

Psychological Science

<http://pss.sagepub.com/>

Political Extremism Is Supported by an Illusion of Understanding

Philip M. Fernbach, Todd Rogers, Craig R. Fox and Steven A. Sloman

Psychological Science published online 25 April 2013

DOI: 10.1177/0956797612464058

The online version of this article can be found at:

<http://pss.sagepub.com/content/early/2013/04/24/0956797612464058>

Published by:



<http://www.sagepublications.com>

On behalf of:



[Association for Psychological Science](#)

Additional services and information for *Psychological Science* can be found at:

Email Alerts: <http://pss.sagepub.com/cgi/alerts>

Subscriptions: <http://pss.sagepub.com/subscriptions>

Reprints: <http://www.sagepub.com/journalsReprints.nav>

Permissions: <http://www.sagepub.com/journalsPermissions.nav>

>> [OnlineFirst Version of Record](#) - Apr 25, 2013

[What is This?](#)

Political Extremism Is Supported by an Illusion of Understanding

Philip M. Fernbach¹, Todd Rogers², Craig R. Fox^{3,4},
and Steven A. Sloman⁵

¹Leeds School of Business, University of Colorado, Boulder; ²Center for Public Leadership, Harvard Kennedy School; ³Anderson School of Management, University of California, Los Angeles; ⁴Department of Psychology, University of California, Los Angeles; and ⁵Department of Cognitive, Linguistic, and Psychological Sciences, Brown University

Psychological Science

XX(X) 1–8

© The Author(s) 2013

Reprints and permissions:

sagepub.com/journalsPermissions.nav

DOI: 10.1177/0956797612464058

pss.sagepub.com



Abstract

People often hold extreme political attitudes about complex policies. We hypothesized that people typically know less about such policies than they think they do (the illusion of explanatory depth) and that polarized attitudes are enabled by simplistic causal models. Asking people to explain policies in detail both undermined the illusion of explanatory depth and led to attitudes that were more moderate (Experiments 1 and 2). Although these effects occurred when people were asked to generate a mechanistic explanation, they did not occur when people were instead asked to enumerate reasons for their policy preferences (Experiment 2). Finally, generating mechanistic explanations reduced donations to relevant political advocacy groups (Experiment 3). The evidence suggests that people's mistaken sense that they understand the causal processes underlying policies contributes to political polarization.

Keywords

explanation, illusion of explanatory depth, political psychology, attitudes, polarization, extremism, moderation, public policy, causal models, mechanism, causality, policymaking, decision making, judgment

Received 6/11/12; Revision accepted 9/17/12

The opinions that are held with passion are always those for which no good ground exists.

—Bertrand Russell (1928/1996, p. 3)

Extremism is so easy. You've got your position and that's it. It doesn't take much thought.

—Clint Eastwood (quoted in Schickel, 2005)

Many of the most important issues facing society—from climate change to health care to poverty—require complex policy solutions about which citizens hold polarized political preferences. A central puzzle of modern American politics is how so many voters can maintain strong political views concerning complex policies yet remain relatively uninformed about how such policies would bring about desired outcomes (for review, see Delli Carpini & Keeter, 1996).

One possible cause of this apparent paradox is that voters believe that they understand how policies work

better than they actually do. In the research reported here, we explored two questions. First, do people really have unjustified confidence in their understanding of how complex policies work? Second, does this illusion of understanding contribute to attitude polarization? We predicted that asking people to explain how a policy works would make them aware of how poorly they understood the policy, which would cause them to subsequently express more moderate attitudes and behaviors.

Rozenblit and Keil (2002) have demonstrated that people tend to be overconfident in how well they understand how everyday objects, such as toilets and combination locks, work; asking people to generate a mechanistic explanation shatters this sense of understanding (see also

Corresponding Author:

Philip M. Fernbach, University of Colorado, Leeds School of Business, 419 UCB, Boulder, CO 80309-0419

E-mail: philip.fernbach@gmail.com

Alter, Oppenheimer, & Zemla, 2010; Keil, 2003). The attempt to explain makes the complexity of the causal system more apparent, leading to a reduction in judges' assessments of their own understanding. Prior research on the illusion of explanatory depth has focused primarily on feelings of understanding, but this phenomenon is likely to have downstream effects on preferences and behaviors. For instance, consumers' willingness to pay for products is influenced by their perceived understanding of how those products work (Fernbach, Sloman, St. Louis, & Shube, 2013). Moreover, people are more likely to change their attitudes about a policy when they have less confidence in their knowledge about it (Krosnick & Petty, 1995). We conjectured, therefore, that extreme policy preferences often rely on people's overestimation of their mechanistic understanding of the complex systems those policies are intended to influence. If this is true, then merely asking people to generate an explanation of relevant mechanisms should decrease their sense of understanding and subsequently lead them to express more moderate political views.

Our prediction is consistent with research on how the complexity with which people think about an object affects the extremity of their evaluation of that object. For instance, Linville (1982) asked participants to evaluate either six or two dimensions of a chocolate-chip cookie (e.g., chewiness, butteriness, number of chocolate chips). Participants who were induced to think about the cookie complexly by rating six dimensions reported less extreme evaluations of the cookie than did participants who were induced to think about the cookie simply by rating only two dimensions. Related work has shown that more complex representations of the self lead to smaller affective swings in the face of stressful events (Linville, 1985) and less vulnerability to depression and illness (Linville, 1987).

On its surface, our prediction appears to contradict research suggesting that people's attitudes become more extreme when they are asked to justify or deliberate about a position (Hirt & Markman, 1995; Ross, Lepper, Strack, & Steinmetz, 1977; Tesser, 1978). Moreover, political discussions among like-minded people typically lead them to become more extreme in their views (Schkade, Sunstein, & Hastie, 2010). We reconcile these opposing predictions by suggesting that the nature of the elaboration is critical in determining whether it will lead to polarization or moderation. For instance, whereas asking people to provide reasons for their position on a policy may cause them to selectively access a supportive rationale, thereby increasing their commitment to the position, asking them to explain the mechanisms by which the policy works may force them to confront their lack of understanding, thereby decreasing their commitment.

Experiment 1: Effect of Explanation on Understanding and Position Extremity

In our first study, we asked participants to rate how well they understand six political policies. After participants judged their understanding of each issue, we asked them to explain how two of the policies work and then to rerate their level of understanding. We expected that asking participants to explain the mechanisms underlying the policies would expose the illusion of explanatory depth and lead to lower ratings of understanding, extending prior findings (e.g., Alter et al., 2010; Rozenblit & Keil, 2002) to the domain of political attitudes.

We predicted further that exposing the illusion of explanatory depth would lead people to express more moderate support for policies. We tested this prediction in two ways. First, we had one group of participants provide ratings of their positions both before and after they generated mechanistic explanations. We examined how their degree of support changed and how this change was associated with self-rated understanding of relevant mechanisms. Recognizing that this within-subjects comparison could give rise to a demand effect, such that some participants may have felt obliged to report a less extreme judgment after providing a poor explanation, we asked a second group to rate their policy support only after generating explanations. This allowed for a between-participant comparison in which we compared the postexplanation ratings of this second group with the preexplanation ratings of the first group.

Method

Participants and design. One hundred ninety-eight U.S. residents were recruited using Amazon's Mechanical Turk and participated in return for a small payment. Participants were 52% male and 48% female, with an average age of 33.3 years. Participants' reported political affiliations were 40% Democrat, 20% Republican, 36% independent, and 4% other.

In the preexplanation-rating conditions ($n = 87$), participants rated their position on policies both before and after generating mechanistic explanations for them. In the no-preexplanation-rating conditions ($n = 111$), participants rated their position only after generating explanations. Each participant generated mechanistic explanations for two of the six policies. The six policies were blocked into three groups of two each so that there were a total of six conditions to which participants were randomly assigned (three preexplanation-rating conditions and three no-preexplanation-rating conditions).

Materials and procedure. After answering demographic questions, participants in the preexplanation-rating

conditions were asked to state their position on six political policies; responses were made using 7-point scales from 1, *strongly against*, to 7, *strongly in favor*. The policies were (a) imposing unilateral sanctions on Iran for its nuclear program, (b) raising the retirement age for Social Security, (c) transitioning to a single-payer health care system, (d) establishing a cap-and-trade system for carbon emissions, (e) instituting a national flat tax, and (f) implementing merit-based pay for teachers. Participants in the no-preexplanation-rating conditions skipped these initial position ratings.

Next, all participants were trained to use a rating scale to quantify their level of understanding of the policies. Instructions were modeled on instructions used in Rozenblit and Keil (2002), but rather than describing different levels of understanding for an object, our instructions described different levels of understanding for a political issue (immigration reform) that was not included as one of the issues in our experiment (see Rating-Scale Instructions in the Supplemental Material available online). After reading the instructions, participants were asked to judge their level of understanding of the six policies (e.g., "How well do you understand the impact of imposing unilateral sanctions on Iran for its nuclear program?"). Responses were made using a 7-point scale, with higher scores indicating greater understanding.

After judging their understanding of all six policies, participants were asked to provide a mechanistic explanation for one of the six policies. Instructions for this measure were also adapted from Rozenblit and Keil (2002; see Example Instructions for Explanation- and Reason-Generation Tasks in the Supplemental Material). Participants were then asked to rerate their understanding of the policy; to rate or rerate their position on the policy; and to rate how certain they were of their position, using a 5-point scale from 1, *not at all certain*, to 5, *extremely certain*.

After completing these measures, participants repeated the process for a second issue. The policies were blocked such that participants explained either (a) the Iran issue followed by the merit-pay issue, (b) the health care issue followed by the Social Security issue, or (c) the cap-and-trade issue followed by the flat-tax issue.

Results

Understanding. We analyzed judgments of understanding using a repeated measures analysis of variance (ANOVA) with timing of judgment (preexplanation vs. postexplanation) and issue number (first issue vs. second issue) as within-subjects factors. All participants provided both preexplanation and postexplanation ratings of understanding, so these analyses used the full data set. Our first prediction was that we would observe a decrease

in understanding judgments following mechanistic explanation. This prediction was confirmed by a significant main effect of judgment timing: Postexplanation ratings of understanding ($M = 3.45$, $SE = 0.12$) were lower than preexplanation ratings ($M = 3.82$, $SE = 0.11$), $F(1, 197) = 34.69$, $p < .001$, $\eta_p^2 = .15$. We found the same pattern across all six policies. To test whether the effect generalized across stimuli, we collapsed over participants and compared average change in understanding due to explanation across the six policies. This effect was also significant, $t(5) = 5.74$, $p < .01$. There was also an unexpected main effect of issue number, such that participants reported having a better understanding of the first issue than of the second, $F(1, 197) = 76.18$, $p < .001$, $\eta_p^2 = .28$. However, issue number did not interact with judgment timing, $F(1, 197) = 1.45$, $p > .23$.

Position extremity. We transformed raw ratings of positions on policies into a measure of position extremity by subtracting the midpoint of the scale (4) and taking the absolute value. We first compared position-extremity scores before and after explanation for participants in the preexplanation-rating conditions. We conducted a repeated measures ANOVA with timing of judgment (preexplanation vs. postexplanation) and issue number (first issue vs. second issue) as within-subjects factors. We predicted that positions would become more moderate following explanation. This prediction was confirmed, with the main effect of judgment timing significant (preexplanation-rating conditions: $M = 1.41$, $SE = 0.07$; postexplanation-rating conditions: $M = 1.28$, $SE = 0.08$), $F(1, 86) = 6.10$, $p = .016$, $\eta_p^2 = .066$. As with understanding, the pattern for position extremity was the same across all six policies, and the test of the moderation effect over the six policies was significant, $t(5) = 3.93$, $p = .011$. However, two of the policies (merit pay and Social Security) showed very small differences. Also consistent with our findings regarding judgments of understanding, results revealed an unexpected main effect of issue number, such that extremity scores for the first issue were lower than those for the second, $F(1, 86) = 10.10$, $p < .01$, $\eta_p^2 = .11$. Again, issue number did not interact with judgment timing, $F(1, 86) = 0.21$, $p > .64$.

We also conducted a between-subjects comparison of extremity of policy support by comparing initial position ratings made by the preexplanation-rating group with the postexplanation ratings made by the no-preexplanation-rating group. We conducted an ANOVA with issue number as a within-subjects factor and judgment timing (before explanation vs. after explanation) as a between-subjects factor. As predicted, there was a significant effect of judgment timing: Judgments made after explanations were less extreme than were judgments made before explanations (preexplanation-rating condition: $M = 1.41$,

$SE = 0.07$; postexplanation-rating condition: $M = 1.19$, $SE = 0.08$), $F(1, 196) = 3.97$, $p < .05$, $\eta_p^2 = .020$, a result that replicated the moderation effect observed for participants who did not give preexplanation ratings of their positions.

Relation between understanding and position extremity. Finally, we assessed correlations between postexplanation position extremity and change in reported understanding to provide evidence that reducing the illusion of depth led participants to express more moderate views. Indeed, an analysis of participant-item pairs revealed a significant negative correlation between the average magnitude of the change in reported understanding and the extremity of the position after explanation, $r = -.19$, $p < .01$. We also examined participants' judgments of how certain they were of their positions after explanation. Uncertainty (i.e., reverse-coded certainty) was negatively correlated with position extremity, $r = -.75$, $p < .001$, and positively correlated with the magnitude of change in understanding, $r = .31$, $p < .001$. Our interpretation of this pattern is that attempting to explain policies made people feel uncertain about them, which in turn made them express more moderate views. This interpretation was supported by mediation analysis (Preacher & Hayes, 2008), which revealed that the effect of change in understanding on extremity was mediated by a significant indirect effect of uncertainty, with a 95% confidence interval excluding 0 [.113, .309].

Discussion

As predicted, asking people to explain how policies work decreased their reported understanding of those policies and led them to report more moderate attitudes toward those policies. We observed these effects both within and between participants. Change in understanding correlated with position extremity, such that participants who exhibited greater decreases in reported understanding also tended to exhibit greater moderation of their positions. Results from a mediation analysis suggested that this relationship was mediated by position uncertainty. Taken together, these results suggest that initial extreme positions were supported by unjustified confidence in understanding and that asking participants to explain how policies worked decreased their sense of understanding, leading them to endorse more moderate positions.

Experiment 2: Generating Mechanistic Explanations Versus Enumerating Reasons

The goal of Experiment 2 was to examine whether the attitude-moderation effect observed in Experiment 1 was

driven specifically by an attempt to explain mechanisms or merely by deeper engagement and consideration of the policies. To induce some participants to deliberate without explaining mechanisms, we asked one group to enumerate reasons why they held the policy attitude they did. Listing reasons why one supports or opposes a policy does not necessarily entail explaining how that policy works; for instance, a reason can appeal to a rule, a value, or a feeling. Prior research has suggested that when people think about why they hold a position, their attitudes tend to become more extreme (for a review, see Tesser, Martin, & Mendolia, 1995), in contrast to the results observed in Experiment 1. Thus, we predicted that asking people to list reasons for their attitudes would lead to less attitude moderation than would asking them to articulate mechanisms.

Method

For participants in the *mechanism* conditions, methods were almost identical to those used for the preexplanation-rating conditions of Experiment 1. For participants in the *reasons* conditions, we modified instructions for the explanation task so that participants were asked to enumerate reasons for their position rather than generate a mechanistic explanation of it (see Example Instructions for Explanation- and Reason-Generation Tasks in the Supplemental Material). We made two additional changes from Experiment 1, omitting the measure of certainty and adding an attention filter to the end of the questionnaire. We also dropped conditions that involved the Iran and merit-pay issues because of a programming error.

One hundred forty-one individuals were recruited using Amazon's Mechanical Turk and participated in return for a small payment. Participants were assigned to the remaining four conditions (two reasons and two mechanism conditions covering either the Social Security and health care issues or the flat-tax and cap-and-trade issues); 112 of these passed the attention filter (mechanism conditions: $n = 47$; reasons conditions: $n = 65$) and were included in the analyses. These participants were 50% male and 50% female, and their average age was 33.9 years. Participants' reported political affiliations were 43% Democrat, 19% Republican, 36% independent, and 4% other.

Results

Replication of Experiment 1. To examine whether results from the mechanism conditions replicated our results from Experiment 1, we submitted judgments of understanding to the same repeated measures ANOVA, which yielded similar results. There was a significant main effect of judgment timing on reported understanding,

such that reported understanding decreased following mechanistic explanation, $F(1, 46) = 20.39, p < .001, \eta_p^2 = .31$. As in Experiment 1, we found the same pattern across all policies. The unexpected main effect of issue number was also significant and, again, there was no significant interaction. Also replicating Experiment 1, results revealed that participants endorsed more moderate positions following mechanistic explanations, $F(1, 46) = 7.32, p < .01, \eta_p^2 = .14$. Finally, change in understanding and extremity change were again significantly correlated, $r = -.34, p < .05$, which suggests that larger reductions in rated understanding following explanation led to less extreme positions.

Mechanistic explanations versus reasons. We next compared the magnitude of change in reported understanding and position extremity across the mechanism and reasons conditions (see Figs. 1a and 1b). We observed a small effect on judgments of understanding in the reasons conditions: Reported understanding slightly decreased after participants enumerated reasons, $F(1, 64) = 7.51, p < .01, \eta_p^2 = .11$. Analysis of the individual reasons given by participants showed that this trend was driven by participants who could provide no reason for their position (see Analysis of Reasons Given in Experiment 2 in the Supplemental Material for further details). More important, and as predicted, the decrement in understanding after enumerating reasons was smaller than the decrement following mechanistic explanation, as reflected by a significant interaction between judgment timing and condition, $F(1, 110) = 6.64, p < .01, \eta_p^2 = .057$. With regard to extremity of positions, there was no change after enumerating reasons, $F(1, 64) < 1$, n.s. Moreover, as predicted, the change in position in the reasons conditions was smaller than in the mechanism conditions, as reflected by a significant interaction between judgment timing and

condition on extremity scores, $F(1, 110) = 3.90, p < .05, \eta_p^2 = .034$.

Discussion

Experiment 2 replicated the results of Experiment 1 and showed further that reductions in rated understanding of policies were less pronounced among participants who enumerated reasons for their positions than among participants who generated causal explanations for them. Moreover, enumerating reasons did not lead to any change in position extremity. Contrary to findings from some previous studies, the results showed that reason generation did not increase overall attitude extremity, although an analysis of individual reasons suggested that it did increase overall attitude extremity when participants provided a reason that was an evaluation of the policy. Other types of reasons led to no change (see Analysis of Reasons Given in Experiment 2 in the Supplemental Material).

Experiment 3: Decision Making

In Experiment 3, we examined whether the moderating effect of mechanistic explanations on political attitudes demonstrated in Experiments 1 and 2 would extend to political decisions. As in Experiment 2, participants first rated their position on a given policy and then provided either a mechanistic explanation of it or reasons why they supported or opposed it. Next, they chose whether or not to donate a bonus payment to a relevant advocacy group. We predicted that participants' initial level of support for the policy would be more weakly associated with their subsequent likelihood of donating in the mechanism condition than in the reasons condition because articulating mechanisms attenuates attitude

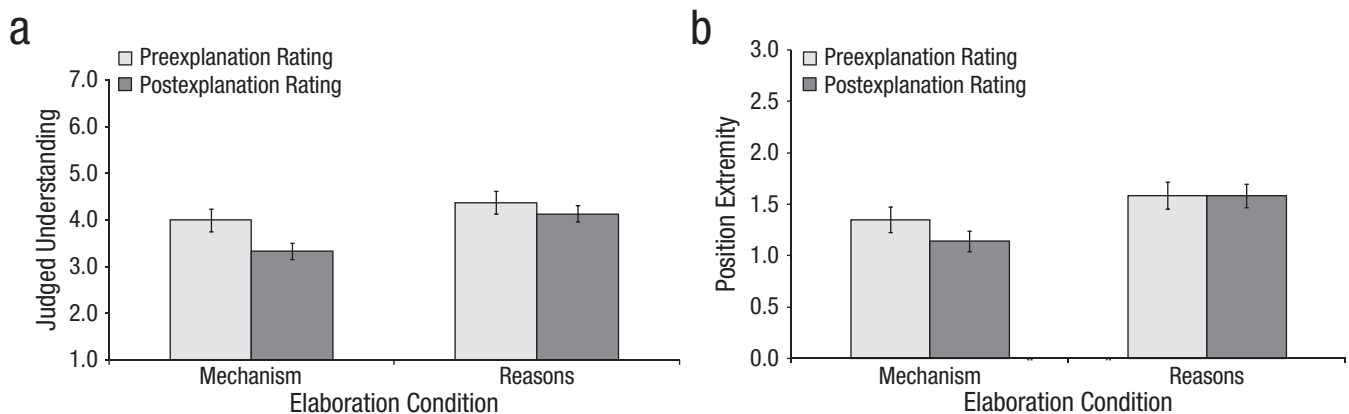


Fig. 1. Results from Experiment 2: (a) judged understanding of policies and (b) extremity of positions on policies as a function of condition and timing of judgment. Understanding was rated on scales from 1 to 7, with higher scores indicating greater understanding. Extremity scores could range from 0 to 3, with higher scores reflecting stronger attitudes in favor of or against a given policy. Error bars represent ± 1 SE.

extremity more than does listing reasons. Thus, we predicted an interaction between the extremity of initial policy support and condition (reasons vs. mechanism) on likelihood of donation.

Method

We recruited 101 U.S. residents (59.0% male, 41% female; average age = 37.3 years) using the same methods used for participant recruitment in Experiment 1. Nine participants did not pass the attention filter and were excluded from subsequent analysis. Participants first provided their position on the six policies, as in the two previous experiments. They were then assigned to one of four conditions and asked to elaborate on one of two policies: cap and trade or flat tax. Depending on condition, participants were asked either to generate a mechanistic explanation ($n = 45$) or to enumerate reasons for their position ($n = 47$), following the same instructions used in Experiment 2. Next, participants were told that they would receive a bonus payment (20 cents; equal to 20% of their compensation for completing the experiment) and that they had four options for what they could do with this bonus payment. They could (a) donate it to a group that advocated in favor of the issue in question, (b) donate it to a group that advocated against the issue, (c) keep the money for themselves (after answering a few additional questions), or (d) turn it down.

Results and discussion

Figure 2 illustrates the likelihood of donating as a function of initial level of policy support for the mechanism and reasons conditions (no participants chose to donate to a group that advocated against their stated position). Our key prediction was that there would be an interaction between initial extremity of policy support and condition,

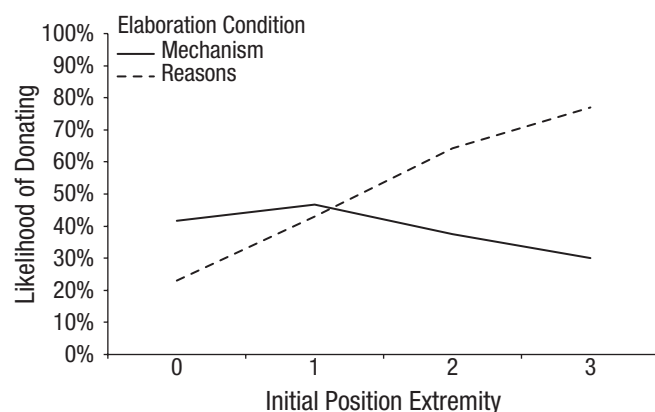


Fig. 2. Results from Experiment 3: likelihood of donating to an advocacy group as a function of condition and initial extremity of position toward a policy.

such that greater extremity would lead to a greater likelihood of donation among participants in the reasons condition but that this tendency would be attenuated in the mechanism condition. We tested this prediction using logistic regression. The dependent variable was whether the participant chose to donate. The independent variables were initial extremity of policy support, condition (reasons vs. mechanism), and their interaction. As predicted, there was a significant interaction between initial extremity of policy support and condition, Waldman's $\chi^2(1) = 6.05$, $p = .014$.

To interpret this interaction, we used spotlight tests (Irwin & McClelland, 2002) at the high and low levels of initial extremity. At the lowest level of initial support, there was no difference in likelihood of donating between the mechanism and reasons conditions, Waldman's $\chi^2(1) = 1.78$, $p > .18$, but at the highest level of initial support, participants in the reasons condition were more likely to donate than were those in the mechanism condition, Waldman's $\chi^2(1) = 6.74$, $p < .01$.

The results of Experiment 3 suggest that among participants who initially held a strong position, attempting to generate a mechanistic explanation attenuated their positions, thereby making them less likely to donate. Consistent with our findings showing a lack of attitude moderation in the reasons condition of Experiment 2, results revealed that initial position extremity was correlated with likelihood of donation in the reasons condition of Experiment 3, which suggests that enumerating reasons did not have the same moderating effect as mechanistic explanation.

General Discussion

Across three studies, we found that people have unjustified confidence in their understanding of policies. Attempting to generate a mechanistic explanation undermines this illusion of understanding and leads people to endorse more moderate positions. Mechanistic-explanation generation also influences political behavior, making people less likely to donate to relevant advocacy groups. These moderation effects on judgment and decision making do not occur when people are asked to enumerate reasons for their position. We propose that generating mechanistic explanations leads people to endorse more moderate positions by forcing them to confront their ignorance. In contrast, reasons can draw on values, hearsay, and general principles that do not require much knowledge.

Previous research has shown that intensively educating citizens can improve the quality of democratic decisions following collective deliberation and negotiation (Fishkin, 1991). One reason for the effectiveness of this strategy may be that educating citizens on how policies

work moderates their attitudes, increasing their willingness to explore opposing views and to compromise. More generally, the present results suggest that political debate might be more productive if partisans first engaged in a substantive and mechanistic discussion of policies before engaging in the more customary discussion of preferences and positions. However, fostering productive discourse among people who have different political stances faces obstacles and can have consequences that fall outside the scope of the current research. Future research should explore the benefits of mechanistic explanation in more ecologically valid civil-discourse contexts.

Our results suggest a corrective for several psychological phenomena that make polarization self-reinforcing. People often are unaware of their own ignorance (Kruger & Dunning, 1999), seek out information that supports their current preferences (Nickerson, 1998), process new information in biased ways that strengthen their current preferences (Lord, Ross, & Lepper, 1979), affiliate with other people who have similar preferences (Lazarsfeld & Merton, 1954), and assume that other people's views are as extreme as their own (Van Boven, Judd, & Sherman, 2012). In sum, several psychological factors increase extremism, and attitude polarization is therefore hard to avoid. Explanation generation will by no means eliminate extremism, but our data suggest that it offers a means of counteracting a tendency supported by multiple psychological factors. In that sense, it promises to be an effective debiasing procedure.

Acknowledgments

The authors thank Julia Kamin, Julia Shube, and Jacob Cohen for help with data collection and John Lynch, Jake Westfall, Donnie Lichtenstein, Pete McGraw, Bart De Langhe, Meg Campbell, and Ji Hoon Jhang for helpful conversations.

Declaration of Conflicting Interests

The authors declared that they had no conflicts of interest with respect to their authorship or the publication of this article.

Supplemental Material

Additional supporting information may be found at <http://pss.sagepub.com/content/by/supplemental-data>

References

- Alter, A. L., Oppenheimer, D. M., & Zemla, J. C. (2010). Missing the trees for the forest: A construal level account of the illusion of explanatory depth. *Journal of Personality and Social Psychology*, 99, 436–451.
- Delli Carpini, M. X., & Keeter, S. (1996). *What Americans know about politics and why it matters*. New Haven, CT: Yale University Press.
- Fernbach, P. M., Sloman, S. A., St. Louis, R., & Shube, J. N. (2013). Explanation fiends and foes: How mechanistic detail determines understanding and preference. *Journal of Consumer Research*, 39, 1115–1131.
- Fishkin, J. S. (1991). *Democracy and deliberation: New directions for democratic reform* (Vol. 217). New Haven, CT: Yale University Press.
- Hirt, E. R., & Markman, K. D. (1995). Multiple explanation: A consider-an-alternative strategy for debiasing judgments. *Journal of Personality and Social Psychology*, 69, 1069–1086.
- Irwin, J. R., & McClelland, G. H. (2001). Misleading heuristics and moderated multiple regression models. *Journal of Marketing Research*, 38, 100–109.
- Keil, F. C. (2003). Folkscience: Coarse interpretations of a complex reality. *Trends in Cognitive Sciences*, 7, 368–373.
- Krosnick, J. A., & Petty, R. E. (1995). Attitude strength: An overview. In R. E. Petty & J. A. Krosnick (Eds.), *Attitude strength: Antecedents and consequences* (pp. 1–24). Mahwah, NJ: Erlbaum.
- Kruger, J., & Dunning, D. (1999). Unskilled and unaware of it: How difficulties in recognizing one's own incompetence lead to inflated self-assessments. *Journal of Personality and Social Psychology*, 77, 1121–1134.
- Lazarsfeld, P. F., & Merton, R. K. (1954). Friendship as social process: A substantive and methodological analysis. In M. Berger, T. Abel, & C. H. Page (Eds.), *Freedom and control in modern society* (pp. 18–66). New York, NY: Octagon Books.
- Linville, P. W. (1982). The complexity-extremity effect and age-based stereotyping. *Journal of Personality and Social Psychology*, 42, 193–211.
- Linville, P. W. (1985). Self-complexity and affective extremity: Don't put all of your eggs in one cognitive basket. *Social Cognition*, 3, 94–120.
- Linville, P. W. (1987). Self-complexity as cognitive buffer against stress-related illness and depression. *Journal of Personality and Social Psychology*, 52, 663–676.
- Lord, C. G., Ross, L., & Lepper, M. R. (1979). Biased assimilation and attitude polarization: The effects of prior theories on subsequently considered evidence. *Journal of Personality and Social Psychology*, 37, 2098–2109.
- Nickerson, R. S. (1998). Confirmation bias: A ubiquitous phenomenon in many guises. *Review of General Psychology*, 2, 175–220.
- Preacher, K. J., & Hayes, A. F. (2008). Asymptotic and resampling strategies for assessing and comparing indirect effects in multiple mediator models. *Behavior Research Methods*, 40, 879–891.
- Ross, L., Lepper, M. R., Strack, F., & Steinmetz, J. (1977). Social explanation and social expectation: Effects of real and hypothetical explanations on subjective likelihood. *Journal of Personality and Social Psychology*, 35, 817–829.
- Rozenblit, L., & Keil, F. C. (2002). The misunderstood limits of folk science: An illusion of explanatory depth. *Cognitive Science*, 26, 521–562.
- Russell, B. (1928/1996). *Sceptical essays*. New York, NY: Routledge Classics.

- Schickel, R. (2005, Feb. 20). Clint Eastwood on "Baby." *Time Magazine*. Retrieved from <http://www.time.com/time/magazine/article/0,9171,1029865,00.html>
- Schkade, D., Sunstein, C. R., & Hastie, R. (2010). When deliberation produces extremism. *Critical Review: A Journal of Politics and Society*, 22, 227–252.
- Tesser, A. (1978). Self-generated attitude change. In L. Berkowitz (Ed.), *Advances in experimental social psychology* (Vol. 11, pp. 289–338). New York, NY: Academic Press.
- Tesser, A., Martin, L., & Mendolia, M. (1995). The impact of thought on attitude extremity and attitude-behavior consistency. In R. E. Petty & J. A. Krosnick (Eds.), *Attitude strength: Antecedents and consequences* (pp. 73–92). Mahwah, NJ: Erlbaum.
- Van Boven, L., Judd, C. M., & Sherman, D. K. (2012). Political polarization projection: Social projection of partisan attitude extremity and attitudinal processes. *Journal of Personality and Social Psychology*, 103, 84–100.