

When in Rome, do as Bayesians do: Statistical learning and parochial norms

Scott Partington^{1,2} (ssp95@cornell.edu)

Shaun Nichols¹ (sbn44@cornell.edu)

Tamar Kushnir² (tk397@cornell.edu)

¹Cornell University, Department of Philosophy, Goldwin Smith Hall, Ithaca, NY 14853 USA

²Cornell University, Department of Human Development, Beebe Hall, Ithaca, NY 14853 USA

Abstract

It's a familiar point in anthropology that many norms are *parochial*, meaning they apply to people in certain groups (e.g., one's ingroup) and not to others (e.g., one's outgroup). One explanation for such parochialism is that people are just innately biased against outsiders. But it's also possible that, given the evidence, people infer the parochiality of norms in statistically appropriate ways. This paper uses a Bayesian learning framework to investigate inferences of normative scope both experimentally and computationally. An experiment in which adult participants ($n = 480$) viewed sample violations of a novel rule among novel groups reveals that both sensitivity to statistical evidence and prior knowledge of relevant social categories are integral to computations of normative scope. In tandem with the experimental results, computational analysis supports the notion that degree of prior inclusivity bias (i.e., an expectation that a norm will be broad, rather than narrow, in scope) is another key factor. Together, these novel insights raise intriguing possibilities for integrating perspectives on norms research.

Keywords: statistical learning, norms, moral psychology, Bayesian inference, computational modeling

Introduction

"When in Rome, do as the Romans do; when elsewhere, do as they do elsewhere."

—proverb attributed to Saint Augustine

Norms can be understood as inclusive or parochial: "Do not harm others" broadly applies to many groups, whereas "Do not harm your ingroup" is narrower. However, it is rarely the case that such boundaries of normative scope are explicitly mapped out. This poses an inductive challenge to naive observers: how does one learn to which group a candidate norm applies? One possible answer comes from findings in social cognition that document seemingly automatic, group-based biases in normative inference (e.g., Dunham, 2018; Roberts, Gelman, & Ho, 2017a). However, another possibility is that, given the evidence, learners infer the scope of norms in statistically appropriate ways.

The present study investigates the role of statistical learning in inferences of normative scope. We propose that both sensitivity to statistical evidence and knowledge of relevant social categories are integral to computations of normative scope. We begin by reviewing the evidence suggesting a central role for automatic group biases in

normative inference, and then detail how a Bayesian inference framework can offer nuance to such accounts.

A consistent theme in social cognition research is the extent to which social category knowledge is a deeply ingrained and highly influential factor in cognition. Indeed, among the field's most striking findings are those which detail the influence of social category knowledge on judgment and behavior under minimal conditions (for an extensive review, see Dunham, 2018). Studies involving the minimal group paradigm (Tajfel, 1970) randomly assign isolated individuals to previously unfamiliar social groups based on arbitrary cues (e.g., a label, "the red group"). Such manipulations elicit ingroup biases in a variety of domains relevant to normative judgment, and many of these tendencies are early-emerging. For example, studies involving the minimal group paradigm have shown that children allocate more resources to ingroup members (Sparks, Schinkel, & Moore, 2017), make more positive trait evaluations of ingroup members (Richter, Over, & Dunham, 2016), and demonstrate greater trust in the testimony of ingroup members (MacDonald, Schug, Chase, & Barth, 2013). This body of work suggests that learners use the available evidence (e.g., perceptual, testimonial, etc.) to rapidly form representations of social categories which in turn shapes judgments about ingroup-relevant norms in a heuristic-like fashion.

In a series of recent studies, Roberts and colleagues (Roberts, Gelman, & Ho, 2017a; Roberts, Gelman, & Ho, 2017b; Roberts, Guo, Ho, & Gelman, 2018) examine a key cue that informs such heuristics by showing that children infer what groups ought to do from descriptions of general group behavior. In this work, Roberts and colleagues use a paradigm whereby participants are introduced to two novel groups, labelled "Hibbles" and "Glerks," who are characterized in terms of morally neutral regularities (e.g., eating a certain kind of berry). If told that Glerks eat green berries and Hibbles eat red berries, participants will say that it is "not okay" when a Glerk eats a red berry. Note that by introducing the novel groups paradigm, Roberts and colleagues probe the influence of *mere group*, as opposed to *ingroup*, representations on normative judgment. In explaining these findings, the authors indeed propose a mechanism by which "group regularities may exert influence by rather automatically fostering a negative evaluative stance," to non-conformity (2017a, p. 593), suggesting that

an automatic group bias influences children's judgments about norms.

Although such findings seem to provide a plausible account of processes likely implicated in normative inference, we propose that incorporating key insights from statistical learning research would offer nuance to such an account. We'll next outline how Bayesian inference provides a rational framework for inferences of normative scope, as well as when and why the Bayesian framework makes diverging predictions from accounts which solely emphasize the role of automatic group biases.

Over the past decade, Bayesian theories of learning have been productive and influential across a number of domains: casual learning (e.g., Gopnik & Wellman, 2012), category discrimination (e.g., Kemp, Perfors, & Tenenbaum, 2007), language acquisition (e.g., Xu & Tenenbaum, 2007) and social inference (e.g., Lucas et al., 2014) among others. Such an approach is especially relevant and compelling here, too, because Bayesian inference provides a rational basis for inferences of normative scope. Indeed, previous work in philosophy has proposed theoretical accounts of Bayesian norm learning (e.g., Colombo, 2013; Muldoon, Lisciandra, & Hartmann, 2014; Nichols, forthcoming), and here we put a simple model to empirical test. Under the Bayesian framework, determining the scope of a norm involves assessing the degree to which the relevant evidence is consistent with competing hypotheses about that norm. In this case, the competing hypotheses are characterized by parochiality (i.e., the norm applies narrowly) and inclusivity (i.e., the norm applies broadly). Such candidate hypotheses are modeled as structured, symbolic representations to which learners assign different levels of certainty given the available evidence (c.f. Goodman, Tenenbaum, Feldman, & Griffiths, 2008).

A key feature of the Bayesian model is the *size principle* (e.g., Perfors, Tenenbaum, Griffiths, & Xu, 2011; Tenenbaum & Griffiths, 2001), which dictates that when all observed evidence is consistent with a smaller hypothesis, that hypothesis should be preferred. To borrow an intuitive example from Nichols and colleagues (Nichols, Kumar, Lopez, Ayars, & Chan, 2016), a sequence of dice rolls 1, 1, 3, 2, 1, 2, 4, 1 could have been generated from a 4-sided dice (H_4) or a 10-sided dice (H_{10}). However, it seems much more likely that H_4 is true, since all observed rolls are less than or equal to 4. In other words, if H_{10} were in fact true, this would be a highly suspicious coincidence in light of the available evidence. The size principle formally captures this intuitive fact: since H_4 is consistent with a smaller set of possible observations, it should be preferred.

Moving back to the case of normative inference, we can represent the size of the competing hypotheses in a nested structure. As shown in Figure 1, a norm being parochial in scope is consistent with a smaller set of subjects than the norm being inclusive in scope. Thus, if we observe members of Group A and Group B engaging in the same behavior, yet only members of Group B are ever sanctioned for violating the norm, the size principle dictates that we should infer the

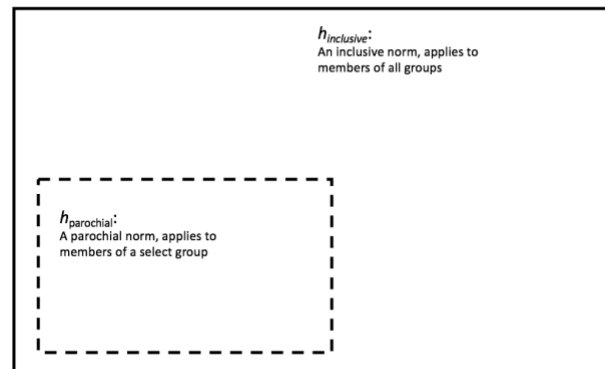


Figure 1: Hypotheses about the potential scope of a norm represented in a subset structure. The hypothesis space picks out individuals who are subject to the candidate norm.

norm narrowly applies to Group B. Note that the degree to which the evidence supports a parochial inference is dependent on the size of Group B relative to total population proportion (cf., Kushnir, Xu, & Wellman, 2010). If all sample violations come from a relatively small minority group, Bayesian learners should infer the norm is parochial in scope. However, if all sample violations come from a relatively large majority group, Bayesian learners should be more likely to infer the norm is inclusive in scope. In contrast, if an automatic group bias based on behavioral regularity determines such inferences of scope, then generalization should be similar in both scenarios, since behavioral regularity is held constant. Thus, incorporating elements of Bayesian inference stands to offer nuance to accounts of normative inference by predicting when and why observed group regularities will lead to parochial or inclusive inferences.

In the present study, we put the Bayesian account to test. Using a novel rule learning paradigm (c.f., Ayars & Nichols, 2017; Nichols, Kumar, Lopez, Ayars, & Chan, 2016) in conjunction with a novel groups paradigm (c.f. Roberts, Gelman, & Ho, 2017a), we hold behavioral regularity constant while varying the size of the target group relative to the total population. As predicted by the size principle, we expect to find parochial generalizations of the rule when sample violations come from small minority groups, but inclusive generalizations when sample violations come from large majority groups.

Experiment

Participants

Adult participants ($n = 480$; 31.9% female, 67.5% male, 0.6% other; $M_{Age} = 35.5$ years, $SD = 10.6$) were recruited from Amazon MTurk to complete a survey for modest compensation. An additional 77 participants were excluded from analyses for failing to complete the survey. All participants included in analyses completed the entire survey and passed a series of comprehension checks.

Procedure

We presented participants with a scenario in which two groups (labelled “Hibbles” and “Glerks”) live together on an island (in all but the 100% condition, see below). Participants were randomly assigned to one of four conditions. Across conditions, the size of the target group relative to the island’s total population was varied (i.e., 20%, 50%, 80%, or 100% of approximately 100 total inhabitants). In all conditions, a fixed proportion of each group (approximately 35%) was shown wearing a morally neutral item of clothing (e.g., Trial 1: ribbons, Trial 2: hats). To provide a concrete example: for Trial 1, participants in the 20% condition would view an island inhabited by 20 Glerks and 80 Hibbles (100 individuals in total), with 7 of the Glerks wearing a ribbon and 28 of the Hibbles wearing a ribbon (35% of the individuals in each group).

Next, participants were told the island has rules and that their task would be to figure out one of the island’s rules. Participants then watched a video highlighting a sample of 4 members of the target group (e.g., “Hibbles”) wearing the clothing item (e.g., a ribbon) as violations (“This is against the rule.”). Afterwards, participants were asked to infer if other individuals on the island were also violating the rule. We solicited judgments about all possible group member/clothing item combinations (order counterbalanced): another target group member with a ribbon, a target group member without a ribbon, a non-sampled group member with a ribbon, and a non-sampled group member without a ribbon. For each individual, participants made a forced-choice judgment of whether “This is against the rule” or “This is not against the rule.” Next, participants made the same judgment about a visitor to the island who was wearing a ribbon. The visitor was called a “Zorg” and its body was purple and spiky.

Participants also articulated their understanding of the rule in an open response item (“What is the rule?”) and provided confidence ratings of this understanding (a 7-point scale, “How confident are you that you know the rule?” with 1 = *Not at all*, 7 = *Very*). Participants then repeated the entire procedure in Trial 2, which was identical to Trial 1 except for the individuals wearing a different clothing item (e.g., hats).

Coding

We scored participants’ judgments as “This is against the rule” = 1 and “This is not against the rule” = 0. For the open response question, responses were coded for whether participants articulated a parochial rule (e.g., “Glerks cannot wear hats”) or an inclusive rule (e.g., “No hats allowed”). Responses that did not articulate a rule (e.g., “I just guessed”) were not counted. Responses that articulated a rule were scored as either ‘parochial’ = 0 or ‘inclusive’ = 1 by a coder who did not know from which condition the responses originated.

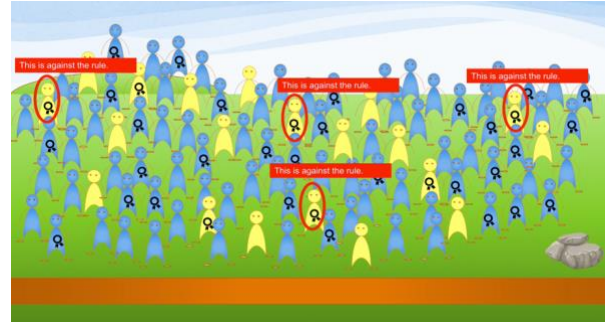


Figure 2: Example of stimuli as presented to participants in the 20% condition. The red text boxes above the sampled individuals read, “This is against the rule.”

Results

Were judgments of normative scope sensitive to statistical evidence? For each group member/clothing item combination, we ran a logistic regression model with judgment score as the dependent variable and condition, trial number, and clothing item as independent variables. As expected, the results suggest that judgments of normative scope were sensitive to statistical evidence. When generalizing the rule to individuals from the non-sampled group ($\beta = 0.640$, $SE = 0.070$, $p < .001$) and the visitor ($\beta = 0.427$, $SE = 0.067$, $p < .001$), the results varied as a function of the relative population proportions. Participants did not apply the rule to the target group at significantly different rates across conditions ($\beta = -0.173$, $SE = 0.101$, $p = .09$).

Post-hoc analyses confirmed this pattern (see Figure 3). In each condition, participants applied the rule most frequently to members of the target group ($M_{\text{Target}} = 0.90, 0.87, 0.92, 0.83$ in the 20%, 50%, 80%, and 100% conditions, respectively), whereas participants were most likely to think the rule applied narrowly when sample violations came from a 20% minority ($M_{\text{Non-sampled}} = 0.28$, $SD_{\text{Non-sampled}} = 0.45$, $M_{\text{Visitor}} = 0.42$, $SD_{\text{Visitor}} = 0.49$), followed by the 50% condition ($M_{\text{Non-sampled}} = 0.39$, $SD_{\text{Non-sampled}} = 0.49$, $M_{\text{Visitor}} = 0.48$, $SD_{\text{Visitor}} = 0.50$) followed by the 80% condition ($M_{\text{Non-sampled}} = 0.48$, $SD_{\text{Non-sampled}} = 0.50$, $M_{\text{Visitor}} = 0.53$, $SD_{\text{Visitor}} = 0.50$). Finally, when the population contained only one group (the 100% condition), participants generalized the rule to all novel individuals ($M_{\text{Non-sampled}} = 0.71$, $SD_{\text{Non-sampled}} = 0.46$, $M_{\text{Visitor}} = 0.72$, $SD_{\text{Visitor}} = 0.45$).

Next, we ran a logistic regression model with open response score as DV and condition as independent variable. As with the judgment scores, open response scores also varied as a function of condition ($\beta = 0.726$, $SE = 0.077$, $p < .001$). Once again, post-hoc analyses show participants most frequently articulated a parochial rule in the 20% condition ($M = 0.30$, $SD = 0.46$) and the 50% condition ($M = 0.33$, $SD = 0.47$), followed by the 80% condition ($M = 0.50$, $SD = 0.50$), and lastly the 100% condition ($M = 0.78$, $SD = 0.41$). Thus, participants’ own articulation of the rule provides further evidence that judgments of normative scope are sensitive to statistical evidence. Participants were most likely to articulate a parochial rule when sample violations

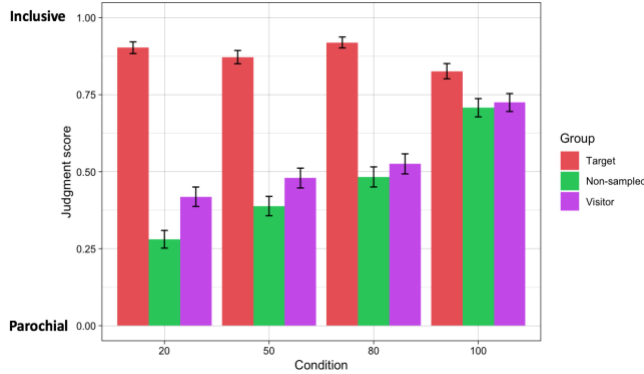


Figure 3: Mean judgment score for each candidate group member. Lower bars indicate subjects more frequently judge the rule to be parochial in scope. Error bars represent the standard error.

were drawn from a minority or equal group, and participants most likely to articulate an inclusive rule when sample violations were drawn from a homogenous population.

Confidence ratings were high on average ($M = 4.73$, $SD = 1.63$). Confidence ratings differed significantly between the 20% condition ($M = 4.38$, $SD = 1.74$) and the 100% condition ($M = 5.05$, $SD = 1.66$) (20% vs. 100%: $t(239) = -3.1$, $p = 0.006$, $d = 0.40$). There was no significant difference between ratings in the 20% condition, the 50% condition ($M = 4.77$, $SD = 1.51$), and the 80% condition ($M = 4.71$, $SD = 1.55$), nor was there a significant difference in ratings between the 50%, 80%, and 100% conditions.

Computational analysis

While the main results are broadly consistent with the proposed Bayesian framework, a formal computational analysis can provide further insight regarding the extent to which such inferences are rational.

Formally, the inference can be defined as learning a rule R from a set of examples D from some known domain U , where $D = d_1, \dots, d_n$. The proposed model assumes the learner has access to a hypothesis space H that contains the set of candidate hypotheses for representing the rule R and a probabilistic model that relates $h \in H$ to the evidence D . The Bayesian learner computes the posterior probabilities $p(h|D)$ for different hypotheses $h \in H$, using Bayes Rule:

$$p(h|D) = \frac{p(D|h)p(h)}{\sum_{h' \in H} p(D|h')p(h')}$$

We can further specify this model with the learner's assumption that violations are sampled at random and independently from the true rule to be learned. This results in the following likelihood function, which constitutes a formal instantiation of the size principle:

$$p(D|h) = \left[\frac{1}{\text{size}(h)} \right]^n$$

where n is the number of violations observed. In order to model this inference, four key assumptions must be made.

First, H is assumed to contain only two hypotheses relating to the rule: $h_{\text{inclusive}}$, the hypothesis the rule applies to all observed ribbon-wearers, and $h_{\text{parochial}}$, the hypothesis the rule only applies to ribbon-wearers from the target group.

Along these lines, the second assumption is that $h_{\text{parochial}} \subseteq h_{\text{inclusive}}$, reflecting the nested structure as specified previously. This allows us to characterize $\text{size}(h_{\text{inclusive}}) = 1$ and express $\text{size}(h_{\text{parochial}})$ as relative percentage of $\text{size}(h_{\text{inclusive}})$.

Third, it is assumed that the rule can only be applied to the observed population of ribbon-wearers (as participants in our study are made to believe). Correspondingly, we can set $\text{size}(h_{\text{parochial}}) = .20$ in the 20% condition, $\text{size}(h_{\text{parochial}}) = .50$ in the 50% condition, $\text{size}(h_{\text{parochial}}) = .80$ in the 80% condition, and $\text{size}(h_{\text{parochial}}) = 1$ in the 100% condition.

Fourth and finally, the prior degree of belief in the rule being inclusive, as opposed to parochial, can be expressed by the ratio of $p(h_{\text{inclusive}})$ to $p(h_{\text{parochial}})$. Thus, for example, prior to observing any sample violations an unbiased learner would have a prior ratio of 1:1, whereas a learner who believes an inclusive rule is twice as likely as a parochial rule has a prior ratio of 2:1. In this simple case, the ratio, rather than exact values, is what matters because there are only two candidate hypotheses (h is $h_{\text{inclusive}}$ and h' is $h_{\text{parochial}}$, if you will), so Bayes Rule reduces the two priors to a ratio when the respective likelihoods are held constant, as is the case here.

With these assumptions in mind, we can model the inference under the different population proportions by setting the sampled violations $n = 4$ and varying inputs for $\text{size}(h_{\text{parochial}})$ and the prior ratio.

As shown in Figure 4, an idealized, unbiased Bayesian learner (prior ratio = 1:1) will always prefer $h_{\text{parochial}}$ as long as it is smaller than $h_{\text{inclusive}}$, as predicted by the size principle. Thus, Bayesian learners must have some degree of *prior inclusivity bias* in order to favor $h_{\text{inclusive}}$. Indeed, such an inclusivity bias is consistently reflected in our experimental results. In comparison to the unbiased Bayesian learner, participants were more likely to infer an inclusive rule across all conditions.

When we consider Bayesian learners with varying degrees of inclusivity bias, a second key trend emerges. No matter the degree of inclusivity bias, Bayesian learners should favor $h_{\text{parochial}}$ when $\text{size}(h_{\text{parochial}}) < 25\%$ of $\text{size}(h_{\text{inclusive}})$. On the flip side, Bayesian learners with inclusivity bias should favor $h_{\text{inclusive}}$ when $\text{size}(h_{\text{parochial}}) > 85\%$ of $\text{size}(h_{\text{inclusive}})$. This trend is also broadly reflected in our experimental results: the majority (71%) of participants in the 20% condition inferred a parochial rule, and the majority (70%) of participants in the 100% condition inferred an inclusive rule.

For the regions in between, when $\text{size}(h_{\text{parochial}}) < 85\%$ of $\text{size}(h_{\text{inclusive}})$ and $> 25\%$ of $\text{size}(h_{\text{inclusive}})$, the degree of inclusivity bias begins to largely differentiate which learners will favor a parochial or inclusive rule. As such, experimental results in the 50% and 80% conditions are consistent with

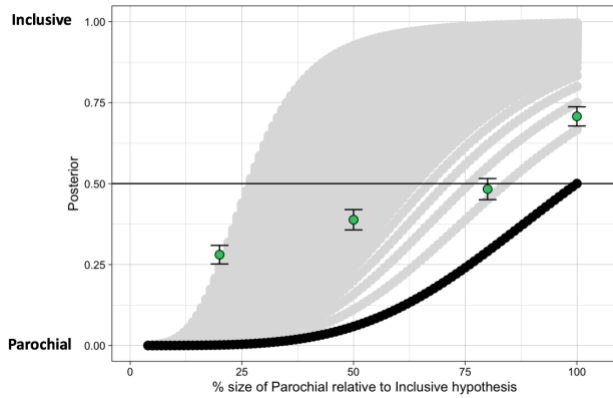


Figure 4: Computed posteriors for $h_{\text{inclusive}}$ (y-axis) plotted against relative size of $h_{\text{parochial}}$ to $h_{\text{inclusive}}$ (x-axis). The black points correspond to an idealized, unbiased Bayesian learner (prior ratio = 1:1). The gray points correspond to Bayesian learners with varying degrees of prior inclusivity bias (range of prior ratios, lowest to highest: 2:1 to 200:1). The green points correspond to the observed percentage of inclusive judgments in each condition, with error bars representing standard error. The horizontal black line denotes when $p(h_{\text{inclusive}}|D) = .50$.

moderate inclusivity bias: 61% of participants in the 50% condition and 52% of participants in the 80% condition inferred a parochial rule.

Next, a probability of generalization function can be specified to predict the results regarding the visitor. Formally, the learner must decide whether any given new individual z belongs to the extension of R , given the observations of D . Thus, a learner must compute a probability of generalization by averaging the predictions of all hypotheses weighted by their posterior probabilities:

$$p(z \in R|D) = \sum_{h \in H} p(z \in R|h)p(h|D)$$

We can model this computation by using our observed posteriors for $h_{\text{inclusive}}$ and $h_{\text{parochial}}$ (i.e., the actual percentage of participants who inferred an inclusive or parochial rule in each condition) and varying the perceived likelihood of the inclusive rule extending to the visitor (i.e., $p(z \in R|h_{\text{inclusive}})$). For the parochial rule, this was set to a low constant ($p(z \in R|h_{\text{parochial}}) = 0.05$) to reflect psychological plausibility—i.e., people who originally infer a parochial rule are not likely to perceive that rule extends to a new group, though it remains a non-negligible possibility. As shown in Figure 5, when the visitor is deemed highly likely to belong to the same social category as the non-sampled group (i.e., ‘subjects of the rule’), computational results approximate the observed data, albeit with a modest inclusivity bias yet again.

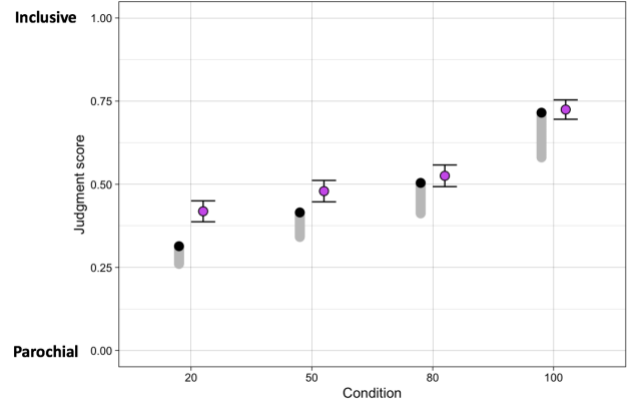


Figure 5: Comparing the expected and observed judgment scores for the visitor in each condition assuming ($p(z \in R|h_{\text{parochial}}) = .05$). The black points correspond to Bayesian learners with maximum certainty that the visitor would be a subject of the rule, if the rule is inclusive ($p(z \in R|h_{\text{inclusive}}) = .99$). The gray lines represent Bayesian learners with a corresponding range of certainties (from .98 to .80). The purple points correspond to the percentage of participants who extended the rule to the visitor in each condition. Error bars represent standard error.

Discussion

One explanation for parochialism is that people are just innately biased against outsiders. However, it’s also possible that, given the available evidence, people infer the parochiality of norms in statistically appropriate ways. Although a great amount of research details the ways in which automatic group biases can influence normative judgment, experimental results here indicate that inferences of normative scope are sensitive to statistical evidence in a manner consistent with Bayesian inference. When sample violations came from a minority group, participants tended to infer parochiality. When sample violations came from a majority group, participants were more likely to infer the norm applied inclusively. This pattern of results suggests that components of rational statistical inference can indeed play a role in normative learning. Formal Bayesian analysis broadly supports this conclusion and further illuminates that the degree of prior inclusivity bias, or favoring a broad generalization of the rule, is an additional key factor.

Combining insights from these empirical and computational perspectives raises intriguing open questions for further investigation. As mentioned, we designed our experiment based on a novel groups paradigm (no real-world social groups) with no role for personal identity (no inclusion of participants as group members). For normative learning in everyday contexts, it is plausible that knowledge of the groups’ typical characteristics and/or one’s own personal identity play a role in judgments of normative scope. There is a clear need for further research to investigate how manipulating these key features effects learning from statistical evidence by shifting expectations about the scope

of the candidate norm. For example, people may be more likely to expect norms are inclusive when considering population-level evidence (such as the evidence in the present study) and more likely to expect norms are parochial when considering essentialist exemplars (such as the evidence provided in Roberts and colleagues' work). Such a shift in expectations about which hypothesis to favor under a given circumstance can be expressed in terms of *overhypotheses* and thus constitutes a possible extension of the Bayesian framework proposed here.

Thus, the proposed Bayesian framework merits future empirical testing and theoretical refinement. Important steps forward include testing the framework under a greater range of experimental manipulations, building from the novel groups paradigm used here, as well as testing the framework in developmental contexts and under conditions of higher external validity. For now, the present findings suggest that certain key components of normative inference can indeed be statistically appropriate, given the evidence available to learners.

Acknowledgments

This research was supported by funding from the Evelyn Edwards Milman grant to T.K. Thanks to three anonymous reviewers and members of the Early Childhood Cognition Lab for helpful comments.

References

- Ayars, A., & Nichols, S. (2017). Moral empiricism and the bias for act-based rules. *Cognition*, 167, 11–24. <https://doi.org/10.1016/j.cognition.2017.01.007>
- Dunham, Y. (2018). Mere Membership. *Trends in Cognitive Sciences*, 22(9), 780–793. <https://doi.org/10.1016/j.tics.2018.06.004>
- Colombo, M. (2014). Two neurocomputational building blocks of social norm compliance. *Biology & Philosophy*, 29(1), 71–88. <https://doi.org/10.1007/s10539-013-9385-z>
- Goodman, N. D., Tenenbaum, J. B., Feldman, J., & Griffiths, T. L. (2008). A Rational Analysis of Rule-Based Concept Learning. *Cognitive Science*, 32(1), 108–154. <https://doi.org/10.1080/03640210701802071>
- Gopnik, A., & Wellman, H. M. (2012). Reconstructing constructivism: Causal models, Bayesian learning mechanisms, and the theory theory. *Psychological Bulletin*, 138(6), 1085–1108. <https://doi.org/10.1037/a0028044>
- Kemp, C., Perfors, A., & Tenenbaum, J. B. (2007). Learning overhypotheses with hierarchical Bayesian models. *Developmental Science*, 10(3), 307–321. <https://doi.org/10.1111/j.1467-7687.2007.00585.x>
- Kushnir, T., Xu, F., & Wellman, H. M. (2010). Young children use statistical sampling to infer the preferences of other people. *Psychological Science*, 21(8), 1134–1140.
- Lucas, C. G., Griffiths, T. L., Xu, F., Fawcett, C., Gopnik, A., Kushnir, T., Markson, L., & Hu, J. (2014). The child as econometrician: A rational model of preference understanding in children. *PloS One*, 9(3), e92160.
- MacDonald, K., Schug, M., Chase, E., & Barth, H. (2013). My people, right or wrong? Minimal group membership disrupts preschoolers' selective trust. *Cognitive Development*, 28(3), 247–259. <https://doi.org/10.1016/j.cogdev.2012.11.001>
- Muldoon, R., Lisciandra, C., & Hartmann, S. (2014). Why are there descriptive norms? Because we looked for them. *Synthese*, 191(18), 4409–4429. <https://doi.org/10.1007/s11229-014-0534-y>
- Nichols, S., Kumar, S., Lopez, T., Ayars, A., & Chan, H.-Y. (2016). Rational Learners and Moral Rules. *Mind & Language*, 31(5), 530–554. <https://doi.org/10.1111/mila.12119>
- Nichols, S. (forthcoming). *Rational Rules: Toward a Theory of Moral Learning*. Oxford University Press.
- Perfors, A., Tenenbaum, J. B., Griffiths, T. L., & Xu, F. (2011). A tutorial introduction to Bayesian models of cognitive development. *Cognition*, 120(3), 302–321. <https://doi.org/10.1016/j.cognition.2010.11.015>
- Richter, N., Over, H., & Dunham, Y. (2016). The effects of minimal group membership on young preschoolers' social preferences, estimates of similarity, and behavioral attribution. *Collabra*. <https://doi.org/10.1525/collabra.44>
- Roberts, S. O., Gelman, S. A., & Ho, A. K. (2017). So It Is, So It Shall Be: Group Regularities License Children's Prescriptive Judgments. *Cognitive Science*, 41(S3), 576–600. <https://doi.org/10.1111/cogs.12443>
- Roberts, S. O., Guo, C., Ho, A. K., & Gelman, S. A. (2018). Children's descriptive-to-prescriptive tendency replicates (and varies) cross-culturally: Evidence from China. *Journal of Experimental Child Psychology*, 165, 148–160. <https://doi.org/10.1016/j.jecp.2017.03.018>
- Roberts, S. O., Ho, A. K., & Gelman, S. A. (2017). Group presence, category labels, and generic statements influence children to treat descriptive group regularities as prescriptive. *Journal of Experimental Child Psychology*, 158, 19–31. <https://doi.org/10.1016/j.jecp.2016.11.013>
- Sparks, E., Schinkel, M. G., & Moore, C. (2017). Affiliation affects generosity in young children: The roles of minimal group membership and shared interests. *Journal of Experimental Child Psychology*, 159, 242–262. <https://doi.org/10.1016/j.jecp.2017.02.007>
- Tajfel, H. (1970). Experiments in Intergroup Discrimination. *Scientific American*, 223(5), 96–103. JSTOR.
- Tenenbaum, J. B., & Griffiths, T. L. (2001). Generalization, similarity, and Bayesian inference. *Behavioral and Brain Sciences*, 24(4), 629–640. <https://doi.org/10.1017/S0140525X01000061>
- Xu, F., & Tenenbaum, J. B. (2007). Word learning as Bayesian inference. *Psychological Review*, 114(2), 245–272. <https://doi.org/10.1037/0033-295X.114.2.245>