

The Same Message, Different Minds: Cognitive Appraisal Tendencies as Moderators of Framing Effects in Pro-Environmental Persuasion

Marion Moranetz*

Symbolic Systems Program

Advised by Professor Noah D. Goodman

Second Readers: Professor Robb Willer and Professor Judith Degen

June 4, 2025

Abstract

The same persuasive message often moves one person and not another. Matching accounts explain this by congruence between a message's framing and a stable property of the recipient, but they usually match on broad personality factors, which sit several inferential steps from the evaluation a frame actually triggers. This thesis asks whether cognitive appraisal tendencies—an individual's habitual way of evaluating situations along dimensions such as certainty and control (Smith & Ellsworth, 1985; Lerner & Keltner, 2001)—moderate framing effects, and whether they do so more precisely than personality traits. I first develop a Rational Speech Act model (Frank & Goodman, 2012) in which a listener interprets a framed appeal under an appraisal-conditioned prior; the model predicts a crossover, in which an individual-agency frame moves high-control listeners more and a collective frame moves low-control listeners more. I then test the prediction in a preregistered online experiment ($N = 471$). The frame had no main effect on attitude change ($d = 0.08$, $p = .38$). The predicted frame $\times$ control interaction was reliable but asymmetric ($b = 0.34$, $p = .004$): the high-control side was clearly significant, the low-control side only marginal. Control appraisal outperformed all five Big Five traits as a moderator. A regularized model predicting individual attitude change from appraisal and linguistic features reached modest out-of-sample accuracy (cross-validated $R^2 = 0.13$; AUC = 0.66). The effects are small and confined to one domain, self-reported attitudes, and an online sample, and the data depart from the model's symmetric prediction. The pattern is consistent with appraisal dimensions being a more legible moderator of framing than personality factors.

*Symbolic Systems Program, Stanford University. Email: mmoranetz@stanford.edu.

I am grateful to Noah Goodman, Robb Willer, and Judith Degen for their mentorship, and to the Computation & Cognition Lab for two years of patience with a project that kept changing shape.

To the Directors of the Program on Symbolic Systems:

I certify that I have read the thesis of Marion Moranetz in its final form for submission and have found it to be satisfactory for the degree of Bachelor of Science with Honors.

Electronically signed

Noah D. Goodman

June 4, 2025

Noah D. Goodman

Departments of Psychology and Computer Science

Electronically signed

Robb Willer

June 4, 2025

Robb Willer

Department of Sociology

Electronically signed

Judith Degen

June 4, 2025

Judith Degen

Department of Linguistics

Contents

1 Introduction	4
2 Background and Related Work	5
2.1 Matching effects in persuasion	5
2.2 The appraisal-tendency framework	5
3 A Model of Appraisal-Conditioned Interpretation	5
3.1 Setup	6
3.2 The predicted crossover	6
4 Methods	7
4.1 Participants	7
4.2 Design and procedure	7
4.3 Materials	7
4.4 Measures	7
4.5 Analysis plan	8
4.6 Deviations from preregistration	8
5 Results	8
6 Discussion	10
6.1 Limitations	10
6.2 Future directions	11
7 Conclusion	11
8 Acknowledgments	11
9 References	12
10 Appendices	13

1 Introduction

Two people read the same paragraph asking them to use less water. One closes the tap; the other turns off the tap while brushing their teeth that night. The message did not change between readings. Something about the readers did. This is the ordinary fact that makes persuasion hard to study: effectiveness is not a property of a message but of a message meeting a particular mind.

The standard response to this fact is the idea of a *matching* or *tailoring* effect. A persuasive appeal tends to work better when its framing is congruent with some stable property of the recipient (O’Keefe, 2016; Teeny, Siev, Briñol, & Petty, 2021). Appeals matched to a person’s regulatory focus (Cesario, Higgins, & Scholer, 2008), their moral foundations (Feinberg & Willer, 2015), or their Big Five profile (Hirsh, Kang, & Bodenhausen, 2012; Matz, Kosinski, Nave, & Stillwell, 2017) move them more than mismatched appeals do. The pattern is robust enough to support “psychological targeting,” in which ads are matched to traits inferred at scale (Matz et al., 2017).

The most common matching variable, though, is an odd one to match on. Traits like extraversion summarize a great deal of behavior but say little, by themselves, about how a person will evaluate a specific appeal. A trait describes a person; a frame succeeds or fails through an act of appraisal at the moment of reading. Matz et al. (2017) are explicit that they match on traits because traits are what can be inferred at scale from digital footprints, not because traits are the level of description a frame operates on. So the field’s workhorse moderator was chosen for measurability, not for its relationship to the appraisal a frame is meant to trigger.

This thesis asks whether a different moderator does better: *appraisal tendencies*, the chronic dispositions individuals carry to evaluate situations along a small set of cognitive dimensions (Lerner & Keltner, 2000, 2001; Smith & Ellsworth, 1985). The dimension I focus on is perceived control. My claim is that appraisal tendencies moderate framing effects, and do so more precisely than Big Five traits, because the appraisal dimensions sit closer to the operation a frame performs.

I argue this in two steps. Section 3 makes the mechanism explicit with a Rational Speech Act model (Frank & Goodman, 2012) in which a listener interprets a framed appeal under a prior set by their appraisal disposition. The model yields a specific, falsifiable prediction: a crossover in which the individual-agency frame moves high-control listeners more and the collective frame moves low-control listeners more. Sections 4 and 5 test that prediction in a preregistered experiment ($N = 471$) and compare appraisal against the Big Five as moderators of the same effect. The results are mixed in a way I try to take seriously. The frame’s main effect is null, as the model expects. The predicted interaction is present but asymmetric: strong on the high-control side, marginal on the low-control side, which the symmetric model does not predict. Control appraisal beats every Big Five trait as a moderator, and a regularized model recovers individual susceptibility at modestly above-chance accuracy. The contribution is the comparison and the asymmetry, not a deployable predictor.

2 Background and Related Work

2.1 Matching effects in persuasion

The Elaboration Likelihood Model (Petty & Cacioppo, 1986) established that the same appeal can be processed through different routes depending on the recipient's motivation and ability, which already implies recipient-dependence. Matching research made the dependence specific: messages persuade more when a structural property of the message lines up with a property of the audience. Regulatory-fit studies show promotion-framed appeals move promotion-focused readers more (Cesario et al., 2008). Moral-reframing work shows conservative audiences shift more on environmental policy when the appeal is framed around purity rather than harm (Feinberg & Willer, 2013, 2015). Personality-matched advertising shows measurable lifts when copy is matched to extraversion or openness (Hirsh et al., 2012; Matz et al., 2017).

These literatures share a logic but differ in what they match on, and the matching variable is rarely justified in terms of the cognition the frame engages. That gap—a moderator chosen for measurability rather than mechanism—is what motivates the model in Section 3.

2.2 The appraisal-tendency framework

Appraisal theories hold that evaluative responses arise from assessments of a situation along a small number of cognitive dimensions—pleasantness, certainty, attentional activity, control, anticipated effort, and responsibility (Smith & Ellsworth, 1985). The appraisal-tendency framework (Lerner & Keltner, 2000, 2001) extends this from transient states to stable dispositions: people differ in the appraisals they habitually bring to ambiguous situations, and these tendencies bias later judgments. Someone with a chronic sense of control reads an uncertain situation as more manageable; someone low in control reads the same situation as something that happens to them.

Control is the dimension with the cleanest mapping onto framing in the environmental domain, where appeals routinely split between "your choices matter" and "we have to act together." I treat certainty as a secondary, exploratory dimension. I do not claim appraisal tendencies are unrelated to personality—control correlates moderately with low neuroticism and with internal locus of control—only that they are a different, and I argue closer, level of description for predicting who a frame moves.

3 A Model of Appraisal-Conditioned Interpretation

Before running anyone, I wanted a reason to expect a particular interaction rather than just a hunch that "appraisal should matter." This section gives one. I adapt the Rational Speech Act (RSA) framework (Frank & Goodman, 2012; Goodman & Frank, 2016), which treats interpretation as Bayesian inference about a speaker's intended meaning, and add a single ingredient: the listener's prior is set by their appraisal disposition. The model is deliberately small. It is meant to fix the *sign* of an interaction, not to be fit to data, and Section 5 reports where the data refuse the model's symmetry.

3.1 Setup

Let $\theta \in [0, 1]$ be a listener's latent attitude toward the target behavior (here, conserving water), with 1 meaning maximally favorable. An appeal uses a frame f drawn from $\{\text{individual}, \text{collective}\}$. Each frame asserts that conservation is worthwhile for a reason located either in personal agency (individual) or in collective outcomes (collective).

A listener carries an appraisal disposition $a \in [0, 1]$, their chronic sense of control, which sets a prior over *where worth is located*: high- a listeners place more prior mass on personally controllable sources of value, low- a listeners on collective ones. Write this prior as a mixture weight $\pi(a)$ on the personal-agency component.

A literal listener updates on the frame's asserted reason:

$$L_0(\theta | f, a) \propto P(f | \theta) p(\theta | a).$$

A pragmatic listener additionally reasons that a cooperative speaker chose f to be relevant given what they could infer about the listener's prior:

$$L_1(\theta | f, a) \propto P_S(f | \theta, a) p(\theta | a), \quad P_S(f | \theta, a) \propto \exp(\lambda U(f; \theta, a)).$$

The key move is that a frame's utility U to the listener depends on the fit between the frame and the appraisal-conditioned prior: an individual-agency frame is more interpretable as relevant for a high-control listener, whose prior already locates value in personal action, and a collective frame is more relevant for a low-control listener. Formally, $U(\text{individual}; \cdot, a)$ increases in a and $U(\text{collective}; \cdot, a)$ decreases in a .

3.2 The predicted crossover

The quantity the experiment measures is the shift in the posterior mean attitude that a frame induces, $\Delta(f, a) = \mathbb{E}[\theta | f, a] - \mathbb{E}[\theta | a]$. Because the individual frame's utility rises in a and the collective frame's falls, the model implies

$$\frac{\partial \Delta(\text{individual}, a)}{\partial a} > 0, \quad \frac{\partial \Delta(\text{collective}, a)}{\partial a} < 0,$$

so the two frames' effect curves cross as control increases—the individual frame wins among high-control listeners, the collective frame among low-control listeners. Two further consequences matter for interpretation. First, averaged over the population the two frames need not differ, so the model predicts a *null main effect of frame* alongside the interaction. Second, in its plain form the model predicts a roughly *symmetric* crossover. I flag this now because it is the prediction the data later strain against: the observed crossover is real but lopsided, which Section 6 treats as the most informative result rather than a nuisance. The model is a scaffold for the sign of the effect; I do not fit its parameters (λ , the prior mixture) to the data, and I return in the limitations to what that restraint costs.

4 Methods

The design, hypotheses, exclusions, and primary analysis were preregistered on the Open Science Framework before data collection (osf.io/7q3kf). Deidentified data and analysis code are in the same project. All procedures were approved by the Stanford University IRB (Protocol #62831), and participants gave informed consent. I report all manipulations, measures, and exclusions (Simmons, Nelson, & Simonsohn, 2011).

4.1 Participants

I recruited 502 U.S. adults through Prolific, compensated at an effective rate of \$12/hour. Thirty-one were excluded per preregistered criteria: 14 failed an instructed-response attention check, 9 finished in under one-third of the median time, and 8 reported afterward that they had not read the appeal. The analyzed sample was $N = 471$ ($M_{\text{age}} = 38.4$, $SD = 13.1$; 51% women, 47% men, 2% nonbinary or undisclosed). A sensitivity analysis indicated 80% power to detect an interaction of $f^2 = 0.017$, a small effect.

4.2 Design and procedure

The design had a single between-subjects factor (frame: individual-benefit vs. collective-benefit), with appraisal tendencies and Big Five traits measured as continuous moderators. Baseline attitudes were measured at the start, before the moderator block, to reduce demand. Participants then completed the moderator measures, were randomly assigned to one appeal, and reported post-attitudes. The dependent measure was attitude change (post minus pre).

4.3 Materials

Both appeals (Appendix A) ran 184–188 words and were matched on reading level (Flesch–Kincaid grade 9.1 vs. 9.3), claim count, and the one statistic cited. They differed in framing: the individual-benefit version emphasized personal agency and household savings, the collective-benefit version emphasized shared stakes and coordinated action. A pilot ($n = 60$) confirmed the appeals differed on perceived individual-vs-collective framing ($t(58) = 6.9$, $p < .001$) but not on perceived argument strength, credibility, or warmth (all $p > .25$). One imperfection survived: the collective appeal was also rated slightly higher on urgency ($t(58) = 2.1$, $p = .04$), which I could not fully equalize without changing the statistic, and which I therefore carry into the analysis as a covariate rather than pretend away.

4.4 Measures

Appraisal tendencies. Chronic control and certainty appraisals were measured with an 8-item dispositional adaptation of Smith and Ellsworth's (1985) dimensions, on 7-point scales (Appendix B; control $\alpha = .78$). Certainty had marginal reliability ($\alpha = .66$), so I treat it as exploratory and read its results cautiously. *Personality.* The Big Five used the 44-item BFI (John & Srivastava,

1999). *Attitudes*. Conservation attitudes were a 4-item composite ($\alpha = .81$; Appendix C). *Linguistic style*. For the predictive model only, participants wrote 2–3 sentences about a recent decision; responses were scored for analytic and clout style with LIWC-22 (Boyd, Ashokkumar, Seraj, & Pennebaker, 2022).

4.5 Analysis plan

The confirmatory test was an OLS regression of attitude change on frame (effect-coded), mean-centered control appraisal, and their interaction, with certainty appraisal and the Big Five as covariates. The exploratory predictive model was an elastic-net regression (Zou & Hastie, 2005) predicting attitude change from all moderator and linguistic features, evaluated with 5-fold cross-validation; I report cross-validated R^2 and, for a median-split classification, AUC.

4.6 Deviations from preregistration

Three. (1) I had preregistered political ideology as a covariate, but a survey-flow error left it missing for 11% of participants, so I dropped it from the primary model and report an ideology robustness check on the reduced sample in the OSF supplement; it does not change the interaction. (2) The urgency covariate was not preregistered; I added it after the pilot flagged the imbalance. (3) I preregistered exclusion on a single attention check and added the self-reported non-reading item after piloting, a deviation toward more conservative exclusion. The frame $\times$ certainty interaction was exploratory throughout.

5 Results

Attitude change was small and positive overall ($M = 0.21$ scale points, $SD = 0.74$), as expected for a single online exposure. The two frames did not differ on average: individual-benefit $M = 0.24$ ($SD = 0.75$), collective-benefit $M = 0.18$ ($SD = 0.73$), $t(469) = 0.88$, $p = .38$, $d = 0.08$. The null main effect is what both the matching literature and the Section 3 model expect.

The confirmatory interaction was reliable. Frame interacted with chronic control appraisal, $b = 0.34$, $SE = 0.12$, 95% CI $[0.11, 0.57]$, $t(461) = 2.90$, $p = .004$, $\Delta R^2 = .018$. Adding the urgency covariate barely moved it ($b = 0.31$, $p = .007$). The crossover predicted in Section 3 appeared, but not symmetrically. Among participants one SD above the control mean, the individual frame beat the collective frame by $\Delta = 0.31$ ($p = .008$); among those one SD below, the collective frame led by only $\Delta = 0.17$ ($p = .09$), a marginal difference. The interaction is carried more by high-control participants than by low-control ones—the model's symmetric prediction is too clean (Figure 1). The exploratory frame $\times$ certainty interaction was in the predicted direction but not significant ($b = 0.13$, $p = .17$), and given certainty's marginal reliability I do not lean on it.

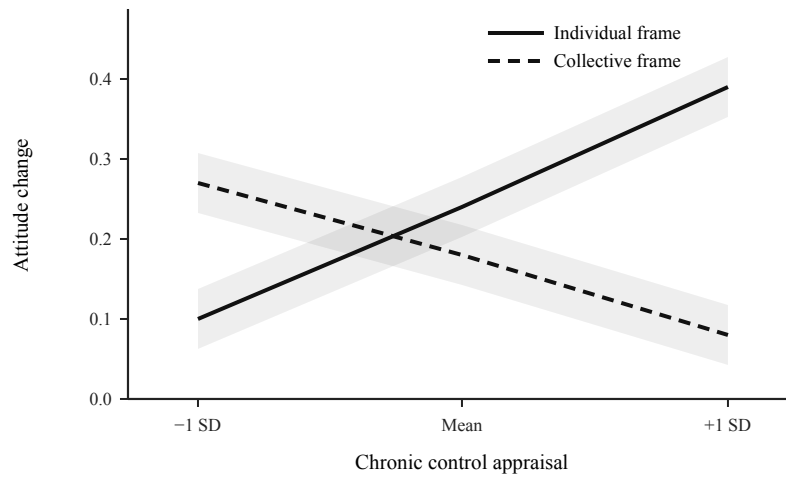


Figure 1: Observed frame $\times$ control-appraisal interaction (shaded bands = 95% CI). The crossover is present but asymmetric: the high-control gap is significant, the low-control gap only marginal.

One unpredicted result is worth reporting rather than burying. Age was a significant covariate: older participants changed more ($b = 0.05$ per decade, $p = .02$). I did not predict this and do not have a clean account for it; it does not interact with frame and leaves the control interaction intact, but it is the kind of thing a tidy story would omit, so I leave it in.

Moderator of frame effect	b	p	ΔR^2
Control appraisal	0.34	.004	.018
Certainty appraisal (exploratory)	0.13	.17	.004
Extraversion (BFI)	0.09	.34	.002
Openness (BFI)	0.11	.21	.003
Conscientiousness (BFI)	0.06	.51	.001

Table 1: Frame $\times$ moderator interactions, each controlling for the remaining moderators. Only control appraisal reaches significance.

Control appraisal also out-performed the Big Five head to head. Adding the frame $\times$ control term to a model that already contained all five frame $\times$ trait interactions improved fit, $\Delta R^2 = .015$, $F(1, 455) = 7.1$, $p = .008$; the reverse—adding the trait interactions to a model already containing frame $\times$ control—did not, $\Delta R^2 = .006$, $p = .43$.

Model (5-fold CV)	CV R^2	AUC
Appraisal + linguistic features	0.13	0.66
Big Five + linguistic features	0.07	0.59
Intercept only	0.00	0.50

Table 2: Out-of-sample prediction of attitude change. The appraisal model beats the trait model but is far from deployable.

In the retained elastic-net solution, control appraisal and its interaction with frame carried the two largest standardized coefficients; LIWC clout contributed a little; three of the five Big Five features were shrunk to zero. The model classifies above- vs below-median movers at $AUC = 0.66$. That is better than the trait model and better than chance, and it is nowhere near good enough to act on for any individual, which is the honest state of individual-level persuadability prediction.

6 Discussion

The same appeal moved different people differently, and the difference tracked chronic control appraisal more than any Big Five trait. The null main effect, the reliable control interaction, and the modest predictive model are each consistent with the Section 3 model, and the head-to-head comparison goes the way the model predicts: when appraisal and traits compete for the same variance, appraisal wins, and it wins on the specific dimension—control—that the individual-vs-collective frame was built to engage. That is the central result.

The more interesting result, to me, is where the data disobeyed the model. The crossover is asymmetric. High-control participants clearly preferred the individual frame; low-control participants preferred the collective frame only marginally. The plain RSA model in Section 3 predicts symmetry, so the asymmetry is information, not noise to apologize for. One reading is that the individual-agency frame is simply a stronger manipulation—personal-benefit language may recruit more elaboration than collective language regardless of appraisal—which would add a main-effect-like component on top of the interaction and flatten the low-control side. A second reading is that low-control individuals are less movable by a single online paragraph in general, compressing any frame effect at that end. The model cannot adjudicate these because I did not fit its parameters; doing so is the obvious next step and I say more about it below. For now the honest summary is that control appraisal is the right moderator and the symmetric mechanism is not quite right.

I want to be careful about size. The interaction explains under two percent of the variance in attitude change. That is a small effect on a weak proxy, and nobody should tailor anything on the strength of it. What the study supports is narrower and methodological: that appraisal dimensions are a better-aimed moderator than personality factors for this kind of frame, which is a claim about where to look next, not a claim about effect size.

6.1 Limitations

Five. *Single domain*: water conservation was chosen because its appeals split cleanly on individual vs. collective framing, and the interaction may not survive in domains where that split is less natural. *Attitudes, not behavior*: change after one exposure is a weak proxy, and the attitude-behavior gap is exactly where persuasion effects tend to disappear. *Self-reported appraisal*: a ques-

tionnaire measure of chronic appraisal inherits the usual concerns about insight and consistency, and certainty's marginal reliability ($\alpha = .66$) shows the measure is not airtight; a behavioral appraisal task would be a stronger test. *Sample*: Prolific participants are more attentive and more WEIRD than a general population, and a single session cannot speak to durability. *Model not fit*: Section 3 fixes the sign of the interaction but is not estimated against the data, so it cannot explain the asymmetry it failed to predict; that is a real gap, not a stylistic choice.

6.2 Future directions

Three steps follow directly. The cleanest is a within-subject design that shows each participant both frames across matched topics, separating the person-level moderator from between-subject noise. The most consequential is to move from attitudes to a behavioral outcome—a pledge, a default, a follow-up—and ask whether the appraisal interaction survives the attitude-behavior gap. The one closest to this project's Symbolic Systems core is to actually fit the model in Section 3: estimate λ and the appraisal-conditioned prior from data, let the fit absorb the asymmetry, and then test whether a frame generated to match an estimated appraisal profile beats a fixed control. I think that is the right next thesis, and I am aware it raises ethical questions about appraisal-level tailoring that I have waved at rather than answered; Vivrekar (2018) treats them at the length they deserve.

7 Conclusion

I started from a small, stubborn observation: the same message moves different people differently. A model said the difference should track how people appraise control, and a preregistered experiment found that it does, more cleanly than it tracks personality. The effect is small, the cross-over is lopsided in a way the model did not predict, and the predictor it yields is nowhere near usable for any individual. What survives is narrow and, I think, worth keeping: for this kind of frame, an appraisal dimension is a better-aimed moderator than a personality factor, because it sits closer to the cognitive operation the frame performs. That is a claim about where to look, and the next study—a fitted model tested against behavior—is the one that would tell whether the re-direction pays.

8 Acknowledgments

This thesis began as a much vaguer question, and Noah Goodman's first piece of advice—write the model down before you run anyone—is most of why it has a spine. He kept asking what I would actually be able to measure, and the answer kept getting smaller and better. Robb Willer's work on reframing is the reason the project exists, and his feedback kept me from over-claiming the result. Judith Degen talked me through the RSA prior in Section 3 over one office hour that I rewrote the whole model after. The Computation & Cognition Lab piloted an early survey and told me, correctly, that it was too long. The Symbolic Systems honors cohort read more drafts of

the introduction than anyone should have to, and Erin and Marcus in particular caught arguments I had talked myself into. The water-conservation framing came from a summer reading utility mailers in Ashland, Oregon; none of them worked on my dad, which is, in a roundabout way, the question. Any errors that remain are mine.

9 References

- Boyd, R. L., Ashokkumar, A., Seraj, S., & Pennebaker, J. W. (2022). *The development and psychometric properties of LIWC-22*. University of Texas at Austin.
- Cesario, J., Higgins, E. T., & Scholer, A. A. (2008). Regulatory fit and persuasion: Basic principles and remaining questions. *Social and Personality Psychology Compass*, 2(1), 444–463.
- Feinberg, M., & Willer, R. (2013). The moral roots of environmental attitudes. *Psychological Science*, 24(1), 56–62.
- 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.
- Frank, M. C., & Goodman, N. D. (2012). Predicting pragmatic reasoning in language games. *Science*, 336(6084), 998.
- Goodman, N. D., & Frank, M. C. (2016). Pragmatic language interpretation as probabilistic inference. *Trends in Cognitive Sciences*, 20(11), 818–829.
- Hirsh, J. B., Kang, S. K., & Bodenhausen, G. V. (2012). Personalized persuasion: Tailoring persuasive appeals to recipients' personality traits. *Psychological Science*, 23(6), 578–581.
- John, O. P., & Srivastava, S. (1999). The Big Five trait taxonomy. In L. A. Pervin & O. P. John (Eds.), *Handbook of personality* (pp. 102–138). Guilford Press.
- Lerner, J. S., & Keltner, D. (2000). Beyond valence: Toward a model of emotion-specific influences on judgement and choice. *Cognition & Emotion*, 14(4), 473–493.
- Lerner, J. S., & Keltner, D. (2001). Fear, anger, and risk. *Journal of Personality and Social Psychology*, 81(1), 146–159.
- Matz, S. C., Kosinski, M., Nave, G., & Stillwell, D. J. (2017). Psychological targeting as an effective approach to digital mass persuasion. *Proceedings of the National Academy of Sciences*, 114(48), 12714–12719.
- O'Keefe, D. J. (2016). *Persuasion: Theory and research* (3rd ed.). Sage.
- Petty, R. E., & Cacioppo, J. T. (1986). The elaboration likelihood model of persuasion. *Advances in Experimental Social Psychology*, 19, 123–205.
- Simmons, J. P., Nelson, L. D., & Simonsohn, U. (2011). False-positive psychology. *Psychological Science*, 22(11), 1359–1366.
- Smith, C. A., & Ellsworth, P. C. (1985). Patterns of cognitive appraisal in emotion. *Journal of Personality and Social Psychology*, 48(4), 813–838.

- Teeny, J. D., Siev, J. J., Briñol, P., & Petty, R. E. (2021). A review and conceptual framework for understanding personalized matching effects in persuasion. *Journal of Consumer Psychology*, 31(2), 382–414.
- Vivrekar, D. (2018). *Persuasive design techniques in the attention economy: User awareness, theory, and ethics* (Master's thesis). Stanford University.
- Zou, H., & Hastie, T. (2005). Regularization and variable selection via the elastic net. *Journal of the Royal Statistical Society: Series B*, 67(2), 301–320.

10 Appendices

Appendix A Persuasive appeals

Both appeals were shown with the same headline ("Your water, this summer") and the same statistic (residential use accounts for ~60% of summer municipal demand). Framing words are not italicized in the version participants saw.

Individual-benefit frame (186 words).

Your water, this summer. Most of us never think about the tap—until the bill arrives or the reservoir map turns brown. Here is the part worth knowing: residential use accounts for about 60% of summer municipal demand, which means the largest lever in this whole system is the one in your own house. The choices you make at home add up faster than people expect. A shorter shower, a full dishwasher instead of a half-empty one, a hose nozzle that actually shuts off—each is small, and together they can cut your household's summer water use by close to a third. That is money back on your bill and a yard that survives August, decisions you control directly and see the result of within a single billing cycle. You do not need anyone's permission or a new policy to start; you need a few habits that stick. The reservoir is not something that happens to you. It is, in a real and measurable sense, the sum of what each household decides to do this week. Start with one change tonight and let it compound.

Collective-benefit frame (184 words).

Your water, this summer. Most of us never think about the tap—until the bill arrives or the reservoir map turns brown. Here is the part worth knowing: residential use accounts for about 60% of summer municipal demand, which means this is a problem no single household solves alone, and one that every household helps solve together. The choices made across a neighborhood add up faster than people expect. Shorter showers, full dishwashers, hose nozzles that actually shut off—each is small, and when a whole community adopts them, regional reserves recover in a way no one home could manage by itself. That is a shared resource protected for the people around you and the summers ahead, the result of many households moving in the same direction at once. You are not being asked to fix this by yourself; you are being asked to be one of the many who make the collective effort work. The reservoir is not something that happens to us. It is, in a real and measurable sense, the sum of what a community decides to do this week. Join the rest of your neighborhood and start.

Appendix B Appraisal-tendency items

Dispositional adaptation of Smith & Ellsworth (1985). Each item rated 1 (*strongly disagree*) to 7 (*strongly agree*). Items marked (R) are reverse-scored. Presented in randomized order, interleaved with BFI items.

Control ($\alpha = .78$)

1. When I face an uncertain situation, I usually feel the outcome is within my power to shape.
2. In most situations I can find something I am able to do about the problem.
3. When things go wrong, it is usually because of forces outside my control. (R)
4. How a difficult situation turns out mostly depends on what I decide to do.

Certainty ($\alpha = .66$)

5. I usually feel I understand what is going on, even in ambiguous situations.
6. I can generally predict how a situation in front of me will unfold.
7. New situations often leave me unsure about what is actually happening. (R)
8. When information is incomplete, I still feel confident I know where things stand.

Appendix C Attitude and open-response measures

Conservation attitudes (4 items, $\alpha = .81$; 1–7). Administered at baseline and post-appeal; attitude change is the post–pre difference of the composite.

1. Cutting my household water use this summer is worth the effort.
2. My own water habits make a meaningful difference.
3. I intend to reduce my water use over the next month.
4. Conserving water is a responsible thing for someone like me to do.

Open-response prompt (LIWC features only).

"In 2–3 sentences, describe a recent decision you made—anything, large or small—and how you went about making it." Responses scored with LIWC-22 for Analytic and Clout; no responses were read for content.