three blocks, where critical trials and fillers were distributed evenly over blocks. Each list began with three filler trials.

Results and Discussion

Proportions of choices are displayed in Fig. 2 (left panel). As expected, participants were close to ceiling in choosing the target on unambiguous filler trials but at chance on ambiguous ones. This confirms that participants understood the task. On critical trials, participants' performance was intermediate between ambiguous and unambiguous filler trials. On simple trials, participants chose the target 66% of the time. On complex trials, the target was chosen less often (50%).

To test whether the observed differences in target choices above were significantly different, we fitted a logistic mixed-effects regression to the data. Trials on which the distractor was selected were excluded to allow for a binary outcome variable (target vs. competitor choice). The model predicted the log odds of choosing a target over a competitor from a Helmert-coded CONDITION predictor. Three Helmert con-

trasts over the four relevant critical and filler conditions were included in the model, comparing each condition with a relatively less skewed distribution against the more skewed distributions (in order: ambiguous fillers, complex, simple, unambiguous fillers). This allowed us to capture whether the differences in choice distributions for neighboring conditions suggested by Fig. 2 were significant. We included the maximal random effects structure, i.e., by-participant random intercepts, by-participant random slopes for CONDITION, and by-item random intercepts.

Of the three contrasts, two reached significance; there were more target choices in the unambiguous filler condition than in the simple condition ($\beta = 4.08, SE = 0.41, p < .0001$) and there were more target choices in the simple than in the complex condition ($\beta = 1.27, SE = 0.47, p < .01$). However, there was no significant difference in target choices between the ambiguous filler and the complex condition ($\beta = 0.38, SE = 0.45, p < .4$). This suggests that participants made simple, but not complex inferences.

Experiment 2 - production

Exp. 2 tested participants' behavior in a *production* task using instantiations of the signaling games described above.

Methods

Participants Using Mechanical Turk, 48 self-reported native speakers of English were paid \$1.20 to participate.

Procedure and Materials The experiment again proceeded in two stages: the *language learning* stage and the *production* stage. The procedure for language learning was the same as in Exp. 1. The procedure for the production stage was the same as on the production trials in Exp. 1. Participants saw 36 trials with a 2:1 ratio of fillers to critical trials. There were 12 critical trials (6 simple and 6 complex situations as in Fig. 1). Half of the fillers used the same displays as the critical trials, but one of the other two objects was highlighted. This meant that the target message was either unambiguous (e.g. when the highlighted object was t_i in Fig. 1(a) the target message was m_c) or entirely ambiguous. The remaining 12 filler trials employed other displays with either entirely unambiguous or ambiguous target messages. Two experimental lists were created and counterbalancing ensured as in Exp. 1.

Results and Discussion

Proportions of choices are displayed in Fig. 2 (right panel). To test whether the observed differences in target choices were different, the same logistic mixed-effects regression was fit to the data as in the Exp. 1 analysis. Trials on which a distractor message was sent were excluded to allow for a binary outcome variable (target vs. competitor choice).

Of the three Helmert contrasts, again only two reached significance; there were more target choices in the unambiguous filler condition than in the simple condition ($\beta = 3.84, SE = 0.47, p < .0001$) and there were more target choices in the simple than in the complex condition ($\beta =$

2.81, $SE = 0.50$, $p < .0001$). However, there was no difference between the ambiguous filler and the complex condition ($\beta = -0.43$, $SE = 0.37$, $p < .3$). This suggests that, as in comprehension, participants made simple, but not complex inferences.

Experiment 3 - comprehension with costs

Exp. 3 tested whether listeners take into account speakers' preferences for producing minimally effortful messages. To this end, we introduced messages for the robot and scarf feature but varied the implicit cost of these messages via word length measured in characters. If listeners take into account their interlocutor's perspective, their behavior should approximate the results from the simple conditions of Exp. 1 (i.e., draw more Quantity inferences) as the message becomes more costly (and thus, more dispreferred/unavailable).

Methods

Participants A total of 240 participants were recruited over Mechanical Turk who were all self-reported native speakers of English. They were paid \$0.80 plus a \$0.10 bonus if they completed the cost estimation stage in under one minute.

Procedure and Materials The experiment proceeded in three stages: the *language learning* stage, the *cost estimation* stage, and the *inference* stage. The procedure in the language learning and inference stage was the same as in Experiment 1 with the following three exceptions: a) the learned language contained two extra *costly* words (to refer to the robot and the scarf feature) in addition to four *free* words (to refer to monsters and hats), b) there were no complex, only simple conditions (Fig. 1) in the inference stage, c) there were only 12 rather than 24 filler trials, of which 6 were completely ambiguous and 6 were completely unambiguous.

The cost estimation stage was introduced to estimate participants' subjective cost function. Each of the nine permutations of feature combinations {robot, green monster, purple monster} $\times$ {scarf, red hat, blue hat} was presented to participants one at a time and they were asked to send one of the Zorx words they had learned to another participant to get them to pick out the presented referent. As in the previous experiments, sending a message required spelling out the word on a virtual keyboard on the screen by clicking on each character individually. In addition, participants were told that they would receive a bonus if they completed this part of the study in under one minute. We hoped these two features of the task would increase participants' subjective costs associated with the objective increase in number of characters and thus encourage a message cost effect.

There were three cost conditions. In the NO-COST condition, the costly messages were of the same length as the free messages (3 characters). In the LOW-COST and HIGH-COST conditions, the costly messages were one and three characters longer than the free messages, respectively. LOW-COST and HIGH-COST were manipulated within participants (we refer to this group as the COST condition, 192 participants) and the

NO-COST condition consisted of a separate group of 48 participants. Thus the languages in the different conditions:

$$\underbrace{\{RAV, ZUB, XEK, KOR\}}_{\text{free messages}} \cup \underbrace{\begin{cases} \{XAB, BAZ\} \\ \{XABI, BAZUZE\} \\ \{BAZU, XABIKO\} \end{cases}}_{\text{costly messages}} \begin{matrix} \text{NO-COST} \\ \text{LOW- and} \\ \text{HIGH-COST} \end{matrix}$$

Results and Discussion

Proportion of choices in the cost estimation stage (messages) and in the inference stage (referents) are shown in Fig. 3a and 3b. We analyzed participants' performance in both stages.

In the cost estimation stage, we analyzed participants' message choices for the four referents with one costly and one free message (i.e., referents with either a robot or a scarf feature). The NO-COST condition served as the baseline in the mixed effects logistic regression predicting the log odds of a costly over a free message choice. Cost condition was dummy-coded and entered as a three-level categorical predictor (NO-COST, LOW-COST, HIGH-COST). The model additionally included random by-participant and by-item intercepts as well as by-participant slopes for feature type (scarf or robot) to account for individual variability in participants' preferences for referring to these features. There was a significant decrease in the log odds of choosing the costly message compared to the NO-COST reference level when the message was HIGH-COST ($\beta = -0.83$, $SE = 0.30$, $p < .01$). For the LOW-COST message, the difference trended in the predicted direction ($\beta = -0.44$, $SE = 0.30$, $p < .14$). Thus, increasing message cost led to a small, but gradient decrease in preference to send the costly message.

Next, we analyzed participants' performance in the inference stage by fitting a mixed effects logistic regression model predicting target over competitor choices. Two Helmert contrasts over the three relevant cost conditions were included in the model, comparing each condition with a relatively lower cost against the higher cost level(s) (in order: no cost, small cost, high cost). The model additionally included by-participant and by-item random intercepts. The difference between the NO-COST and other conditions did not reach significance, though it trended in the predicted direction ($\beta = -0.08$, $SE = 0.05$, $p < .14$). However, there were significantly more target choices in the HIGH-COST than in the LOW-COST condition ($\beta = 0.25$, $SE = 0.12$, $p < .05$). Thus, the gradient effect of message cost on message choice is in turn reflected in listener inferences: as the cost of the unambiguous message increases, listeners make more scalar inferences and begin to approximate the behavior displayed in Exp. 1, where robot/scarf messages were entirely unavailable.

The Iterated Quantal Response Model

The observed production and comprehension behavior can be predicted by a parameterized model that returns a quantitative description of speaker and listener behavior. The iterated quantal response (IQR) model combines key features of so-called *cognitive hierarchy models* from behavioral economics

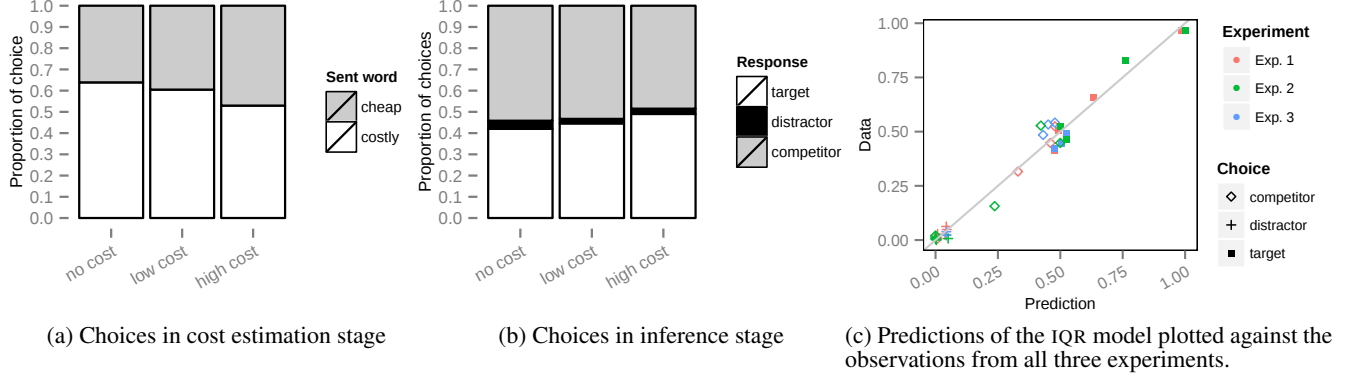


Figure 3: Experiment 3 results.

(Camerer et al., 2004; Rogers et al., 2009) with game theoretic models of pragmatic reasoning (Franke, 2011; Jäger, 2013). The resulting model is also very similar to, but slightly more general than recently popular Bayesian models (Frank & Goodman, 2012; Bergen, Levy, & Goodman, 2012).

We consider two parameters. Parameter τ models the depth of strategic reasoning that language users engage in. Parameter λ models how successful language users are at making rational choices. The output of the model is a prediction of probabilistic speaker and listener behavior.

Signaling Games We model our referential tasks as signaling games. For our purposes, a signaling game is just a structure $\langle T, M, B, c \rangle$ with $T = \{t_1, \dots, t_a\}$ a set of a different states (referents), $M = \{m_1, \dots, m_b\}$ a set of b relevant descriptions, B is a Boolean (a, b) -matrix with $B_{ij} = 1$ if description m_j is true of referent t_i , and $c = (c_1, \dots, c_b)$ a vector of costs.³

Strategies A sender strategy σ is a row-stochastic (a, b) -matrix, mapping each state onto a probability distribution over messages. A sender strategy describes how likely an average speaker would choose a message given that they want to talk about a given state. Likewise, since the receiver chooses states in T as interpretations of an observed message, a receiver strategy ρ is a row-stochastic (b, a) -matrix.

Expected Utilities A sender who believes that his listener plays ρ has an (a, b) -matrix of expected utilities $\text{EU}(\rho) = \mathbf{T}(\rho) - c$.⁴ A receiver who believes that his opponent plays σ has a unique posterior belief μ_σ derived from Bayes' formula iff σ has at least one non-zero entry in each column, i.e., each

message is expected to be sent with some positive probability (guaranteed by the quantal response function introduced below). This unique $\mu(\sigma)$ is just $\mathbf{N}(\mathbf{T}(\sigma))$. The receiver's expected utilities are then $\text{EU}(\sigma) = \mu(\sigma)$.

Best & Quantal Responses Generally speaking, a response function maps expected utilities to choice probabilities. Rational choices maximize expected utility. In case of ties, rational agents are indifferent. If U is an expected utility matrix, the rational best response function is $\text{BR}(U) = \mathbf{N}(\max \text{row}(U))$. In contrast, the quantal response function assumes that agents make small mistakes in implementing the $\text{BR}(\cdot)$ function. For given U , quantal response $\text{QR}_\lambda(U)$ is the unique row-stochastic matrix with $\text{QR}_\lambda(U)_{ij} \propto \exp(\lambda U_{ij})$. Here λ is a rationality parameter, with entirely random choices for $\lambda = 0$ and $\lim_{\lambda \rightarrow \infty} \text{QR}_\lambda(U) = \text{BR}(U)$.⁵

IQR The IQR model defines a hierarchy of player types. Unsophisticated level-0 behavior is anchored in the given semantics. Level- $(k+1)$ players play quantal responses to a belief that the interlocutor is of a lower type. Concretely, level-0 senders and receivers simply try to implement the semantic meaning: $\sigma_0 = \text{QR}_\lambda(B - c)$ and $\rho_0 = \text{QR}_\lambda(\mathbf{T}(B)I_a)$. Level- $(k+1)$ player behavior is defined as a quantal response to a belief that the other player is at most of level k . Following the relevant literature in behavioral game theory (Camerer et al., 2004; Rogers et al., 2009), we subscribe to the simplifying assumption that the actual distribution of strategic types is a Poisson distribution $\text{Pois}(\tau, k) = \tau^k / k! \exp(-\tau)$ with parameter τ , and that agents know this. So, level- $(k+1)$ players' beliefs are derived by conditioning the underlying population distribution by the event that the opponent is less sophisticated: $f_\tau^{\leq k}(l) = \text{Pois}(\tau, l) / \sum_{i=1}^k \text{Pois}(\tau, i)$ if $l \geq k$ and 0 otherwise. This yields the following definition of level- $(k+1)$ players:

$$\begin{aligned} \sigma_{k+1} &= \text{QR}_\lambda(\text{EU}(\rho_{\leq k})) \quad \text{with } \rho_{\leq k} = \sum_{l \leq k} f_\tau^{\leq k}(l) \times \rho_l \\ \rho_{k+1} &= \text{QR}_\lambda(\text{EU}(\sigma_{\leq k})) \quad \text{with } \sigma_{\leq k} = \sum_{l \leq k} f_\tau^{\leq k}(l) \times \sigma_l \end{aligned}$$

³Normally one would specify prior probabilities of states, but we assume that all referents are (believed to be) equiprobable. One would also normally specify utilities, but since we assume interlocutors want to cooperatively identify the referent, utilities are given by identity matrices that cancel out where normally they'd be relevant.

⁴As for notation, if A is a matrix, let $\mathbf{T}(A)$ be its transpose, and $\mathbf{N}(A)$ its row-normalization. If A and C are matrices, AC is their matrix product. We will also use a non-standard operation on matrix A , namely $\max \text{row}(A)$ which returns a binary matrix with the same dimensions as A , such that $\max \text{row}(A)_{ij} = 1$ if $A_{ij} = \max_k A_{ik}$ and 0 otherwise. We abuse notation by assuming that vectors are implicitly coerced if combined with matrices in standard arithmetic operations. So, for instance, $B - c$ is obtained by subtracting c in each row.

⁵The quantal response function is also known as *logit choice rule*, as *soft-max* function (Sutton & Barto, 1998) or, if $\lambda = 1$ as *Luce's choice rule* (Luce, 1959).

Given λ and τ , the model's behavioral prediction is the pair of strategies $\sigma^* = \sum_{k=1}^{\infty} f_{\tau}(k) \times \sigma_k$ and $\rho^* = \sum_{k=1}^{\infty} f_{\tau}(k) \times \rho_k$.

Model fitting As stated above, we fitted a mixed effects logistic regression to participants' behavior in the cost estimation task to estimate the difference x between the log odds of costly vs. cheap message. Assuming that x is the result of a quantal choice rule, we can compute the average subjective costs for a given fixed λ as $c = x/2\lambda$.⁶ Using this, we determined a pair of parameters λ and τ separately for the data on comprehension (Exps 1 and 3) and production (Exp. 2) using a least squares regression ($\lambda = 4.825$, $\tau = 0.625$, $r = 0.99$ for comprehension; $\lambda = 8.853$, $\tau = 0.818$, $r = 0.99$ for production). The prediction-to-data plot is given in Figure 3c.

These results are interesting in many respects. First, they serve as a proof-of-concept that a rather general game theoretic framework can predict behavioral data on language use and interpretation rather well. Second, the small but non-negligible τ indicates that participants in our experiment were able to perform one but not necessarily more steps of best response reasoning including considerations of production costs. Third, production behavior is better explained by higher values of λ and τ . This suggests that the model of Frank and Goodman (2012), which assumes that listeners perform two steps of optimization, while speakers perform exactly one, might be too inflexible. More relevant data is pending, but at present the more general model of Bergen et al. (2012) or the IQR model seem more realistic.

Conclusion

The empirical contributions of this paper are two-fold. First, we provided evidence that language users can draw simple *ad hoc* scalar inferences in an artificial language paradigm. Second, we showed that even when the cost of an unambiguous message is only implicit (i.e., without telling participants explicitly that a message has a certain dollar value (Bergen et al., 2012; Rohde et al., 2012)), scalar inferences were drawn with increased frequency as costs of competing unambiguous messages increased. The theoretical value of this paper lies in the proposal for a probabilistic model of pragmatic reasoning that synthesizes previous Bayesian and game theoretic approaches. The model provides a good fit to the choice distributions of both speakers and listeners; and within listeners, for message cost effects, thus constituting a powerful model of pragmatic inference.

We conclude that not only do listeners take into account available utterances a speaker could have made, they also maintain a gradient estimate of production cost and take this into account in interpretation.

Acknowledgments

We thank M. Tanenhaus, D. Kleinschmidt, and C. Kurumada and our three anonymous reviewers. This work was partially

supported by a EURO-XPRAG grant awarded to the authors and NIH grant HD-27206 to M. Tanenhaus.

References

- Bergen, L., Levy, R., & Goodman, N. D. (2012). That's what she (could have) said: How alternative utterances affect language use. In *Proceedings of the 34th annual meeting of the cognitive science conference*.
- Camerer, C. F., Ho, T.-H., & Chong, J.-K. (2004). A cognitive hierarchy model of games. *The Quarterly Journal of Economics*, 119(3), 861-898.
- Clark, H. H., & Marshall, C. R. (1981). Definite reference and mutual knowledge. In A. K. Joshi, B. Webber, & I. A. Sag (Eds.), *Elements of discourse understanding* (p. 10-63). Cambridge University Press.
- Degen, J., & Franke, M. (2012). Optimal Reasoning About Referential Expressions. In S. Brown-Schmidt, J. Ginzburg, & S. Larsson (Eds.), *Proceedings of SemDial* (p. 2-11).
- Frank, M. C., & Goodman, N. D. (2012). Predicting pragmatic reasoning in language games. *Science*, 336(6084), 998.
- Franke, M. (2011). Quantity implicatures, exhaustive interpretation, and rational conversation. *Semantics & Pragmatics*, 4(1), 1-82.
- Grice, H. (1975). Logic and conversation. *Syntax and Semantics*, 3, 41-58.
- Hanna, J., Tanenhaus, M. K., & Trueswell, J. C. (2003). The effects of common ground and perspective on domains of referential interpretation. *Journal of Memory and Language*, 49, 43-61.
- Jaeger, T. F., & Tily, H. (2011). On language 'utility': processing complexity and communicative efficiency. *Wiley Interdisciplinary Reviews: Cognitive Science*, 2(3), 323-335.
- Jäger, G. (2013). Rationalizable signaling. *Erkenntnis*.
- Keysar, B., Barr, D. J., & Brauner, J. S. (2000). Taking perspective in conversation: The role of mutual knowledge in comprehension. *Psychological Science*, 11, 32-37.
- Lewis, D. (1969). *Convention. a philosophical study*. Harvard University Press.
- Luce, D. R. (1959). *Individual choice behavior: A theoretical analysis*. New York: Wiley.
- Rogers, B. W., Palfrey, T. R., & Camerer, C. (2009). Heterogeneous quantal response equilibrium and cognitive hierarchies. *Journal of Economic Theory*, 144(4), 1440-1467.
- Rohde, H., Seyfarth, S., Clark, B., Jäger, G., & Kaufmann, S. (2012). Communicating with Cost-based Implicature: a Game-Theoretic Approach to Ambiguity. In *Proceedings of SemDial 16* (pp. 107-116).
- Stiller, A., Goodman, N. D., & Frank, M. C. (2011). Ad-hoc scalar implicature in adults and children. In *Proceedings of the 33rd annual conference of the cognitive science society*.
- Sutton, R. S., & Barto, A. G. (1998). *Reinforcement learning*. MIT Press.

⁶This is because $x = \log \frac{P(\text{costly})}{P(\text{cheap})} - \log \frac{P(\text{cheap})}{P(\text{costly})}$ reduces to $x = \log(\exp(-\lambda c)) - \log \frac{1}{\exp(-\lambda c)} = -2\lambda c$.