

Conversational receptiveness: Improving engagement with opposing views

Michael Yeomans^{a,*}, Julia Minson^b, Hanne Collins^a, Frances Chen^c, Francesca Gino^a

^a Harvard Business School, United States

^b Harvard Kennedy School, United States

^c University of British Columbia, Canada

ARTICLE INFO

Keywords:

Conflict management

Natural language processing

Communication

Receptiveness

Disagreement

ABSTRACT

We examine “conversational receptiveness” – the use of language to communicate one’s willingness to thoughtfully engage with opposing views. We develop an interpretable machine-learning algorithm to identify the linguistic profile of receptiveness (Studies 1A–B). We then show that in contentious policy discussions, government executives who were rated as more receptive – according to our algorithm and their partners, but not their own self-evaluations – were considered better teammates, advisors, and workplace representatives (Study 2). Furthermore, using field data from a setting where conflict management is endemic to productivity, we show that conversational receptiveness at the beginning of a conversation forestalls conflict escalation at the end. Specifically, Wikipedia editors who write more receptive posts are less prone to receiving personal attacks from disagreeing editors (Study 3). We develop a “receptiveness recipe” intervention based on our algorithm. We find that writers who follow the recipe are seen as more desirable partners for future collaboration and their messages are seen as more persuasive (Study 4). Overall, we find that conversational receptiveness is reliably measurable, has meaningful relational consequences, and can be substantially improved using our intervention (183 words).

1. Introduction

Disagreement is a fundamental feature of social life, permeating organizations, families and friendships. In the course of any complex relationship or organizational endeavor, however, stakeholders need to coordinate their behavior, even when holding diametrically opposing beliefs. At least a minimal ability to co-exist with disagreeing others is a requirement of social cooperation and democratic society (Iyengar & Westwood, 2015; Milton, 1644/1890; Westfall, Van Boven, Chambers, & Judd, 2015). And although the need to confront the opposing views of others is particularly notable in civic life, it also permeates professional organizations (Baron, 1991; Fiol, 1994; Schweiger, Sandberg, & Rechner, 1989; Todorova, Bear, & Weingart, 2014) and domestic relationships (Cutrona, 1996; Gottman & Levenson, 1999; Gottman, 1994).

While encountering opposing viewpoints seems inevitable, in practice, people do not seem to handle disagreement well. An extensive body of research has shown that the presence of contradictory opinions gives rise to avoidance (Gerber, Huber, Doherty, & Dowling, 2012; Chen & Rohla, 2018), negative affect (Gottman, Levenson, & Woodin, 2001; Wojcieszak, 2012), biased information processing (Frey, 1986; Hart et al., 2009; Nickerson, 1998), reactance (Brehm 1966; Goldsmith, 2000; Fitzsimons and Lehmann, 2004; John, Jeong, Gino, & Huang,

2019; Blunden, Logg, Brooks, John, & Gino, 2019), and negative inferences about the other side (Minson, Liberman, & Ross, 2011; Pacilli, Roccato, Pagliaro, & Russo, 2016; Ross & Ward, 1995; 1996).

Prior research has demonstrated that disagreement can lead to conflict spirals and harm the relationship (Ferrin, Bligh, & Kohles, 2008; Kennedy & Pronin, 2008; McCroskey & Wheelless, 1976; Weingart, Behfar, Bendersky, Todorova, & Jehn, 2015). Furthermore, conflict in one relationship can spill over into other relationships (Neff & Karney, 2004; Repetti, 1989; Story & Repetti, 2006; King & DeLongis, 2014; Keltner, Ellsworth, & Edwards, 1993; Goldberg, Lerner, & Tetlock, 1999).

Yet, despite these risks, there are many potential benefits of engagement with disagreeing others. This is especially true when individuals have interdependent goals that are more important than resolving the disagreement itself, which is common in many professional, organizational, and personal relationships. For example, across many domains, other viewpoints can help us increase the accuracy of our own beliefs by exposing us to new information and perspectives (Liberman et al., 2012; Soll & Larrick, 2009; Sunstein & Hastie, 2015; Tost, Gino, & Larrick, 2013). Organizations depend on an active and healthy conversation among diverse perspectives for making decisions and giving voice to underrepresented views (De Dreu & Van Vianen, 2001; Edmondson & Lei, 2014; Hirschman, 1970; Shi, Teplitskiy, Duede, &

* Corresponding author.

E-mail address: myeomans@hbs.edu (M. Yeomans).

Evans, 2019; Van Knippenberg & Schippers, 2007). Importantly, disagreements that are persistently discussed are those where neither person has yet been (or is likely to be) persuaded (Sears & Funk, 1999; Hillygus, 2010). In these cases, civil engagement with disagreeing others might foster a mutually beneficial space for understanding and improve parties' ability to achieve other cooperative goals. These benefits of disagreement, however, are foregone when constructive engagement spirals into conflict. In the present research, we examine the choices that individuals make in the course of conversation that enable productive dialogue around opposing views.

There have been many interventions designed to foster constructive engagement and resolve conflict. The majority of this research has focused on individuals' attitudes toward conflict counterparts and their willingness to interact with them, but has not closely examined the content of those interactions (Weingart et al., 2015). For example, several streams of research have attempted to change people's beliefs regarding members of outgroups and their claims (e.g. Bruneau & Saxe, 2012; Hameiri, Porat, Bar-Tal, & Halperin, 2016; Schroeder, Kardas, & Epley, 2017; Schroeder & Risen, 2016; Turner & Crisp, 2010), or else have encouraged people to seek out opposing viewpoints in their media consumption (e.g. Bail et al., 2018; Dorison, Minson, & Rogers, 2019). There is also a long literature on intergroup contact as a mode for engagement across difference (Allport, 1979; Pettigrew & Tropp, 2006), although there is wide variation in how effective intergroup contact can be (MacInnis & Page-Gould, 2015; Mannix & Neale, 2005; Paluck, Green, & Green, 2018). The inconsistency in these results suggests that perhaps the interpersonal outcomes of engagement may be critically moderated by the way in which counterparts treat one another during their interactions. Previous research, however, has mostly left the content of these conversations unexplored.

In the present research, we examine whether people can improve the ways in which they communicate with holders of opposing views. Specifically, we test whether it is possible to communicate "receptiveness to opposing views" (Minson, Chen, & Tinsley, 2019) in the course of a conversation between people who disagree and the interpersonal consequences of communicating receptiveness. Specifically, we test whether communicating in a more receptive manner helps foster cooperative goals between disagreeing others, such as willingness to work together in the future, interpersonal trust, and conflict de-escalation. We address this question by identifying the specific signals that people recognize when they judge a partner's receptiveness. We use these cues to create an intervention to encourage conversation partners to communicate their willingness to engage with opposing views.

1.1. Communication in conflict

Interpersonal conflicts of all kinds - in organizations, in politics, in families - often includes parties who assert that the other side has not listened to them carefully, considered their arguments thoughtfully, and evaluated their position in a fair-minded way. The belief that the other side is failing to be sufficiently open to or respectful of one's views features prominently in the clinical psychology literature on conflict (Gottman, 1993, 1994; Gottman, 2008). Additionally, the feeling that one's partner is not "listening" or working with enough diligence to understand one's perspective has been repeatedly shown to stand in the way of conflict resolution, relationship satisfaction and positive intergroup relations (Cohen, Schulz, Weiss, & Waldinger, 2012; Gordon & Chen, 2016; Livingstone, Fernández, & Rothers, 2019).

The benefits of receptive listening have also been well documented. Employees who feel that their boss listens to them experience less emotional exhaustion and report being more willing to stay in their positions (Lloyd, Boer, Keller, & Voelpel, 2015); and open-mindedness among supervisors and employees leads to more efficient solutions to workplace conflict (Tjosvold, Morishima, & Belsheim, 1999). Even in the medical context, patients who feel that their doctors listen to them show higher medication adherence (Shafan-Tikva & Kluger, 2016).

Across domains, the belief that one's counterpart is open to hearing what one has to say has important implications.

Interventions to improve perceptions of listening have yielded positive results. For example, Chen, Minson, and Tormala (2010) instructed experimental confederates to ask "elaboration questions" in the course of disagreement. Participants who received a question requesting that they further elaborate on the bases for their views were more willing to engage in future conversations with their counterpart and evaluated their counterpart more positively. Relatedly, conflict resolution practitioners advise engaging partisans in an exercise of restating each other's positions in order to demonstrate that counterparts have indeed "heard" each other (e.g., Coltri, 2010).

Some interventions have focused more specifically on the language used during conflictual dialogue. For example, "I-statements" (e.g., "I feel overwhelmed when you don't help with chores"), rather than "you-statements" (e.g., "You are so lazy; you never help with chores"), have been recommended as a means of attenuating conflict by reducing accusation or contempt (Gottman, 1994, 2008; Gottman, Notarius, Gonso, & Markman, 1976; Simmons, Gordon, & Chambless, 2005). Indeed, when presented with hypothetical conflict scenarios, lay individuals report believing that such strategies will reduce hostility during conflict (Rogers, Howieson, & Neame, 2018). But this research does not observe how people communicate with disagreeing counterparts in open, unstructured conversation, whether such cues are perceived or reciprocated by counterparts, and whether they lead to positive downstream consequences.

1.2. Dispositional receptiveness to opposing views

An inherent challenge in investigating receptiveness in an interpersonal context is establishing the "ground truth" of what receptiveness is and who is being receptive. A large prior literature demonstrates that at baseline, parties in conflict do in fact process opposing views in a biased manner and sometimes fail to process them at all (Frey, 1986; Lord, Ross, & Lepper, 1979; Nickerson, 1998; Ross & Ward, 1995, 1996; Hart et al., 2009). In other words, balanced and thoughtful consideration of opposing perspectives is the exception rather than the rule. Recent research by Minson, Chen, and Tinsley (2019), however, has identified an individual difference, "receptiveness to opposing views," that reliably predicts individuals' dispositional willingness to engage with opposing views and linked this tendency to various measures of information consumption.

Minson et al. present an 18-item scale to measure dispositional receptiveness and show that it predicts behavior during three stages of information processing. First, individuals high in receptiveness are more willing to seek out belief-disconfirming information (i.e., demonstrate less "selective exposure"). Second, individuals high in receptiveness exposed to belief-disconfirming information pay more attention to it. Finally, receptive individuals offer assessments of that information that appear less biased by their prior beliefs (i.e., demonstrate less "biased assimilation," Lord et al., 1979).

Minson et al. established the predictive validity of the receptiveness scale over a variety of behaviors related to information processing, even months in advance, supporting a dispositional model of receptiveness. However, these behaviors all reflected an individual's passive exposure - reading or watching opposing views far removed from the person espousing them. Whether and how receptiveness affects the ways in which people interact with conflict counterparts in unstructured conversation remains an open question.

1.3. Conversational receptiveness

In the present research, we focus on *conversational receptiveness* - the extent to which parties in disagreement communicate their willingness to engage with each other's views. We approach this construct from four measurement perspectives. First, we ask individuals to introspect

on and report their dispositional willingness to engage with opposing views using the existing Minson et al. scale. Second, we ask individuals to write a response to a position they disagree with and evaluate the extent to which they believe they communicated such willingness in the text. Third, we ask other people who disagree with that person (including the original position writer) to assess the text on the same dimension. Finally, we develop a natural language-processing algorithm for conversational receptiveness – i.e., a model that identifies the specific features of natural language that lead conflict counterparts to perceive each other as receptive. Comparing where the four measures of conversational receptiveness align versus diverge gives us insight into this complex but crucial facet of conflictual communication.

It is possible that a speaker's willingness to thoughtfully engage with opposing views is easily and transparently communicated. For example, research on sentiment analysis using the tools of natural language processing has demonstrated that people can clearly communicate their evaluations of goods and experiences (Liu, 2012). By contrast, it could be the case that communicating receptiveness is quite difficult even when one has every intention of doing so. Self-help books, the advice of clinicians, and our own rumination provide ample guidance regarding how to come across as an open-minded, attentive, and respectful partner in conflict. However, to our knowledge, no experimental research has examined whether our collective intuitions are in fact accurate.

Capturing the specific linguistic markers of conversational receptiveness allows us to begin designing simple, scalable interventions to improve conversation. Although an extensive prior literature has tested a variety of conflict-resolution strategies, many are not easily scalable or applicable “in the heat of the moment.” For example, the research on conflict-resolution interventions often features in-depth, time-intensive trainings (e.g. Gottman, 2008; Heydenberk, Heydenberk, & Tzenova, 2006; Sessa, 1996), a tightly controlled context in which participants receive direct instruction from an experimenter (e.g., Gutenbrunner & Wagner, 2016; Halperin, Porat, Tamir, & Gross, 2013; Johnson, 1971; Liberman, Anderson, & Ross, 2010; Page-Gould, Mendoza-Denton, & Tropp, 2008), or complex laboratory inductions requiring writing or listening exercises (e.g., Bruneau & Saxe, 2012; Finkel, Slotter, Luchies, Walton, & Gross, 2013).

Furthermore, while many studies measure how beliefs are developed and changed (about issues, or about other people), very few directly measure the interpersonal behavior itself (i.e., the language exchanged by conflicting parties). Indeed, the majority of the research on barriers to conflict resolution do not allow opposing partisans to interact at all, but instead attempt to change the perceptions of the members of one group regarding the members, or claims of the members, of a rival outgroup (Bruneau, Cikara, & Saxe, 2015; Bruneau & Saxe, 2012; Galinsky & Moskowitz, 2000; Wang, Kenneth, Ku, & Galinsky, 2014).

Our work is based on the hypothesis that successful discussion of opposing views is in part hampered by individuals' inability to clearly communicate their willingness to thoughtfully engage with their opponents' views (i.e., exhibiting low conversational receptiveness). In our proposed model, conversational receptiveness is a behavioral construct that mediates the effect of conversational decisions by one person on their conversation partner. In other words, when one individual decides to signal willingness to have a thoughtful and respectful conversation, the words they use to communicate that willingness ultimately determine whether their partner does or does not accurately interpret their intentions. This interpretation in turn affects the partner's reciprocal behavior toward the speaker, both in the current conversation and in the future. In this way, receptiveness expressed in language offers a new mediational pathway for all pre-conversation interventions that shape a speaker's attitudes and intentions toward their counterpart.

To test this theory, we develop a model of conversational receptiveness consisting of a set of simple conversational strategies, and

measured directly from natural language. We show that this model predicts two important outcomes in discussions regarding a point of disagreement: conflict escalation in the current conversation and collaboration intentions regarding the future. Importantly, we find that speakers misjudge their conversational receptiveness, creating an easy opportunity for intervention.

By documenting the challenges in communicating receptiveness and providing empirically tested and easily implementable strategies for overcoming it, we hope to both illuminate an important barrier to conflict resolution and provide an easily scalable intervention. Our findings regarding markers of receptiveness in conversation and individuals' misperceptions of those markers can be readily applied to creating more productive conversations in organizations, families, and civic settings.

1.4. Overview of the present research

We report the results of five studies wherein participants were exposed to statements on controversial issues written by people with whom they disagree. Some participants were then asked to write a response that communicates a willingness to thoughtfully consider the original writer's position (i.e., to be receptive). In our first study, we instructed some participants to respond in a receptive manner and compared their writing to that of participants who were instructed to respond naturally. A separate group of participants read and evaluated the written responses, a process that allowed us to test the effectiveness of simply instructing individuals to be more receptive and to train a natural language-processing algorithm to model conversational receptiveness.

In Study 2, we tested this construct in conversations between senior government executives who were paired to discuss controversial policy topics with a disagreeing partner. After the conversation, participants rated their own and their partner's receptiveness to their arguments. In Study 3, we examined discussions in an organizational context where disagreement naturally arises and where people have more freedom to choose discussion topics and respond to others: the editorial process of correcting Wikipedia articles.

In Study 4, we test a short intervention designed to teach conversational receptiveness. Specifically, we teach writers a “receptiveness recipe” and examine how counterparts rate them.

For each study in which we collected our own data, we report how we determined our sample size, all data exclusions, all manipulations, and all measures. All data (both our own and from other sources), analysis code, stimuli, and preregistrations from each study are available in the Online Supplemental Material, stored on OSF at <https://bit.ly/2QwyiuL>.

2. Study 1: Developing an algorithm to measure conversational receptiveness

In the first set of studies, we use asynchronous conversations to establish how conversational receptiveness can be measured and manipulated. In Study 1A, we randomly assigned participants (“responders”) to write a response to a political statement on one of two issues (campus sexual assault or police treatment of minorities) that expressed a point of view opposed to their own. Some of these responders were randomly assigned to first receive instructions on how to be receptive, while others received instructions to respond “naturally.” This approach gave us a manipulated ground truth measure of intended receptiveness. In Study 1B, we then recruited a new set of participants (“raters”) who evaluated the responses blind to condition, giving us a ground truth measure of rated receptiveness.

Using the text of the responses and these ground truth measures, we developed a machine learning algorithm to detect conversational receptiveness in natural language. This allowed us to estimate the malleability of this construct and generate a prescriptive model for how

people can better express it. Because groups of responders wrote about two different issues, we can evaluate whether conversational receptiveness is generalizable across domains.

2.1. Study 1A methods

Sample. All participants were recruited from Amazon's Mechanical Turk (mTurk) to participate in a study on "Political Issues and Discussions" ($M_{\text{age}} = 38.0$; 45% Male). In line with our pre-registered exclusion criteria, we only excluded participants who did not complete our attention checks or reported "no opinion" on their assigned issue, or who did not complete the study. This left a final sample of 1,102 participants.

Protocol. In the first phase of this study, participants were randomly assigned to one of two conditions (see Appendix A for full details). In the receptive condition, participants read a description of the construct of receptiveness to opposing views as "being willing to read, deeply think about, and fairly evaluate the views of others, even if you disagree." Then participants were shown a sample of questions from the 18-item, Receptiveness to Opposing Views scale (Minson, Chen, & Tinsley, 2019) and were told how someone who was high vs. low in receptiveness would respond to these items. Finally, participants were presented with four new "Agree" or "Disagree" items from the receptiveness scale and were asked to answer the questions the way a person who is receptive would respond.

In the control condition, participants were asked to read a passage about the discovery of a new species of fish. After reading this passage, participants were asked to answer four comprehension questions about the passage they just read. In both conditions, participants who offered an incorrect response to the quiz items were prompted to answer the item again. This quiz was designed to be irrelevant to the main writing task but still produce a similar level of effort as the quiz in the receptive condition.

Next, all participants were randomly assigned to consider one of two issue statements. One statement dealt with the issue of policing and minority suspects: "The public reaction to recent confrontations between police and minority crime suspects has been overblown." The other statement dealt with the issue of campus sexual assault: "When a sexual assault accusation is made on a college campus, the alleged perpetrator should be immediately removed from campus to protect the victim's well-being." Participants were asked to state their agreement with the statement to which they were assigned on a scale from "-3: Strongly Disagree" to "+3: Strongly Agree."

We used participants' stated position on the policy issue to which they were assigned to match them to an opinion statement from a previous participant expressing the opposing view. For each participant, we randomly selected a target statement from a pool of 20 (five agreeing and five disagreeing with each of the two issue statements) generated in our previous studies on conflict in political domains. The statements regarding police treatment of minorities were generated in a previous laboratory study in which participants, who were government employees, interacted with a disagreeing peer over a chat platform. The statements regarding campus sexual assault were generated in a study of affective reactions in conflict conducted on mTurk.

Participants in the receptive condition were told, "Imagine that you are having an online conversation with this person. In your response, try to be as receptive and open-minded as you can." Participants in the control condition were told, "Imagine that you are having an online conversation with this person. How would you respond?" Participants were asked to spend at least two minutes writing their response. Finally, they answered demographic questions about their age, gender and political orientation.

2.2. Study 1B methods

Sample. All participants were recruited from Amazon's Mechanical

Turk (mTurk) to participate in a study on "Political Issues and Discussions" ($M_{\text{age}} = 37.8$; 45% Male). In this study, as in all studies with mTurk participants, we precluded participation by individuals who had participated in a similar study (like Study 1A) in the past. In line with our pre-registered exclusion criteria, we only excluded participants who did not complete our attention checks or reported "no opinion" on their assigned issue, or who did not complete the study. We collected data in several waves to achieve our target sample size of at least three raters for every responder from Study 1A. This left a final sample of 1,322 raters, and an average of 4.9 ratings for each response.

Protocol. First, participants reported their positions on our two target issue statements (police relations and sexual assault) on a scale ranging from "-3: Strongly Disagree" to "+3: Strongly Agree." These positions were used to assign the rater to two responders. These responders were randomly assigned, with the condition that the raters and responders always held opposing positions on the issue—thus, raters always agreed with the original writer to whom the responder was responding. Additionally, both responders were always responding to the same original statement, and the rater saw this original statement along with the responders' replies.

Each statement-response pair was presented separately, and participants were asked to read the exchange carefully and rate how receptive the responder was in their response to the writer. We modified the Minson, et al. receptiveness scale to refer to a target individual's behavior in a particular conversation, as opposed to one's own dispositional receptiveness. Thus, an item that originally read "I am willing to have conversations with individuals who hold strong views opposite to my own" became "On this issue, the respondent seems willing to have conversations with individuals who hold strong views opposite to their own" (Appendix B). Finally, the raters answered demographic questions, including about their age, gender, and political orientation.

2.3. Study 1 results

As we hoped, there was no difference in post-treatment attrition among respondents across conditions ($\chi^2(1) = 0.16, p = .69$). This suggests that the task of learning about the receptiveness construct and passing the relevant quiz was roughly as onerous as the control condition task of reading and taking the quiz about the new fish species.

Rated receptiveness. Overall, raters tended to agree on participants' level of receptiveness. No two raters evaluated the same pair of responses, so we instead evaluated rater agreement using a nonparametric measure that mirrored the structure of our experimental design. Individual raters evaluated two responses, and we compared how often each rater ranked their two responses in the same order as the consensus of the other raters. Each pair was composed of two responses to the same writer, removing some variance. However, on average, each rater still agreed with the other raters 66.8% of the time (95% CI = [64.3%, 69.3%]). A traditional test of inter-rater reliability (that does not account for our sample procedure) also suggests that the pooled ratings reveal a stable and universal component of receptiveness, even though raters do tend to have some disagreements individually ($\text{ICC}(3,1) = 0.40$; $\text{ICC}(3,k) = 1$). Given the inherent subjectivity of these judgments, we decided this was an acceptable level of agreement, and decided to pool the raters and focus on the signals of conversational receptiveness that were largely agreed upon by third parties.

The average ratings for each responder also revealed a treatment effect: Participants who were told to be receptive were rated as more receptive ($M = 0.85, SD = 0.93$) than participants in the control condition ($M = 0.39, SD = 0.97$; $t(541) = 5.6, p < .001$; Cohen's $d = 0.47$). This confirms that conversational receptiveness can, in principle, be intentionally manipulated. Although manipulated and rated conversational receptiveness were correlated, our data suggests they are distinct. The median text response in the receptive condition

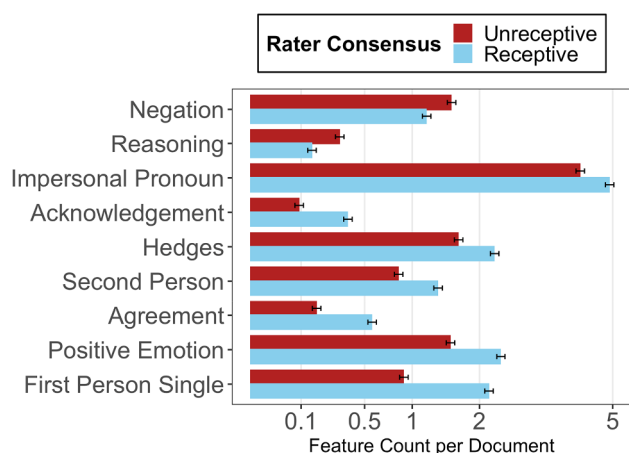


Fig. 1. Model of receptiveness trained with ratings from Study 1B collected for written responses from Study 1A. Bars compare the top-third and bottom-third of responses, in terms of how receptive they were rated, on average. Each datapoint represents a group mean (+/- 1 SE).

was rated as more receptive than only two-thirds (67%) of the responses in the control condition. For the focus of this study, we treat the average rated receptiveness as the primary measure of our construct, and collapse across randomized conditions.

Algorithmic receptiveness. To model the linguistic features of receptiveness, we used the politeness R package (Yeomans, Kantor & Tingley, 2019). This package builds off of state-of-the-art pre-trained natural language processing (NLP) models to calculate a set of syntactic and social markers from natural language (e.g., gratitude, apologies, acknowledgment, commands). We chose this tool because we anticipated that many of the signals of conversational receptiveness were in structural elements of the text that might generalize across domains.

In Fig. 1, we plot the features that most clearly distinguish the responses rated as receptive from those rated as not receptive. Each bar represents the average usage of a feature among the messages people rated as being in the highest and lowest terciles of receptiveness. Some features are quite consistent; for example, positive emotional words are a predictable sign of receptiveness. Additionally, “I” statements are more common from receptive writers than from unreceptive ones. In particular, the “I” statements that seem to make the most difference are explicit acknowledgement (e.g., “I understand,” “I see your point”) and explicit agreement (e.g., “I agree,” “you’re right”)—even though the responders disagreed with the writers by design. Additionally, hedges (e.g., “somewhat,” “might”) were useful to soften factual statements. On the opposite end of the spectrum, specific features were more common in unreceptive responses, included negations (e.g., “no,” “wrong”) and a focus on explanatory reasoning (e.g., “because,” “therefore”).

Domain specificity. Because we collected data on two different issues – police relations and campus sexual assault – we could empirically assess how well our model transfers from one domain to the other. That is, we trained our data on one issue and tested it on the other to see whether the model had been learning generalizable rules about linguistic receptiveness. Using human receptiveness ratings as ground truth, we were able to predict the receptiveness of the police relations responses using a model trained only on sexual assault responses ($r = 0.39$, $t(260) = 6.9$, $p < .001$) and vice versa ($r = 0.506$, $t(279) = 9.8$, $p < .001$).

Linguistic benchmarks. To estimate the overall accuracy of the receptiveness detection algorithm, we performed a 20-fold nested cross-validation across the entire dataset, using a LASSO algorithm to classify responses based on their rated receptiveness (Friedman, Hastie, & Tibshirani, 2010; Stone, 1974). Overall, the results confirmed that the receptiveness detector captured many relevant features of receptiveness

and correlated with the average human ratings ($r = 0.45$, $t(541) = 11.7$, $p < .001$). Using the pairwise validity test above, we found the algorithm could discriminate, from a pair of responses to the same statement, which would be rated as more receptive ($M = 65.2\%$, $95\% \text{ CI} = [62.6\%; 67.8\%]$) at similar rates to the average human rater (66.8%).

We compared some related constructs that could also be calculated from text, and found that our receptiveness model was uniquely predictive of our raters’ average judgments. For example, the most common text analysis benchmarks like word count ($r = 0.07$, $t(541) = 1.7$, $p = .08$) and sentiment ($r = 0.20$, $t(541) = 4.8$, $p < .001$) do not approach the accuracy of our model in these data. We also applied dictionaries to measure use of moral foundations in the responses (Graham, Haidt, & Nosek, 2009; Feinberg and Willer, 2015; Day, Fiske, Downing, & Trail, 2014) and found no relationship between use of words focusing on opponents’ moral values and rated receptiveness ($r = 0.022$, $t(541) = 0.5$, $p = .605$). We also compared our receptiveness model to a recent paper which used the same politeness package to train a model of “communicated warmth” in distributed negotiations (Jeong, Minson, Yeomans, & Gino, 2019). This was somewhat more accurate ($r = 0.167$, $t(541) = 3.9$, $p < .001$), which is unsurprising as it put weight on some similar features (e.g. hedges). However, it was not close in accuracy to the receptiveness detector.

We tested a more open-ended text analysis model – tokenizing each response into one-, two- and three-word sequences from stemmed words, including stop words (Manning, Manning, & Schütze, 1999). We found that according to the same nested cross-validation test, the receptiveness detector had similar performance ($r = 0.46$, $t(541) = 11$, $p < .001$). However, one potential advantage of our conversational receptiveness detector is that it is designed to focus on structural and stylistic features in text rather than content. Accordingly, we found that the “bag-of-ngrams” model could not perform as well as the conversational receptiveness detection model when transferred from one topical domain to another (sexual assault test: $r = 0.36$, $t(260) = 6.2$, $p < .001$; police relations test: $r = 0.43$, $t(279) = 8.0$, $p < .001$).

2.4. Study 1 discussion

These results provide strong initial evidence for the construct of conversational receptiveness. That is, when a person converses with someone who holds an opposing viewpoint, it is possible to intentionally modulate conversational receptiveness, and its signals are reliably detected. Furthermore, conversational receptiveness can be communicated with a discrete set of short linguistic features that are content- and topic-agnostic, and that could be incorporated into almost any conversation.

Our algorithm also yields novel insights regarding the linguistic features perceived as receptive. For example, “I” statements are indeed seen as receptive, primarily because they often precede an affirmation of understanding or at least partial agreement, as in “I understand that...” or “I agree with.” However, “you” statements are actually also seen positively, even though they have long been maligned in the interpersonal literature (Hahlweg et al., 1984; Simmons et al., 2005). In our data, “you” was often a signal of understanding because the word was used when speakers were restating their partner’s position or beliefs. However, this may not hold when people are talking about one another’s actions, rather than their beliefs.

Interestingly, the relationship between instructions to be receptive and perceived receptiveness was far from perfect. One explanation for this is that people in the control condition are already quite receptive. Alternatively, it could be that people have a broken mental model of receptiveness – that is, they may not know how receptively they are perceived, and thus might not know how to change own their communicative behavior to appear more receptive.

3. Study 2: Receptiveness in professional disagreement

In Study 2, we demonstrate the interpersonal benefits of receptiveness during conversation and validate our linguistic model of receptiveness in an observational sample of high-level government professionals. The participants were recruited from a continuing education program for state and local government executives, consisting of an ideologically diverse group of experienced public service professionals ($M_{\text{age}} = 46.6$ years; 70% Male). As part of an in-class exercise, they engaged in an online discussion regarding a hot-button policy topic with a disagreeing partner from the same class. Afterwards, they evaluated both their own and their partner's receptiveness. Importantly, participants also evaluated their partner on several workplace-related dimensions that they knew would be used to form teams for a later class exercise.

3.1. Study 2 methods

Sample. The data here are pooled across eight sessions from the two years in which this exercise was included in the program, in the summers of 2017 and 2018. In each session, everyone in the program was asked to participate, and everyone agreed. Thus, 270 people began the study on the first day of the program. Eighteen participants (nine dyads) were excluded from our analyses because they either did not complete the second day of the study or had missing data due to technical difficulties. Another fourteen participants (seven dyads) were excluded because they did not actually disagree with their partner on the issue to which they were assigned. All our analyses are based on the remaining 238 participants, or 119 dyads.

Protocol. We conducted this two-part study on the first two days of the program in order to ensure that participants did not yet know each other. On the first day of class, participants consented to participate and reported their gender, age, political orientation, and dispositional receptiveness to opposing views (Appendix C).

Participants also reported their views on nine controversial socio-political issues on a scale from -3 : "Strongly Disagree" to $+3$: "Strongly Agree," as well as the extent to which each issue was "personally important" to them from 1 : "Not at all important to me" to 5 : "Very important to me." Only three of these issues were used in the current study: "The death penalty should be abolished in all US States"; "On balance, public sector unions should be reined in"; and "The public reaction to recent confrontations between police and minority crime suspects has been overblown." We chose these three issues because they jointly provided a balanced but polarized distribution of views, as established based on survey data from an earlier sample of the same population.

The participants' responses from the first day were used to create dyads for the second day and to assign those dyads to one of the three issues. We wanted data to be collected evenly across all three issues and to make sure that as many people as possible were paired with a partner who held an opposing view on their assigned issue. This is a non-trivial computational problem: matching algorithms typically optimize for similarity within matched pairs, rather than difference. After the first year, when dyads were created by hand, we wrote our own matching algorithm to automatically generate dyads (described in Appendix D).

On the second day of their class, participants were brought into the laboratory as a group and seated at individual computer terminals so that they could not identify their partner. Participants were told that for approximately 20 min, they would have the opportunity to discuss a controversial topic via an online chat with another member of their program and would answer some questions about the interaction. They were assured that their writing would remain anonymous. To start the conversation, participants first saw their assigned issue statement and received the following instructions:

"Think about whether you agree or disagree with this statement and why.

Please write out the best arguments that you can to support your opinion. Please be as persuasive as possible and try to come up with the best explanation to support your beliefs. Please make sure your statement focuses on the issue and does not contain identifying information."

After participants finished writing out their point of view, the chat software exchanged these opening statements between dyad partners, who were each asked to separately write a response to each other's initial statement. These responses were then exchanged again, so that each person could write a response to their partner's response. Participants then continued these message exchanges, so that their conversations lasted for a total of five rounds (one opening statement plus four responses) from each person.

After the interaction was over, participants completed a survey (see Appendix E for exact items). In a counterbalanced order, they rated their partner's receptiveness toward them during the interaction and their own receptiveness toward their partner. We modified the original 18-item dispositional receptiveness scale to address receptiveness during a specific interaction (i.e., "situational receptiveness"). Participants then again reported their views on the three issues (death penalty, police relations, public sector unions) discussed during the experiment.

We then (truthfully) informed participants that during a future session of the program they would engage in a team decision-making exercise. We then asked them to rate their counterparts on a number of characteristics that would be used as inputs into the formation of those teams. Specifically, participants rated the extent to which (1) they would like to have their discussion partner on a work team, (2) how much they trusted their partner's judgement, and (3) how much they would like their partner to represent their organization in a professional context. These ratings were made using 7-point Likert scales that ranged from "1: Strongly Dislike/Distrust" to "7: Strongly Like/Trust."

Finally, participants reported their beliefs regarding the value of disagreement in policy contexts. Specifically, they reported (1) how useful discussion between disagreeing others is for developing good public policies, (2) how valuable political deliberation is to the democratic process, and (3) to what extent hearing arguments from both sides of an issue ultimately leads to better decisions (1: "Not at all" to 5: "Extremely").

3.2. Study 2 results

Modeling framework. Because the data in Study 2 are dyadic, we must account for within-dyad correlations to make proper assessments of the uncertainty of each estimate. Accordingly, we estimate all effects as standardized regression coefficients and adjust all the standard errors in all of these regressions for clustering within dyads, using the multiwayvcov R package (Graham, Arai, & Hagströmer, 2016). None of the models reported here control for other covariates (except where explicitly noted). However, we confirm that all of these results are robust when we include controls for assigned issue, amount of disagreement, or session/class fixed effects (or all of these controls at once). The details of these analyses are presented in our OSF repository.

Receptiveness ratings. Participants' day one ratings of their own dispositional receptiveness and their partner's dispositional receptiveness were not positively correlated, confirming that random assignment was successful ($\beta = -0.113$, $SE = 0.087$, cluster-robust $t(236) = -1.3$, $p = .199$). The construct seemed reliable within-judge, as well, as participants' evaluation of their own behavior after the conversation (situational receptiveness) were strongly correlated with their earlier dispositional self-evaluation ($\beta = 0.51$, $SE = 0.06$, cluster-robust $t(236) = 8.9$, $p < .001$).

However, perceptions of receptiveness seemed less reliable. For example, day one dispositional receptiveness did not predict how someone was evaluated by their partner after day two ($\beta = -0.05$, $SE = 0.05$, cluster-robust $t(236) = 1.0$, $p = .32$). In fact, the correlation between self-reported and partner-evaluated receptiveness after the

conversation was relatively small ($\beta = 0.13$, $SE = 0.07$, cluster-robust $t(236) = 1.8$, $p = .08$). There also was perhaps a slight egocentric bias in the day-two evaluations, as participants rated themselves somewhat higher on receptiveness ($M = 5.39$, $SD = 0.95$) than they did their partners ($M = 5.25$, $SD = 1.26$), paired cluster-robust $t(237) = 2.0$, $p = .047$.

Linguistic receptiveness. We first used the detector model trained on Study 1 data to algorithmically identify the level of receptiveness in each participant's conversational behavior. Specifically, we concatenated all five rounds of each person's conversations to create a single document, which was then fed into the model. The resulting algorithmic labels were strongly correlated with partner-rated conversation receptiveness ($\beta = 0.29$, $SE = 0.06$, cluster-robust $t(236) = 4.6$, $p < .001$). Interestingly, they were not correlated at all with self-rated receptiveness, either situational ($\beta = 0.05$, $SE = 0.08$, cluster-robust $t(236) = 0.6$, $p = .54$) or dispositional ($\beta = -0.08$, $SE = 0.074$, cluster-robust $t(236) = 1.1$, $p = .275$). These findings are in line with the poor correlations between self-report and partner ratings. While others' evaluations of receptiveness relied on constant and reliable behavioral cues that our algorithm was able to also recognize, self-evaluations did not seem to track any behavioral pattern that we could observe in language.

We used the conversational data to improve our model in order to determine which conversational features best predicted how participants would be evaluated by their partners. In Fig. 2, we plot the linguistic features most closely correlated with conversational receptiveness in two panels, showing the features associated with self-evaluations and partner evaluations. As the validation results imply, the model of partner evaluations is much closer to the model we trained in Study 1. The model trained on self-evaluations suggests that there are in fact distinct and consistent features that lead people to believe they have

been receptive (even when they had not come across as such). Participants mainly seemed to give themselves credit for formalities (e.g., using titles, expressing gratitude, absence of swearing) rather than on demonstrable engagement with their partners' point of view.

Time course of receptiveness. While we only collected human receptiveness ratings before and after the conversation, one advantage of the linguistic model is that it allows us to examine how receptiveness modulated during the course of the conversation. We used the model to measure conversational receptiveness in each of the five rounds separately. In general, receptiveness increased over time (round 1: $M = 4.43$, $SD = 0.33$; round 5: $M = 4.61$, $SD = 0.39$; paired cluster-robust $t(237) = 6.1$, $p < .001$), although the correlation between receptiveness as measured by the algorithm and self-reported receptiveness does not increase (linear interaction: $\beta = 0.007$, $SE = 0.021$, cluster-robust $t(1186) = 0.3$, $p = .758$). Instead, across every round, we find a robust relationship between algorithmically-measured receptiveness and partner-rated receptiveness (all $p < .005$).

This design also allowed us to observe the dyadic and interactive nature of conversational receptiveness. After the conversation, partners' ratings of one another's receptiveness were correlated with one another within dyads ($\beta = 0.225$, $SE = 0.09$, cluster-robust $t(236) = 2.6$, $p = .009$), even though their day-one dispositional receptiveness had been uncorrelated. Interestingly, this pattern did not hold for self-rated conversation receptiveness: People in dyads together did not tend to rate themselves similarly after the conversation ($\beta = 0.010$, $SE = 0.09$, cluster-robust $t(236) = 0.1$, $p = .91$).

Our round-by-round measures of linguistic receptiveness help us shed some light on this process. Looking only at opening statements in the first round, we find no correlation between partners' linguistic receptiveness ($\beta = -0.02$, $SE = 0.09$, cluster-robust $t(236) = 0.2$, $p = .86$), which confirms the success of random assignment. However, in the remaining four rounds, partners clearly affected one another's language, as the algorithm's measure of their conversational receptiveness was highly correlated within dyads ($\beta = 0.378$, $SE = 0.096$, cluster-robust $t(236) = 3.9$, $p < .001$). Overall, these results show that the language dyad members used during the conversation influenced their partners' language.

Consequences of receptiveness. At the end of the interaction, we asked participants survey questions designed to measure two potential behavioral consequences of conversational receptiveness. First, expressed receptiveness could affect whether someone is willing to collaborate with a specific partner in the future. We tested this by combining three measures: desire to have the partner on a team, trust in their professional judgment, and willingness to have them represent your organization ($\alpha = 0.85$). Second, conversational receptiveness could affect how someone feels about engaging with disagreeing others more broadly, which we tested by combining the three measures regarding perceived value of disagreement ($\alpha = 0.81$; data missing for one person).

In Fig. 3, we compare several attitudinal and behavioral predictors of these consequences. Specifically, we estimate how each outcome measure is correlated with: day-one dispositional receptiveness (self and partner's); own and partner's self-reported situational receptiveness; the participant's ratings of their partner's receptiveness; and own and partner's conversational receptiveness, as measured by our model.

The best predictors of participants' collaboration intentions toward their partner tend to be their own responses to other questions. For example, a person's day-one self-rated dispositional receptiveness predicts both their desire for future collaboration with their partner ($\beta = 0.18$, $SE = 0.07$, cluster-robust $t(236) = 2.6$, $p = .01$), and their stated belief in the value of disagreement ($\beta = 0.27$, $SE = 0.075$, cluster-robust $t(235) = 3.6$, $p < .001$). This, of course, is not surprising, as people who report being more dispositionally receptive should be more open to contact with disagreeing others and value disagreement more than their less dispositionally receptive counterparts.

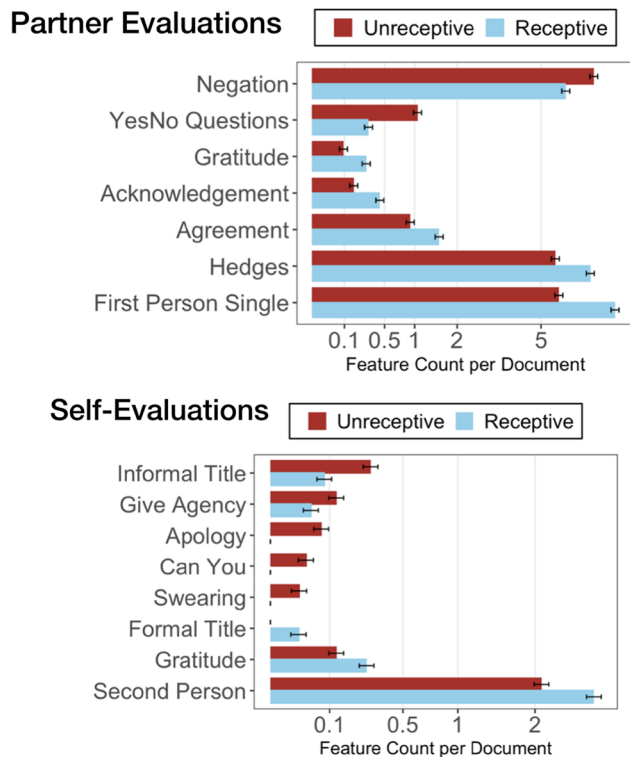


Fig. 2. Model of receptiveness trained on conversational data in Study 2. Groups were chosen to compare the top-third and bottom-third of responses, in terms of how receptive they were rated, on average. Each panel uses a different ground truth rating – either a person's rating of their own conversational receptiveness or the rating they received from their partner. Each datapoint represents a group mean (± 1 SE).

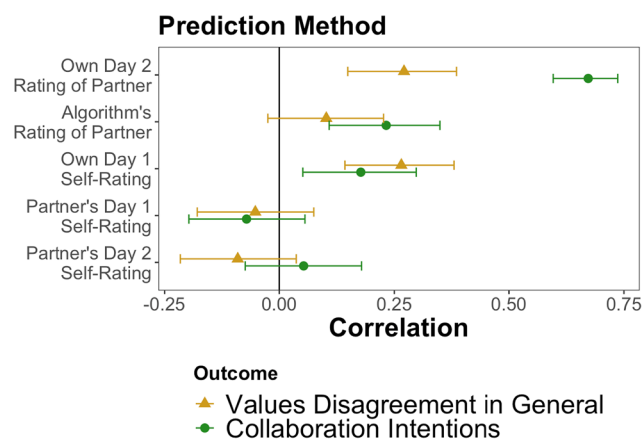


Fig. 3. Comparison of behavioral and dispositional predictors of the outcome measures in Study 2. Each datapoint represents a group mean (+/-1 SE).

The random assignment of partners does allow us to estimate some causal relationships from these data. For instance, a participant's willingness to collaborate with a partner in the future was strongly predicted by the algorithmic measure of the partner's receptiveness ($\beta = 0.23$, $SE = 0.06$, cluster-robust $t(236) = 4.0$, $p < .001$). This was also true controlling for the partner's day-one dispositional receptiveness (standardized $\beta = 0.23$, $SE = 0.06$, cluster-robust $t(235) = 4.1$, $p < .001$). By contrast, desire for future collaboration with a partner was not related to the partner's self-reported situational receptiveness ($\beta = 0.05$, $SE = 0.07$, cluster-robust $t(236) = 0.8$, $p = .42$) or the partner's self-reported dispositional receptiveness ($\beta = -0.07$, $SE = 0.06$, cluster-robust $t(236) = 1.2$, $p = .23$).

We find that a participant's belief about the value of disagreement is only weakly related to the algorithmically-rated receptiveness of their partner ($\beta = 0.10$, $SE = 0.06$, cluster-robust $t(235) = 1.9$, $p = .06$), including when controlling for the partner's dispositional receptiveness ($\beta = 0.10$, $SE = 0.05$, cluster-robust $t(234) = 1.9$, $p = .06$). These results suggest that when individuals express receptiveness through language, disagreeing conversation partners see them more positively across a variety of organizationally-relevant dimensions. Importantly, these perceptions are limited to the specific dyad partner exhibiting receptive behavior and do not inspire a more general appreciation for disagreement.

Interestingly, like partner-rated receptiveness, desire for future collaboration within a dyad tended to converge at the end of the conversation ($\beta = 0.193$, $SE = 0.09$, cluster-robust $t(236) = 2.2$, $p = .03$). Participants who expressed greater desire to collaborate with their partners had partners who also reported higher desire for collaboration with them. This pattern, however, was again not significant for expressed preference for disagreement between two dyad members ($\beta = 0.09$, $SE = 0.082$, cluster-robust $t(234) = 1.0$, $p = .3$).

3.3. Study 2 discussion

Study 2 data offer several important insights regarding the nature of conversational receptiveness. First, these results offer more validation of the linguistic model we developed in Study 1. The model transfers successfully to a new conversational context involving high-level government professionals, as well as to new issues. Furthermore, it could successfully distinguish between the behavioral cues associated with self-rated and partner-rated conversational receptiveness, illuminating the mistakes that people make in evaluating their own communication.

The results from this sample also demonstrate the important consequences of conversational receptiveness among policy professionals who regularly manage conflict. Participants who were rated as more receptive were perceived as more desirable team members, were seen as having better professional judgment, and were rated as more

desirable organizational representatives. Importantly, these results held using partner-rated receptiveness, as well as algorithmically-rated receptiveness. Whereas partner-rated receptiveness might reflect the rater's own characteristics, algorithmically-rated receptiveness reflects concrete and measurable features of language that can be exported across topics and individuals.

The rich two-party conversational data from this study also provided insights into the considerable differences in how the same conversational behavior is evaluated by the speaker and the listener. While partner-rated receptiveness was strongly correlated with algorithmically-rated receptiveness, self-rated receptiveness had almost no correlation with either of those measures. This suggests that people have a broken mental model of receptive conversational behavior – that is, one of the reasons so many people failed to be more receptive in Study 1 is because they were unaware of what linguistic behaviors would lead to being perceived as receptive. It is important to know that, in practice, receptiveness may be difficult to execute. In our final study, we unpack this broken mental model in greater detail.

4. Study 3: Consequences of receptiveness in the field

In Study 3 we turn to examining the consequences of conversational receptiveness in observational data in a domain where conflict naturally arises in pursuit of longer-term goals in globally distributed teams. Specifically, we investigate receptiveness in the talk pages on Wikipedia, one of the world's largest reference websites, where editors propose and debate revisions to existing pages.

Wikipedia is structured as a massively collaborative organization, with hundreds of thousands of active contributors editing articles in hundreds of languages (Kittur, Suh, Pendleton, & Chi, 2007; Reagle, 2010). By design, every Wikipedia article is considered a living document that can be modified by any of its editors. An elaborate system of checks and balances among the editors ensures that these changes are fair and accurate. In particular, modifications to longstanding articles are typically discussed and agreed upon beforehand on the “talk page” associated with the article.

These talk page discussions typically exemplify productive disagreement by soliciting alternative viewpoints and building consensus among attentive editors (Shi et al., 2019). However, discussions can also become hostile. In particular, Wikipedia has specific policies regarding personal attacks, a class of behaviors that “harm the Wikipedia community and the collegial atmosphere needed to create a good encyclopedia.”¹ When an editor's post is flagged as a personal attack, the page is evaluated to see if the interactions violate Wikipedia guidelines. Confirmed personal attacks can lead to temporary loss of editorial privileges, with repeated violations leading to more serious consequences.

In Study 3 we use a dataset of Wikipedia discussion threads to test whether conversational receptiveness can predict the occurrence of personal attacks in this setting. We leverage our algorithmic model of receptiveness to precisely measure receptiveness in naturally-occurring text and examine whether our construct can predict the occurrence of organizational conflict before it starts.

4.1. Study 3 methods

To study the relationship between receptiveness and personal attacks, we use a previously published dataset (Zhang, Chang, Danescu-Niculescu-Mizil, Dixon, Hua, Taraborelli, & Thain, 2018) in which the authors scraped Wikipedia to find examples of threads with personal attacks in the talk pages for popular articles. Crucially, for each thread containing a personal attack, the authors also identified a different thread for the same article (with a similar length and date) that did not

¹ From https://en.wikipedia.org/wiki/Wikipedia:No_personal_attacks

contain a personal attack. Thus, for each attack thread they created a synthetic control in a matched-pairs design (Iacus, King, & Porro, 2012). These threads were then verified by a research assistant to ensure each pair contained one (and only one) attack, leaving a final sample of 585 pairs of talk threads.

The dataset contained the entire contents of each thread, including messages posted after the attack itself. To create comparable “pre-attack” conversations within each thread pair, the full threads were shortened to the same number of turns ($M = 2.72$, $SD = 1.07$). That number was determined separately for each pair, as the minimum of either (a) the number of messages before the attack in the attack thread, or (b) the total number of messages in the non-attack thread. The resulting pre-attack conversation pairs contained a similar number of words in the attack threads ($M = 58.4$, $SD = 44.1$) as their matched control threads ($M = 59.4$, $SD = 50.2$; paired $t(585) = 0.3$, $p = .727$), as well as a similar number of users in the conversation (attack: $M = 2.19$, $SD = 0.44$; control: $M = 2.21$, $SD = 0.50$; paired $t(585) = 1.2$, $p = .216$).

Using the algorithm from Study 1, we calculated the receptiveness of every turn in every thread. We consider posts by the “original poster” who initially proposed an edit to a page, versus posts by other editors who are responding to the poster. In the threads that contain a personal attack, that attack can turn out to have been launched by the original poster against one of the other editors (where the original poster is the “attacker,”) or by one of the other editors against the original poster (where the original poster is the “victim”). We test whether the algorithmically-rated receptiveness of the first exchange in the thread predicts the ultimate occurrence of the attack. We further examine whether the conversational receptiveness of the attacker or the victim is the strongest predictor of conflict.

4.2. Study 3 results

In 48% of attack threads, the eventual attack was launched by the person who started the thread (the “original poster”) against someone else in the thread, while in the remaining 52% of threads, the attack was launched by someone else against the original poster. Accordingly, when the attacker is an original poster, we label the original poster's first post as the “attacker” turn, and the first post from the other editors as the “victim” turn. Likewise, when the attacker is another editor, the original poster's first post is labelled the “victim” turns, and the first post from the other editors is labeled as the “attacker” turns. Crucially, we align these labels in each control thread to the attack thread with which it is paired.

Using the algorithm from Study 1, we calculated the receptiveness of every turn in every thread. We found that 53.6% percent of the time, attackers expressed less receptiveness than their matched controls in the pre-attack conversation (95% CI = [51.6%, 55.6%]; $\chi^2(1) = 3.0$, $p = .082$). This suggests that receptive people are less likely to themselves escalate conflict later on. However, we also found that 60.3% of the time, the victims expressed less receptiveness than their matched controls (95% CI = [58.3%, 62.2%]; $\chi^2(1) = 24$, $p < .001$). Indeed, the difference in receptiveness between the victim and her matched control was consistently larger than the difference in receptiveness between the attacker and her matched control (paired $t(583) = 2.6$, $p = .009$). This suggests that editors who wrote less receptive posts were more likely to trigger conflict escalation in others. As a robustness check, we find the same basic results when we only look at the first post in every thread (attacker: $M = 54.1\%$, 95% CI = [52.1%, 56.1%], $\chi^2(1) = 3.9$, $p = .047$; victim: $M = 58.2\%$, 95% CI = [56.2%, 58.2%], $\chi^2(1) = 16$, $p < .001$; paired $t(583) = 2.4$, $p = .019$). In other words, our algorithm was able to predict organizational conflict during conversation before it occurs.

4.3. Study 3 discussion

We examined conversational receptiveness in a rich dataset from an online community where discussion covers a wide variety of topics, and where engagement with opposing views is central to the purpose of the organization. We found that people who wrote receptive editorial posts were less likely to be personally attacked later on in the thread. Thus, we show that conversational receptiveness can protect discussions between people who disagree from conversational conflict spirals (Weingart et al., 2015). Importantly, we find this result in a field setting, with posts written on a broad variety of topics. These results build on Study 2 and show that one of the most immediate benefits of expressing conversational receptiveness is prevention of interpersonal conflict in the midst of a disagreement.

5. Study 4: Testing a conversational receptiveness intervention

In our earlier studies, we demonstrated that variations in conversational receptiveness can be reliably detected by interaction partners, and that perceptions of higher receptiveness lead to positive interpersonal outcomes. However, our manipulation of receptiveness in Study 1 relied on a lengthy set of instructions about a complex psychological construct. In Studies 4a and 4b, we expand on our understanding of conversational receptiveness by testing four important questions: (1) How easy is it to manipulate receptiveness? (2) Does manipulated receptiveness lead to the same benefits observed with naturally-occurring conversational receptiveness? (3) Does receptiveness undermine persuasion? (4) Are people willing to use receptiveness in practice?

In Study 4A, we use a methodology similar to that used in Study 1A. “Responders” were asked to reply to a statement expressing a point of view opposed to their own on one of two randomly assigned social issues. We were interested in whether individuals would be able to enact receptiveness based on a simple set of instructions regarding the linguistic strategies identified by our algorithm. Thus, some of the responders were randomly assigned to first receive our focal intervention (“the receptiveness recipe”), which provided participants with instructions about the primary linguistic markers of receptiveness. The control condition was similar to Study 1A.

In Study 4B, we showed the text responses to a set of “raters” who evaluated both the responses themselves and the individuals who wrote them, while blind to the condition to which the respondents had been assigned. This set of studies allowed us to investigate the efficacy and practicality of a simple and scalable intervention targeting discrete linguistic markers for improving conflictual dialogue.

5.1. Study 4A methods

Sample. We recruited participants from mTurk for a study on “Political Issues and Discussions” ($M_{\text{age}} = 36.2$, 38.7% Male). As per our pre-registration, we only excluded participants who did not correctly complete our attention checks, reported “no opinion” on the issue they were asked to write a response on, or did not complete the study. This left a final sample of 771 responders.

Protocol. First, participants considered the two issue statements used in Study 1A (sexual assault on campus and police relations) and reported their agreement with the relevant policy statements on a scale from “-3: Strongly Disagree” to “+3: Strongly Agree.”

Participants were then randomly assigned to one of two conditions. In the control condition, participants again read about the discovery of a new species of fish and answered four comprehension questions. Once again, participants who offered an incorrect response to the quiz items were prompted to answer the item again and were dropped if they failed the same item three times. All conditions thus involved a quiz of some form to balance attrition and noncompliance across conditions. The quiz in the control condition was shorter than the quiz in Study 1a

to match the shorter intervention condition.

In the “receptiveness recipe” condition, participants read a four-item description of the linguistic markers that people can use to signal receptiveness to opposing views (as identified by our algorithm). Namely, participants read that responses judged to be receptive often include: positive statements, rather than negations; explicit acknowledgement of understanding; finding points of agreement; and hedging to soften claims. Participants were then presented with four comprehension questions on an unrelated topic (one’s preference for dogs versus cats) and were asked to identify which of two example statements better communicates receptiveness to the opposing view.

In the next phase of the experiment, participants were randomly assigned to read and respond to a statement on one of the two issues written by someone who holds the opposing view. These initial statements were the same as those used in Study 1A. Participants in the intervention condition were told to “Imagine that you are having an online conversation with this person. In your response, try to be as receptive and open-minded as you can.” Participants in the control condition were asked to “Imagine that you are having an online conversation with this person. How would you respond?” Participants spent at least two minutes writing their response to this disagreeing other.

After writing their response, participants in both conditions evaluated the communication style they had just used in their response. Specifically, they were asked to consider the conversations that they have in daily life with holders of opposing views and report (1) how hard it would be for them to adopt a similar communication style in future interactions and (2) how likely they were to adopt a similar communication style in these interactions. Participants responded on a scale from “1: Not at all” to “7: Extremely.”

Next, participants were (truthfully) told that we would ask a group of “raters” (who disagree with them on the issue) to read what they wrote and answer some questions about them. Participants were asked to take the perspective of these raters and predict how they would be evaluated on the 18-item receptiveness scale and the three-item measure of future collaboration intentions from Study 2. We also asked participants to predict how persuasive their written message would be on a –3 to +3 slider scale. Finally, participants reported basic demographic information including age, gender, and political orientation.

5.2. Study 4B methods

Sample. All participants were recruited from Amazon’s Mechanical Turk to participate in a study on “Political Issues and Discussions” ($M_{\text{age}} = 35.6$, 40% Male). By our pre-registered exclusion criteria, we only excluded participants who did not complete our attention checks, or reported “no opinion” on their assigned issue, or who did not complete the study. As in Study 1B, we collected data in several waves to achieve our target sample size of at least three raters for every responder. This left a final sample of 1,548 raters and an average of 4.0 raters for each response.

Protocol. The protocol for this study was nearly identical to that of Study 1B. Raters read messages from two different responders with whom they disagreed on the target issue (randomly assigned). Both responders were always answering relative to the same target statement on the same issue, and raters and responders always held opposing positions on the issue. The raters separately evaluated each responder using the 18-item receptiveness scale modified as in Study 1B (Appendix B; Minson, Chen, & Tinsley, 2019). In this study, we also added the three-item measure of future collaboration intentions used in Study 2 (i.e., desire to have the partner on a team, trust in their professional judgment, and willingness to have them represent your organization). Raters also evaluated the persuasiveness of the response on a –3 to +3 slider scale.

5.3. Study 4 results

Post-treatment attrition. We included irrelevant content in the control condition to try to equate post-treatment attrition and intervention time. However, we did not quite succeed in that goal, as there was a slight imbalance across both metrics in opposite directions. Post-treatment attrition was slightly greater in the intervention condition (17.2%) than the control condition (12.6%; $\chi^2(1) = 3.6$, $p = .056$), although this was not merely due to duration as the intervention itself was not longer than the control condition $\beta = -11.5$ s, $SE = 7.4$ s, $t(769) = 1.6$, $p = .121$). Still, this slight attrition bias could possibly affect our results, and smaller effect sizes should be interpreted with caution. However, we remain confident in effects that are too large to be caused by this bias. Accordingly, we still conduct all of our analyses among the people in each condition who passed our pre-registered exclusion criteria.

Receptiveness measures. Our raters again tended to agree on receptiveness, and the average of the raters’ evaluations for any piece of text was also correlated with the algorithm’s measure of receptiveness, as trained on Study 1 data ($r = 0.337$, $t(769) = 9.9$, $p < .001$). But the raters’ evaluations were not well predicted by either word count ($r = 0.032$, $t(769) = 0.9$, $p = .375$) or sentiment ($r = 0.120$, $t(769) = 3.4$, $p < .001$).

The main effects of the intervention are plotted in Fig. 4. The recipe instructions significantly affected the responses: participants who were instructed to be receptive were rated as more receptive by human raters than participants in the control condition (standardized $\beta = 0.57$, $SE = 0.07$, $t(769) = 8.3$, $p < .001$). The algorithm concurred: receptiveness as rated by the algorithm was significantly higher in the intervention condition than in the control condition (standardized $\beta = 0.51$, $SE = 0.07$, $t(769) = 7.3$, $p < .001$). Furthermore, the response writers anticipated this increase in perceived receptiveness (standardized $\beta = 0.36$, $SE = 0.07$, $t(751) = 5.0$, $p < .001$). Thus, our simple intervention successfully manipulated conversational receptiveness.

Collaboration Intentions. Importantly, the results of our measure of future collaboration intentions were similar to the receptiveness ratings. There was a difference in overall collaboration intentions expressed by raters who evaluated messages written by recipe condition versus control participants (standardized $\beta = 0.28$, $SE = 0.07$, $t(769) = 3.9$, $p < .001$). This difference was similar in magnitude to the difference in rated receptiveness. The response writers also anticipated this increase in raters’ desire to collaborate with them in the future (standardized $\beta = 0.25$, $SE = 0.07$, $t(753) = 3.4$, $p < .001$). Finally, across conditions, there was a consistent prediction bias: writers overestimated ($M = 4.81$, $SD = 1.36$) how much judges would

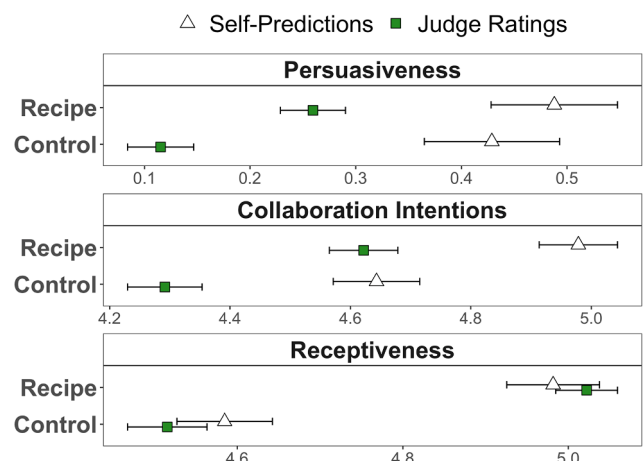


Fig. 4. Comparing actual (and predicted) ratings of responders in Study 4, by condition. Each datapoint represents a group mean (+/- 1 SE).

want to collaborate with them ($M = 4.45$, $SD = 1.18$; paired $t(754) = 5.5$, $p < .001$).

Persuasiveness. The results of the persuasiveness ratings were surprising, both to us and to the participants. We found that responses from the recipe condition were rated as more persuasive by ideological opponents than responses from the control condition (standardized $\beta = 0.24$, $SE = 0.07$, $t(769) = 3.3$, $p < .001$). However, raters did not seem to predict this effect (standardized $\beta = 0.05$, $SE = 0.07$, $t(753) = 0.7$, $p = .506$).

Intervention difficulty. After writing their responses, participants reported their own comfort with the interventions. As compared to the control condition, the intervention was rated as harder to execute in a future interaction with a disagreeing other ($\beta = 0.47$, $SE = 0.12$, $t(769) = 4.0$, $p < .001$). Likewise, participants reported that they were somewhat less likely to use the receptiveness strategy again, as compared to the control condition ($\beta = -0.17$, $SE = 0.10$, $t(769) = 1.7$, $p = .081$). However, the intervention did not lead participants to spend more time writing ($\beta = 9.2$ s, $SE = 13.4$ s, $t(769) = 0.7$, $p = .493$).

5.4. Study 4 discussion

Study 4 tested a new method for inducing conversational receptiveness. We found that our intervention clearly enhanced receptiveness, as perceived by disagreeing counterparts and by our algorithm. Furthermore, we found that when individuals were given specific instructions on how to enact conversational receptiveness, they were more able to do so effectively.

Similar to our Study 2 results, we found that rated receptiveness strongly predicted future collaboration intentions across all conditions. Importantly, our intervention led to both significantly higher levels of receptiveness and more positive collaboration intentions. This result is important in that it demonstrates that conversational receptiveness *causes* positive interpersonal outcomes that are crucial to organizational functioning. Participants who were instructed to enact four simple conversational strategies were seen as better team-mates, better organizational representatives, and as having better judgment.

Yet, participants also reported that—as compared to communicating naturally—following the recipe is difficult, and they would be less likely to communicate in a receptive manner in the future. This suggests that taking on a receptive conversational style can be challenging, or can perhaps compete with other conversational goals, such as appearing resolute in one's convictions. Future research should examine the extent to which training individuals in conversational receptiveness can be done with minimal effort, thus enhancing individuals' interpersonal experiences.

6. General discussion

Disagreement and conflict are an inevitable part of life. Whether discussing politics, business, science, or even mundane topics like where to go to dinner, we often find ourselves disagreeing with others' beliefs. Being able to engage with people who think differently than us is a critical decision-making and learning skill. Though people generally understand the value of considering opposing viewpoints with an open mind, handling disagreement is a difficult task. As a result, prior scholars have examined various interventions for improving conflictual dialogue.

In this paper, we extend this research by focusing on the construct of conversational receptiveness. Across five studies, we find that when a person converses with someone who holds an opposing viewpoint, it is possible to both communicate and modulate conversational receptiveness. Importantly, receptiveness can be communicated with a discrete set of linguistic features that are content- and topic-agnostic, and that could be incorporated into almost any conversation. These features are readily recognized and used by others to evaluate individuals' fitness for future cooperative goal pursuit. Interestingly, people are not as good

at recognizing these features in their own language.

In our initial study, we found that while there was some individual variation, there was considerable stable consensus on what people perceive as receptive. We were also able to build an interpretable natural language-processing algorithm that identified key features that lead to the perception of receptiveness. In Study 2, we saw that experienced government executives involved in a computer-based conversation evaluated their partners' professional characteristics more positively when they saw those partners as receptive. Importantly, the perceptions of receptiveness in Study 2 closely aligned with the levels of linguistic receptiveness identified by our algorithm. In Study 3, we find that conversational receptiveness prevents conflict escalation, suggesting that receptive dialogue can support productive and inclusive collaboration despite differences. In Study 4, we showed that an intervention that taught the results of our NLP model as a “recipe” could improve interpersonal outcomes in the midst of a disagreement. Overall, our results suggest that conversational receptiveness is measurable and has meaningful interpersonal consequences, but can be underutilized, in part, because speakers systematically misjudge their own receptiveness.

6.1. Theoretical contributions

Our research contributes to a budding literature on the psychology of conversation. While language has primarily been used as a measure (to communicate attitudes, values, personality, etc.), recent research has shown that the language expressed in conversation can be treated as a behavioral indicator of interpersonal goal pursuit (Huang, Yeomans, Brooks, Minson, & Gino, 2017; Jeong et al., 2019). Like so many other conversational goals, receptiveness is fraught with both errors of commission (saying things you shouldn't, such as explaining or contradicting) and errors of omission (not saying the things you should, such as acknowledgement or hedging). By using natural language processing, we build both a descriptive model of conversational receptiveness as well as a prescriptive one. In this sense, Fig. 1 provides a digestible and empirically informed template for improving conversations.

More broadly, our work contributes to a growing body of evidence that suggests people may hold their own behavior and others' behavior to different standards. This discrepancy offers an intriguing avenue for future research on the extent to which people can evaluate their own behavior from other people's perspectives (Vazire, 2010; Wilson & Dunn, 2004). Receptiveness seems particularly likely to generate such asymmetry: Other people are almost definitionally the best qualified to evaluate how we are treating them. However, the emotions and misunderstanding in conflict itself can also prevent people from knowing how they are seen by others. It is also possible that self-rated receptiveness is a stable but distinct construct, primarily determined by features of their internal mental states – e.g., what someone chooses not to say – that are unobservable to outsiders. By measuring linguistic behavior within the conversation itself, we could render a much clearer verdict on the sources of these misperceptions.

Finally, our results highlight an under-discussed element of recent efforts to improve civic discourse. Extensive writing has addressed “echo chambers,” in which the forces of selective exposure and homophily lead us away from engagement with opposing views (Sunstein & Hastie, 2015). Most recently, the importance of these ideas has been highlighted by technological advances that allow individuals to choose to expose themselves to a wider variety of views, or to even more selectively curate them (Barberá, Jost, Nagler, Tucker, & Bonneau, 2015; Flaxman, Goel, & Rao, 2016; Gentzkow & Shapiro, 2011; Yeomans et al., 2018). While academics often recommend more exposure to opposing viewpoints (Mendelberg, 2002; Pettigrew & Tropp, 2006), our results suggest that the effectiveness of these recommendations will be tempered by the contents of the resulting conversations (see Bail et al., 2019; Paluck et al., 2018). Simply choosing to

engage with opposing views may not lead to greater understanding or cooperation if the language of that engagement is unreceptive.

6.2. Limitations and future directions

We draw data from several domains and issues to confirm the generalizability of our model. However, our data only represent a small subset of the kinds of conversations that disagreeing people have. For example, in this paper we only looked at disagreements based on written text. However, many disagreements take place face-to-face, which complicates the decision environment. Synchronous conversation removes much of the time required for deliberation between turns on what conversational strategies to employ. Given the results of Study 4, which suggest that receptiveness is an effortful strategy, this raises the potential concern that people may respond more impulsively (and thus less receptively) in close conversations. On the other hand, there are many other potential channels that could signal receptiveness, including tone of voice, posture, backchannels, eye contact, accommodation, and so on. We believe that understanding the many other dimensions of receptiveness is an important goal for future research.

The interactions we studied were also somewhat limited, in that we only observed relationships that lasted for one conversation. However, longer relationships allow for more meaningful demonstrations of other kinds of behavior. When words do not align with actions, receptive language may seem like cheap talk. Alternatively, receptiveness may be a crucial step toward relational repair after a conflict episode. This is particularly relevant when the disagreement is asymmetric – for example, in customer service or human resources mediation, where one person involved in the disagreement is acting on behalf of an organization.

Appendix A. Description of interventions in Study 1

We review the content of the receptiveness interventions we gave to participants. The exact text and implementation (as exported Qualtrics files) are given in our online OSF repository.

Long Receptiveness Condition

Stimuli. First, participants read the following passage about receptiveness to opposing views:

“We are interested in how people interact with each other when discussing current “hot-button” policy and social topics. Specifically, we are interested in whether people can be receptive to others’ views. By “receptive” we mean being willing to read, deeply think about, and fairly evaluate the views of others, even if you disagree. In other words, can people be open-minded?

For example, imagine you are at a family dinner and your uncle starts going off about his views regarding immigration, which you strongly disagree with. A person who is not receptive might find this so aversive that they would leave the room or change the subject. If they were stuck listening to arguments they disagree with, they might mentally “tune out” and think about other things, or work really hard to find holes in the argument.

By contrast, a person who is being receptive would try to listen carefully to figure out where the uncle is coming from. They would not try to avoid the conversation, but would be genuinely curious and think hard about the reasons why somebody might hold these views in good faith.

We have developed a questionnaire to measure people’s receptiveness to opposing views. On the next page you will see examples of the questions we ask to understand how willing someone is to read, deeply think about, and fairly evaluate the views of others, even if they might disagree.

It is very important that you read this information carefully. Later you will be asked to answer questions based on what you learned and respond to another participant using the mindset that we describe.

According to our questionnaire, people who are receptive to opposing views would tend to agree with the following statements:
On the other hand, people who are receptive to opposing views would tend to disagree with the following statements:”

Quiz. Next, participants completed a comprehension quiz about the passage:

“How would someone who is receptive to opposing views feel about each of the following statements? Please answer the questions below the way a person who is receptive would answer.”

	Disagree	Agree
Listening to people with views that strongly oppose mine tends to make me angry.	☐	☐
I often get annoyed during discussions with people with views that are very different from mine	☐	☐
People who have views that oppose mine often base their arguments on emotion rather than logic.	☐	☐
I am generally curious to find out why other people have different opinions than I do.	☐	☐

Implementation instructions. Finally, participants read the following instructions before writing their response:

“In this next part of the survey, we are interested in how people interact with each other when discussing current “hot-button” policy and social topics.

7. Conclusion

The linguistic behavior that people exhibit in conversation can powerfully affect their partners’ perceptions, engagement, and willingness to cooperate with them. This research lays the groundwork for rigorously quantifying these linguistic decisions and their effects on social relationships. It also provides concrete, empirically based recommendations for more effective communication between people who disagree with each other. Four simple conversational behaviors can enhance conflictual conversations across organizations, families and friendships.

CRedit authorship contribution statement

Michael Yeomans: Conceptualization, Methodology, Data curation, Software, Formal analysis, Visualization, Validation, Writing - review & editing, Investigation, Project administration, Writing - review & editing. **Julia Minson:** Conceptualization, Methodology, Writing - review & editing, Investigation, Project administration, Writing - review & editing, Resources, Funding acquisition, Supervision. **Hanne Collins:** Writing - original draft, Investigation, Project administration, Writing - review & editing. **Frances Chen:** Writing - review & editing. **Francesca Gino:** Writing - original draft, Investigation, Project administration, Writing - review & editing, Resources, Funding acquisition, Supervision.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Participants in a previous survey were asked to write out statements describing their point of view on a topic assigned to them. We will show you one of these statements, along with an open-ended text box.

Please read the statement carefully and write a response to this person. In writing your response, try to be as receptive as possible. What would you say to this other person if you were trying to show that you are thinking hard about their point of view and trying to be fair in how you evaluate their perspective? What would you say if you were trying to show that you are working to understand where they are coming from?"

Control Condition

Stimuli. First, participants read the following passage:

"Donald Stewart's efforts to reclassify a giant Amazonian fish as representing several distinct species, rather than just one, are still ongoing. Stewart's latest work has just been published in a scientific journal, and marks official identification of *Arapaima leptosoma*, the first entirely new species of arapaima - a giant Amazonian fish - since 1847.

Among the world's largest freshwater fish, arapaimas, live in South America (Brazil and Guyana). They can grow up to 3 m long and weigh 200 kg. They breathe air through a primitive lung, and tend to live in oxygen-poor backwaters.

Arapaimas have long been an important food source for Amazonian peoples. They continue to be hunted and biologists have concerns about their status, although they are not endangered.

Getting the new species named is important "because it brings attention to the diversity of arapaimas that is out there and that need to be collected and studied," said Stewart.

For a century and a half, the prevailing view among scientists had been that there was only one species of arapaima, but Stewart has shown that there are actually at least five. In March 2014 he published a paper that renamed a species of arapaima that had been suspected in the 1800 s, before scientists decided to roll it up into one species.

The newest species, *Arapaima leptosoma*, had not been suspected before. It is more slender than other arapaimas. Its name, *leptosoma*, is a reference to its characteristic of slenderness. Stewart explained that the new species also has a horizontal black bar on the side of its head, which is a unique series of sensory organs.

The new species was described from a specimen kept at the Instituto Nacional de Pesquisas de Amazonia in Manaus, Brazil. That animal had been collected in 2001 near the confluence of the Solimoes and Purus rivers in Amazona State, Brazil.

Stewart added that he suspects there may be even more species of arapaima. "We keep finding other things out there," he said.

After publishing his paper in March 2014, Stewart received a great deal of media attention for his discovery of the newest species of fish, *Arapaima leptosoma*. In fact, this discovery could not have occurred at a better time. A recent study has revealed that these river giants are already extinct in some areas of Brazil.

Currently, the Purus River in Brazil is being aggressively fished to meet the demand of local markets. Consequently, few of these fish species are making their way to scientists or to aquariums where they can be preserved and studied.

More research about this new species of fish is needed in order to encourage conversation. As one scientist wrote, "This advance in taxonomic knowledge combined with the encouraging effectiveness of limiting numbers of mature fish allowed to be caught offer hope for these amazing river giants."

Quiz. Next, participants completed a comprehension quiz about the passage:

"Now that you have read about this new scientific discovery and considered how the information was communicated, please answer the following questions about the content of this passage.

1. Where do these arapaimas typically reside?
 - South America
 - Central America
 - North America
2. What best describes the body type of *Arapaima Leptosoma*?
 - They are very large in width, but very short in length
 - They are moderate in size at less than 1 m long, and weighing between 30 and 50 kg
 - They are exceptionally large at over 3 m long, and weighing 200 kg
3. What was the common belief of scientists in regards to Arapaimas - until now?
 - There was only one species of them
 - They are a significant food source for Amazonian peoples
 - They have become endangered
4. Why is it important to classify a new species of these fish?
 - Bring attention to the diversity of Arapaimas for future study
 - Bring attention to their endangered status
 - To have a complete taxonomy of Amazonian fish"

Implementation instructions. Finally, participants read the following instructions before writing their response:

"Now that you have read and considered one type of information communication, we would like you to read and consider another type of communication. Specifically, communication in the context of opposing views.

In this next part of the survey, we are interested in how people interact with each other when discussing current "hot-button" policy and social topics. Participants in a previous survey were asked to write out statements describing their point of view on a topic assigned to them. We will show you one of these statements, along with an open-ended text box."

Appendix B. Situational receptiveness rating scale

The questions below address the manner in which the respondent deals with contrary views and opinions on social and political issues. When

answering these questions, think specifically about the issue and discussion you just read. Please indicate the extent to which **you agree or disagree with [each] statement**.

1	2	3	4	5	6	7
Strongly Disagree	Somewhat Disagree	Slightly Disagree	Neither Agree nor Disagree	Slightly Agree	Somewhat Agree	Strongly Agree
1. On this issue, the respondent seems willing to have conversations with individuals who hold strong views opposite to their own. 2. On this issue, the respondent seems to like reading well thought-out information and arguments supporting viewpoints opposite to their own. 3. On this issue, the respondent seems like a person who finds listening to opposing views on this issue informative. 4. The respondent seems like a person who values interactions with people who hold strong views opposite to their own on this issue. 5. On this issue, the respondent seems generally curious to find out why other people have different opinions. 6. On this issue, the respondent seems to feel that people who have opinions that are opposite to theirs often have views which are too extreme to be taken seriously. 7. On this issue, the respondent seems to feel that people who have views that oppose theirs don't present compelling arguments. 8. On this issue, the respondent seems to feel that information from people who have strong opinions that oppose theirs is designed to mislead less-informed listeners. 9. On this issue, the respondent seems to feel that some points of view are too offensive to be equally represented in the media. 10. The respondent seems to feel that this issue is just not up for debate. 11. On this issue, the respondent seems to feel that some ideas are simply too dangerous to be part of public discourse. 12. The respondent seems to consider their views on this issue to be sacred. 13. On this issue, the respondent seems to feel that people who have views that oppose theirs are biased by what would be best for them and their group. 14. On this issue, the respondent seems to feel that people who have views that oppose theirs base their arguments on emotion rather than logic. 15. It seems that listening to people with views that strongly oppose their own views tends to make the respondent angry. 16. It seems that the respondent feels disgusted by some of the things that people with views that oppose theirs say regarding this issue. 17. On this issue, it seems that the respondent often feels frustrated when they listen to people with social and political views that oppose their own. 18. It seems that the respondent often gets annoyed during discussions of this issue with people with views that are very different from their own.						

Appendix C. Dispositional receptiveness scale

The questions below address the manner in which you deal with contrary views and opinions on social and political issues that are important to you. When answering these questions, think about the hotly contested issues in current social and political discourse (for example: universal healthcare, abortion, immigration reform, gay rights, gun control, environmental regulation, etc.). Consider especially the issues that you care about the most.

Please select the number below each statement to indicate the extent to which you agree or disagree with that statement.

1	2	3	4	5	6	7
Strongly Disagree	Somewhat Disagree	Slightly Disagree	Neither Agree nor Disagree	Slightly Agree	Somewhat Agree	Strongly Agree
1. I am willing to have conversations with individuals who hold strong views opposite to my own. 2. I like reading well thought-out information and arguments supporting viewpoints opposite to mine. 3. I find listening to opposing views informative. 4. I value interactions with people who hold strong views opposite to mine. 5. I am generally curious to find out why other people have different opinions than I do. 6. People who have opinions that are opposite to mine often have views which are too extreme to be taken seriously. 7. People who have views that oppose mine rarely present compelling arguments 8. Information from people who have strong opinions that oppose mine is often designed to mislead less-informed listeners. 9. Some points of view are too offensive to be equally represented in the media. 10. Some issues are just not up for debate. 11. Some ideas are simply too dangerous to be part of public discourse. 12. I consider my views on some issues to be sacred. 13. People who have views that oppose mine are often biased by what would be best for them and their group. 14. People who have views that oppose mine often base their arguments on emotion rather than logic. 15. Listening to people with views that strongly oppose mine tends to make me angry. 16. I feel disgusted by some of the things that people with views that oppose mine say. 17. I often feel frustrated when I listen to people with social and political views that oppose mine. 18. I often get annoyed during discussions with people with views that are very different from mine.						

Appendix D. Opposing views assignment algorithm

Our research questions centered around conversations between people who have opposing views. Studies 1 and 3 were asynchronous, allowing us to initially collect a common pool of participants, and then recruit their partners in the second phase of the study, matching them based on their

reported positions. Furthermore, we were not too concerned about whether an issue was imbalanced (e.g., 75% support, 25% opposed), because we could oversample in the second phase to make sure we collected enough partners holding the minority position to match every person holding the majority position in the first phase.

The design of Study 2 had no such affordances. This was part of an educational program, so every person needed one (and only one) partner. And the pool of people from which to match was not large (a few dozen per session). Furthermore, if we chose only a single issue, then any imbalance in positions would necessarily mean that some participants would not be matched with a disagreeing other. Although these matches can be generated by hand (as they were in the very first run of the program), this is a frantic and unaccountable protocol. Instead, we developed our own matching algorithm to produce a robust, systematic solution to this common experimental design problem.

There is a rich literature on two-sided matching algorithms (Becker, 1973; Roth & Sotomayor, 1992). Conceptually, this framework was useful for us, although we did not have any concerns that our participants were being strategic about reporting their issue sets. In practice, our case is closer to covariate matching in applied econometrics, which is used to control bias in treatment effect estimation, by selecting pairs of observations – one in treatment, one in control – that are otherwise balanced across one or several pre-treatment covariates (Dehejia & Wahba, 2002; Iacus et al., 2012). But to the best of our understanding, none of these algorithms apply to our current case: Whereas most markets feature positive assortative matching (i.e., high types with high types, and low with low), our objective was to choose pairs that are different from one another.

These algorithms typically define two objectives: “match quality,” or the similarity (or in our case, dissimilarity) of the matched pairs, and “yield,” or the number of matches. These two objectives often trade off against one another along a frontier; intuitively, a higher threshold for match quality will usually mean discarding data that does not have a suitable match. In our application, the primary matching objective was yield: We wanted to assign as many people as possible to a pair in which they met a threshold for disagreement on at least one of the issues. We had collected issue positions on a seven-point scale, so we defined our threshold for disagreement as a difference of at least four points on the scale – for example, a “1” could be paired with “5,” “6,” or “7”; a “2” could be paired with a “6” or “7”; and so on.

Intuitively, it should be obvious that some people will be easier to match than others. For example, someone who holds extreme minority views (e.g., a “7” when most people are “1,” “2,” or “3”) would have many potential matches, whereas someone in that majority has fewer potential matches, and someone who has no opinion (the “4” on the scale) cannot be matched with anyone. The most difficult matching cases are avoided by using three (relatively balanced) issues. However, this does not guarantee that any pool of people can be arranged to produce a complete set of satisfactory matches, and in any case, the computational complexity of generating this arrangement is not trivial.

Formally, our algorithm defines the “matchability” of a person as the total number of other people in the pool to whom they could be matched. At each iteration, the algorithm sorts everyone in the pool by matchability. Then the algorithm selects a pair from the pool, composed of (a) the person in the pool who is least matchable and (b) the least-matchable remaining person who can also be a match for (a). If the selected pair sufficiently disagrees on more than one issue, they are assigned to discuss the issue that has been assigned to the fewest number of pairs thus far. The two people in this generated pair are then removed from the pool. Then the whole process is repeated: Matchability is re-calculated for the remaining pool, the least-matchable remaining person is found and paired, and so on, until the whole pool has been partnered up. This guarantees that as the set of remaining matches dwindles, the last few people remaining will be the most-matchable, while also encouraging some balance across the three issues along the way.

Appendix E. Conversational receptiveness scales

Please select the number below each statement to indicate the extent to which **you agree or disagree with that statement**.

1	2	3	4	5	6	7
Strongly Disagree	Somewhat Disagree	Slightly Disagree	Neither Agree nor Disagree	Slightly Agree	Somewhat Agree	Strongly Agree

Receptiveness of self

The questions below address the manner in which you deal with contrary views and opinions on social and political issues. When answering these questions think specifically about the issue you just discussed.

1. On this issue, I am willing to have conversations with individuals who hold strong views opposite to my own.
2. On this issue, I like reading well thought-out information and arguments supporting viewpoints opposite to mine.
3. I find listening to opposing views on this issue informative.
4. I value interactions with people who hold strong views opposite to mine on this issue.
5. On this issue, I am generally curious to find out why other people have different opinions than I do.
6. I feel that people who have opinions that are opposite to mine on this issue often have views which are too extreme to be taken seriously.
7. On this issue, I feel that people who have views that oppose mine rarely present compelling arguments
8. I feel that information regarding this issue from people who have strong opinions that oppose mine is often designed to mislead less-informed listeners.
9. I feel that on this issue some points of view are too offensive to be equally represented in the media.
10. I feel that this issue is just not up for debate.
11. I feel that some ideas with regard to this issue are simply too dangerous to be part of public discourse.
12. I consider my views on this issue to be sacred.
13. I feel that people who have views that oppose mine on this issue are often biased by what would be best for them and their group.
14. I feel that people who have views that oppose mine on this issue often base their arguments on emotion rather than logic.
15. On this issue, listening to people with views that strongly oppose mine tends to make me angry.
16. I feel disgusted by some of the things that people with views that oppose mine say regarding this issue.
17. On this issue, I often feel frustrated when I listen to people with social and political views that oppose mine.

18. I often get annoyed during discussions of this issue with people with views that are very different from mine.

Receptiveness of partner

The questions below address the manner in which your partner deals with contrary views and opinions on social and political issues. When answering these questions think specifically about the issue you just discussed.

1. On this issue, my partner seems willing to have conversations with individuals who hold strong views opposite to their own.
2. On this issue, my partner seems to like reading well thought-out information and arguments supporting viewpoints opposite to their own.
3. On this issue, my partner seems like a person who finds listening to opposing views on this issue informative.
4. On this issue, my partner seems like a person who values interactions with people who hold strong views opposite to their own on this issue.
5. On this issue, my partner seems generally curious to find out why other people have different opinions than they do.
6. On this issue, my partner seems to feel that people who have opinions that are opposite to theirs often have views which are too extreme to be taken seriously.
7. On this issue, my partner seems to feel that people who have views that oppose theirs rarely present compelling arguments
8. On this issue, my partner seems to feel that information from people who have strong opinions that oppose theirs is often designed to mislead less-informed listeners.
9. On this issue, my partner seems to feel that some points of view are too offensive to be equally represented in the media.
10. My partner seems to feel that this issue is just not up for debate.
11. On this issue, I would guess that my partner seems to feel that some ideas are simply too dangerous to be part of public discourse.
12. My partner seems to consider their views on this issue to be sacred.
13. On this issue, my partner seems to feel that people who have views that oppose theirs are often biased by what would be best for them and their group.
14. On this issue, my partner seems to feel that people who have views that oppose theirs on this issue often base their arguments on emotion rather than logic.
15. It seems that listening to people with views that strongly oppose their own views tends to make my partner angry.
16. It seems that my partner feels disgusted by some of the things that people with views that oppose theirs say regarding this issue.
17. On this issue, it seems that my partner often feels frustrated when they listen to people with social and political views that oppose their own.
18. It seems that my partner often gets annoyed during discussions of this issue with people with views that are very different from their own.

Appendix F. Supplementary data

Supplementary data to this article can be found online at <https://doi.org/10.1016/j.obhdp.2020.03.011>.

References

- Allport, G. W. (1954/1979). *The nature of prejudice*. Reading, MA: Addison-Wesley.
- Bail, C., Argyle, L., Brown, T., Bumpus, J., Chen, H., Fallin Hunzaker, M., . . . Volfovsky, A. (2018). Exposure to opposing views on social media can increase political polarization. *Proceedings of the National Academy of Sciences of the United States of America*, 115(37), 9216–9221.
- Barberá, P., Jost, J. T., Nagler, J., Tucker, J. A., & Bonneau, R. (2015). Tweeting from left to right: Is online political communication more than an echo chamber? *Psychological science*, 26(10), 1531–1542.
- Baron, R. A. (1991). Positive effects of conflict: A cognitive perspective. *Employee Responsibilities and Rights Journal*, 2, 25–36.
- Becker, G. S. (1973). A theory of marriage: Part I. *Journal of Political economy*, 81(4), 813–846.
- Blunden, H., Logg, J. M., Brooks, A. W., John, L. K., & Gino, F. (2019). Seeker beware: The interpersonal costs of ignoring advice. *Organizational Behavior and Human Decision Processes*, 150, 83–100.
- Brehm, J. W. (1966). A theory of psychological reactance.
- Bruneau, E. G., Cikara, M., & Saxe, R. (2015). Minding the gap: Narrative descriptions about mental states attenuate parochial empathy. *PloS one*, 10(10).
- Bruneau, E. G., & Saxe, R. (2012). The power of being heard: The benefits of “perspective-giving” in the context of intergroup conflict. *Journal of experimental social psychology*, 48(4), 855–866.
- Chen, F. S., Minson, J. A., & Tormala, Z. L. (2010). Tell me more: The effects of expressed interest on receptiveness during dialogue. *Journal of Experimental Social Psychology*, 46, 850–853.
- Chen, M. K., & Rohla, R. (2018). The effect of partisanship and political advertising on close family ties. *Science*, 360(6392), 1020–1024.
- Cohen, S., Schulz, M. S., Weiss, E., & Waldinger, R. J. (2012). Eye of the beholder: The individual and dyadic contributions of empathic accuracy and perceived empathic effort to relationship satisfaction. *Journal of Family Psychology*, 26(2), 236.
- Coltri, L. (2010). *Alternative dispute resolution: A conflict diagnosis approach* (2nd ed.). Boston: Prentice Hall.
- Cutrona, C. E. (1996). *Social support in couples: Marriage as a resource in times of stress*. Thousand Oaks, CA: Sage.
- Day, M. V., Fiske, S. T., Downing, E. L., & Trail, T. E. (2014). Shifting liberal and conservative attitudes using moral foundations theory. *Personality and Social Psychology Bulletin*, 40(12), 1559–1573.
- De Dreu, C. K., & Van Vianen, A. E. (2001). Managing relationship conflict and the effectiveness of organizational teams. *Journal of Organizational Behavior: The International Journal of Industrial, Occupational and Organizational Psychology and Behavior*, 22(3), 309–328.
- Dehejia, R. H., & Wahba, S. (2002). Propensity score-matching methods for nonexperimental causal studies. *Review of Economics and statistics*, 84(1), 151–161.
- Dorison, C. A., Minson, J. A., & Rogers, T. (2019). Selective exposure partly relies on faulty affective forecasts. *Cognition*. In press.
- Edmondson, A. C., & Lei, Z. (2014). Psychological safety: The history, renaissance, and future of an interpersonal construct. *Annual Review of Organizational Psychology and Organizational Behavior*, 1(1), 23–43.
- Feinberg, M., & Willer, R. (2015). From gulf to bridge: When do moral arguments facilitate political influence? *Personality and Social Psychology Bulletin*, 41(12), 1665–1681.
- Ferrin, D. L., Bligh, M. C., & Kohles, J. C. (2008). It takes two to tango: An interdependence analysis of the spiraling of perceived trustworthiness and cooperation in interpersonal and intergroup relationships. *Organizational Behavior and Human Decision Processes*, 107(2), 161–178.
- Fiol, C. M. (1994). Consensus, diversity, and learning in organizations. *Organization Science*, 5(3), 403–420.
- Finkel, E. J., Slotter, E. B., Luchies, L. B., Walton, G. M., & Gross, J. J. (2013). A brief intervention to promote conflict reappraisal preserves marital quality over time. *Psychological Science*, 24(8), 1595–1601.
- Fitzsimons, G. J., & Lehmann, D. R. (2004). Reactance to recommendations: When unsolicited advice yields contrary responses. *Marketing Science*, 23(1), 82–94.
- Flaxman, S., Goel, S., & Rao, J. M. (2016). Filter bubbles, echo chambers, and online news consumption. *Public Opinion Quarterly*, 80(S1), 298–320.
- Frey, D. (1986). Recent research on selective exposure to information. *Advances in Experimental Social Psychology*, 19, 41–80. [https://doi.org/10.1016/S0065-2601\(08\)60212-9](https://doi.org/10.1016/S0065-2601(08)60212-9).
- Friedman, J., Hastie, T., & Tibshirani, R. (2010). Regularization paths for generalized linear models via coordinate descent. *Journal of statistical software*, 33(1), 1.
- Galinsky, A. D., & Moskowitz, G. B. (2000). Perspective-taking: Decreasing stereotype expression, stereotype accessibility, and in-group favoritism. *Journal of personality and social psychology*, 78(4), 708.
- Gentzkow, M., & Shapiro, J. M. (2011). Ideological segregation online and offline. *The Quarterly Journal of Economics*, 126(4), 1799–1839.
- Gerber, A. S., Huber, G. A., Doherty, D., & Dowling, C. M. (2012). Disagreement and the avoidance of political discussion: Aggregate relationships and differences across personality traits. *American Journal of Political Science*, 56(4), 849–874.
- Goldberg, J. H., Lerner, J. S., & Tetlock, P. E. (1999). Rage and reason: The psychology of

- the intuitive prosecutor. *European Journal of Social Psychology*, 29(5–6), 781–795.
- Goldsmith, D. J. (2000). Soliciting advice: The role of sequential placement in mitigating face threat. *Communications Monographs*, 67(1), 1–19.
- Gordon, A. M., & Chen, S. (2016). Do you get where I'm coming from? Perceived understanding buffers against the negative impact of conflict on relationship satisfaction. *Journal of Personality and Social Psychology*, 110(2), 239.
- Gottman, J. M. (1993). A theory of marital dissolution and stability. *Journal of family psychology*, 7(1), 57.
- Gottman, J. M. (1994). *What predicts divorce?* Hillsdale, NJ: Erlbaum.
- Gottman, J. M. (2008). Gottman method couple therapy. *Clinical handbook of couple therapy*, 4(8), 138–164.
- Gottman, J. M., & Levenson, R. W. (1999). Rebound from marital conflict and divorce prediction. *Family Process*, 38(3), 287–292.
- Gottman, J., Levenson, R., & Woodin, E. (2001). Facial expressions during marital conflict. *Journal of Family Communication*, 1(1), 37–57.
- Gottman, J. M., Notarius, C., Gonso, J., & Markman, H. (1976). *A couple's guide to communication*. Champaign, IL: Research Press.
- Graham, N., Arai, M., & Hagströmer, B. (2016). multiwayvcov: Multi-way standard error clustering. *R package version*, 1(2), 3.
- Graham, J., Haidt, J., & Nosek, B. A. (2009). Liberals and conservatives rely on different sets of moral foundations. *Journal of personality and social psychology*, 96(5), 1029.
- Gutenbrunner, L., & Wagner, U. (2016). Perspective-taking techniques in the mediation of intergroup conflict. *Peace and Conflict: Journal of Peace Psychology*, 22(4), 298–305.
- Hahlweg, K., Reisner, L., Kohli, G., Vollmer, M., Schindler, L., & Revenstorf, D. (1984). Development and validity of a new system to analyze interpersonal communication: *Kategoriensystem für Partnerschaftliche Interaktion*. Marital interaction: Analysis and modification 182–198.
- Halperin, E., Porat, R., Tamir, M., & Gross, J. J. (2013). Can emotion regulation change political attitudes in intractable conflicts? From the laboratory to the field. *Psychological Science*, 24(1), 106–111.
- Hameiri, B., Porat, R., Bar-Tal, D., & Halperin, E. (2016). Moderating attitudes in times of violence through paradoxical thinking intervention. *Proceedings of the National Academy of Sciences*, 113(43), 12105–12110.
- Hart, W., Albarracín, D., Eagly, A. H., Brechan, I., Lindberg, M. J., & Merrill, L. (2009). Feeling validated versus being correct: A meta-analysis of selective exposure to information. *Psychological Bulletin*, 135(4), 555–588. <https://doi.org/10.1037/a0015701>.
- Heydenberk, R. A., Heydenberk, W. R., & Tzenova, V. (2006). Conflict resolution and bully prevention: Skills for school success. *Conflict Resolution Quarterly*, 24(1), 55–69.
- Hillygus, D. S. (2010). Campaign effects on vote choice. *The Oxford Handbook of American elections and political behavior*, 326–345.
- Hirschman, A. O. (1970). Exit, voice, and loyalty: Responses to decline in firms, organizations, and states (Vol. 25). Harvard university press.
- Huang, K., Yeomans, M., Brooks, A. W., Minson, J. A., & Gino, F. (2017). It doesn't hurt to ask: Question-asking increases liking. *Journal of Personality and Social Psychology*, 113, 430–452.
- Iacus, S. M., King, G., & Porro, G. (2012). Causal inference without balance checking: Coarsened exact matching. *Political analysis*, 20(1), 1–24.
- Iyengar, S., & Westwood, S. J. (2015). Fear and loathing across party lines: New evidence on group polarization. *American Journal of Political Science*, 59(3), 690–707.
- Jeong, M., Minson, J. A., Yeomans, M., & Gino, F. (2019). *In high offers I trust: The effect of first offer value on economically vulnerable behaviors*. Psychological Science, under review.
- John, L. K., Jeong, M., Gino, F., & Huang, L. (2019). The self-presentational consequences of upholding one's stance in spite of the evidence. *Organizational Behavior and Human Decision Processes*, 154, 1–14.
- Johnson, D. W. (1971). Effectiveness of role reversal: Actor or listener. *Psychological Reports*, 28(1), 275–282.
- Keltner, D., Ellsworth, P. C., & Edwards, K. (1993). Beyond simple pessimism: Effects of sadness and anger on social perception. *Journal of personality and social psychology*, 64(5), 740.
- Kennedy, K. A., & Pronin, E. (2008). When disagreement gets ugly: Perceptions of bias and the escalation of conflict. *Personality and Social Psychology Bulletin*, 34(6), 833–848.
- King, D. B., & DeLongis, A. (2014). When couples disconnect: Rumination and withdrawal as maladaptive responses to everyday stress. *Journal of Family Psychology*, 28(4), 460.
- Kittur, A., Suh, B., Pendleton, B. A., & Chi, E. H. (2007, April). He says, she says: Conflict and coordination in Wikipedia. In Proceedings of the SIGCHI conference on Human factors in computing systems (pp. 453–462). ACM.
- Liberman, V., Anderson, N. R., & Ross, L. (2010). Achieving difficult agreements: Effects of positive expectations on negotiation processes and outcomes. *Journal of Experimental Social Psychology*, 46(3), 494–504.
- Liberman, V., Minson, J. A., Bryan, C. J., & Ross, L. (2012). Naïve realism and capturing the “wisdom of dyads”. *Journal of Experimental Social Psychology*, 48(2), 507–512.
- Liu, B. (2012). Sentiment analysis and opinion mining. *Synthesis lectures on human language technologies*, 5(1), 1–167.
- Livingstone, A. G., Fernández, L. R., & Rothers, A. (2019). “They just don't understand us”: The role of felt understanding in intergroup relations. *Journal of personality and social psychology*, in press.
- Lloyd, K. J., Boer, D., Keller, J. W., & Voelpel, S. (2015). Is my boss really listening to me? The impact of perceived supervisor listening on emotional exhaustion, turnover intention, and organizational citizenship behavior. *Journal of Business Ethics*, 130(3), 509–524. <https://doi.org/10.1007/s10551-014-2242-4>.
- 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(11), 2098–2109. <https://doi.org/10.1037/0022-3514.37.11.2098>.
- MacInnis, C. C., & Page-Gould, E. (2015). How can intergroup interaction be bad if intergroup contact is good? Exploring and reconciling an apparent paradox in the science of intergroup relations. *Perspectives on Psychological Science*, 10(3), 307–327.
- Mannix, E., & Neale, M. A. (2005). What differences make a difference? The promise and reality of diverse teams in organizations. *Psychological science in the public interest*, 6(2), 31–55.
- Manning, C. D., Manning, C. D., & Schütze, H. (1999). *Foundations of statistical natural language processing*. MIT press.
- McCroskey, J. C., & Wheelless, L. R. (1976). *Introduction to human communication*. Allyn and Bacon.
- Mendelberg, T. (2002). The deliberative citizen: Theory and evidence. *Political decision making, deliberation and participation*, 6(1), 151–193.
- Milton, J. (1644/1890) *Areopagitica: A speech of Mr. John Milton for the liberty of unlicensed printing to the parliament of England*. New York, NY: The Grolier Club.
- Minson, J., Chen, F., & Tinsley, C. H. (2019). Why won't you listen to me? Measuring receptiveness to opposing views. *Management Science*, in press.
- Minson, J. A., Liberman, V., & Ross, L. (2011). Two to tango: Effects of collaboration and disagreement on dyadic judgment. *Personality and Social Psychology Bulletin*, 37(10), 1325–1338.
- Neff, L. A., & Karney, B. R. (2004). How does context affect intimate relationships? Linking external stress and cognitive processes within marriage. *Personality and Social Psychology Bulletin*, 3.
- Nickerson, R. S. (1998). Confirmation bias: A ubiquitous phenomenon in many guises. *Review of General Psychology*, 2(2), 175–220. <https://doi.org/10.1037/1089-2680.2.2.175>.
- Pacilli, M., Roccato, M., Pagliaro, S., & Russo, S. (2016). From political opponents to enemies? The role of perceived moral distance in the animalistic dehumanization of the political outgroup. *Group Processes & Intergroup Relations*, 19(3), 360–373.
- Page-Gould, E., Mendoza-Denton, R., & Tropp, L. R. (2008). With a little help from my cross-group friend: Reducing anxiety in intergroup contexts through cross-group friendship. *Journal of personality and social psychology*, 95(5), 1080.
- Paluck, E. L., Green, S. A., & Green, D. P. (2018). The contact hypothesis re-evaluated. *Behavioural Public Policy*, 1–30.
- Pettigrew, T. F., & Tropp, L. R. (2006). A meta-analytic test of intergroup contact theory. *Journal of personality and social psychology*, 90(5), 751.
- Reagle, J. M. (2010). *Good faith collaboration: The culture of Wikipedia*. MIT Press.
- Repetti, R. L. (1989). Effects of daily workload on subsequent behavior during marital interaction: The roles of social withdrawal and spouse support. *Journal of personality and social psychology*, 57(4), 651.
- Ross, L., & Ward, A. (1995). Psychological barriers to dispute resolution. In M. Zanna (Vol. Ed.), *Advances in experimental social psychology: vol. 27*, (pp. 255–304). San Diego: Academic Press.
- Ross, L., & Ward, A. (1996). Naïve realism in everyday life: Implications for social conflict and misunderstanding. In E. S. Reed, E. Turiel, & T. Brown (Eds.), *Values and knowledge: The Jean Piaget symposium series* (pp. 103–135). Hillsdale, NJ, England: Lawrence: Erlbaum Associates Inc.
- Rogers, S., Howieson, J., & Neame, C. (2018). I understand you feel that way, but I feel this way: The benefits of I-language and communicating perspective during conflict. *PeerJ*, 6(5), E4831.
- Schroeder, J., & Risen, J. (2016). Befriending the enemy: Outgroup friendship longitudinally predicts intergroup attitudes in a coexistence program for Israelis and Palestinians. *Group Processes & Intergroup Relations*, 19(1), 72–93.
- Roth, A. E., & Sotomayor, M. (1992). *Two-sided matching. Handbook of game theory with economic applications, Vol. 1*, 485–541.
- Schroeder, J., Kardas, M., & Epley, N. (2017). The humanizing voice: Speech reveals, and text conceals, a more thoughtful mind in the midst of disagreement. *Psychological science*, 28(12), 1745–1762.
- Schweiger, D., Sandberg, W., & Rechner, P. (1989). Experiential effects of dialectical inquiry, devil's advocacy, and consensus approaches to strategic decision making. *Academy of Management Journal*, 32, 745–772.
- Sears, D. O., & Funk, C. L. (1999). Evidence of the long-term persistence of adults' political predispositions. *The Journal of Politics*, 61(1), 1–28.
- Sessa, V. I. (1996). Using perspective taking to manage conflict and affect in teams. *The Journal of applied behavioral science*, 32(1), 101–115.
- Shafraan-Tikva, S., & Kluger, A. N. (2016). Physician's listening and adherence to medical recommendations among persons with diabetes. *International Journal of Listening*, 1–10.
- Shi, F., Teplitskiy, M., Duende, E., & Evans, J. A. (2019). The wisdom of polarized crowds. *Nature Human behaviour*, 3(4), 329.
- Simmons, R. A., Gordon, P. C., & Chambliss, D. L. (2005). Pronouns in marital interaction: What do “you” and “I” say about marital health? *Psychological Science*, 16(12), 932–936.
- Soll, J. B., & Larrick, R. P. (2009). Strategies for revising judgment: How (and how well) people use others' opinions. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 35(3), 780.
- Stone, M. (1974). Cross-validatory choice and assessment of statistical predictions. *Journal of the Royal Statistical Society: Series B (Methodological)*, 36(2), 111–133.
- Story, L. B., & Repetti, R. (2006). Daily occupational stressors and marital behavior. *Journal of Family Psychology*, 20(4), 690.
- Sunstein, C. R., & Hastie, R. (2015). *Wiser: Getting Beyond Groupthink to Make Groups Smarter*. Harvard Business Press.
- Tjosvold, D., Morishima, M., & Belsheim, J. A. (1999). Complaint handling on the shop floor: Cooperative relationships and open-minded strategies. *International Journal of Conflict Management*, 10(1), 45–68.
- Todorova, G., Bear, J. B., & Weingart, L. R. (2014). Can conflict be energizing? A study of

- task conflict, positive emotions, and job satisfaction. *Journal of Applied Psychology*, 99(3), 451.
- Tost, L. P., Gino, F., & Larrick, R. P. (2013). When power makes others speechless: The negative impact of leader power on team performance. *Academy of Management Journal*, 56(5), 1465–1486.
- Turner, R. N., & Crisp, R. J. (2010). Imagining intergroup contact reduces implicit prejudice. *British Journal of Social Psychology*, 49(1), 129–142.
- Van Knippenberg, D., & Schippers, M. C. (2007). Work group diversity. *Annual Review of Psychology*, 58, 515–541.
- Vazire, S. (2010). Who knows what about a person? The self–other knowledge asymmetry (SOKA) model. *Journal of Personality and Social Psychology*, 98(2), 281.
- Wang, C. S., Kenneth, T., Ku, G., & Galinsky, A. D. (2014). Perspective-taking increases willingness to engage in intergroup contact. *PloS one*, 9(1), e85681.
- Weingart, L. R., Behfar, K. J., Bendersky, C., Todorova, G., & Jehn, K. A. (2015). The directness and oppositional intensity of conflict expression. *Academy of Management Review*, 40(2), 235–262.
- Westfall, J., Van Boven, L., Chambers, J. R., & Judd, C. M. (2015). Perceiving political polarization in the United States: Party identity strength and attitude extremity exacerbate the perceived partisan divide. *Perspectives on Psychological Science*, 10(2), 145–158.
- Wilson, T. D., & Dunn, E. W. (2004). Self-knowledge: Its limits, value, and potential for improvement. *Annual Review of Psychology*, 55, 493–518.
- Wojcieszak, M. E. (2012). On strong attitudes and group deliberation: Relationships, structure, changes, and effects. *Political Psychology*, 33(2), 225–242.
- Yeomans, M., Kantor, A., & Tingley, D. (2018). The politeness package: Detecting politeness in natural language. *R Journal*, 10(2).
- Yeomans, M., Stewart, B. M., Mavon, K., Kindel, A., Tingley, D., & Reich, J. (2018). The civic mission of MOOCs: Engagement across political differences in online forums. *International journal of artificial intelligence in education*, 28(4), 553–589.
- Zhang, J., Chang, J., Danescu-Niculescu-Mizil, C., Dixon, L., Hua, Y., Taraborelli, D., & Thain, N. (2018). Conversations gone awry: Detecting early signs of conversational failure. In Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics, 1, 1350–1361.