

**Credibility of More vs. Less Precise Predictions Depends on
the Perceived Nature of Uncertainty**

Eitan D. Rude,¹ Amin Shiri,² Felipe M. Affonso,³ Hal E. Hershfield,^{1,4} Craig R. Fox^{1,4}

¹ University of California, Los Angeles – Anderson School of Management

² Texas A&M University – Mays Business School

³ Oklahoma State University – Spears School of Business

⁴ University of California, Los Angeles – Department of Psychology

Conditionally Accepted at the *Journal of Experimental Psychology: General*

Working Draft as of June 27, 2026

(see [here](#) for latest version)

Author Note

Eitan D. Rude  <https://orcid.org/0000-0001-5182-1421>

Amin Shiri  <https://orcid.org/0000-0003-1766-9344>

Felipe M. Affonso  <https://orcid.org/0000-0001-8928-5871>

Hal E. Hershfield  <https://orcid.org/0000-0002-1507-7022>

Craig R. Fox  <https://orcid.org/0000-0001-9057-112X>

We have no known conflicts of interest to disclose.

All preregistrations, materials, data, and code are available via ResearchBox at the following link: https://researchbox.org/3688?PEER_REVIEW_passcode=NBBUPL

We are grateful to attendees of the following conferences for their helpful feedback on earlier versions of this work: the Society for Judgment and Decision Making (2023), the California Schools Conference (2024), the Association for Consumer Research (2024), the Society for Consumer Psychology (2025), and the Haring Symposium (2025).

Correspondence regarding this article should be addressed to Eitan D. Rude, University of California, Los Angeles – Anderson School of Management, 110 Westwood Plaza, Room B401, Los Angeles, CA 90095, United States. E-mail: eitan.rude.phd@anderson.ucla.edu

Abstract

People expressing uncertain values frequently offer point predictions or intervals of varying widths (e.g., “I’d say your investment will earn \$6,000” versus “I’d say your investment will earn \$5,000–\$7,000”). Interestingly, prior research has drawn conflicting conclusions concerning the credibility of more versus less precise predictions. We resolve this inconsistency by identifying a key moderator: the perceived nature of uncertainty. Specifically, we find that more precise predictions are more credible under epistemic (“knowable”) uncertainty because they convey experts’ degree of confidence. Under aleatory (“random”) uncertainty, however, less precise predictions are more credible — if they are within reasonable bounds — because they convey experts’ assessments of outcome variability. These results inform literatures on both the communication of uncertainty and the interpretation of confidence intervals.

Keywords: Uncertainty, Forecasting, Credibility, Expertise

Public Significance Statement

When people make predictions (e.g., the future value of a stock), is it more credible for them to offer more precision (e.g., a point estimate, such as \$100) or less precision (e.g., a range of possible values, such as \$90–\$110)? Past research has offered conflicting answers to this question. We resolve this conflict by showing that the answer depends on recipients' intuitions concerning the source of uncertainty. Specifically, when recipients think uncertainty stems from insufficient knowledge or skill, they tend to find more precise predictions to be more credible. In contrast, when people think uncertainty stems from randomness, they tend to tolerate less precise predictions — though only up to a point. Thus, when experts choose a level of precision for expressing their forecasts, they should account for their audiences' perceptions of the source of relevant uncertainty, as this will critically influence their credibility.

Introduction

Experts often assess uncertain values for others: contractors estimate how long renovations will take, financial advisors forecast the performance of investments, meteorologists predict daily temperatures, and so forth. When conveying such numeric judgments, experts sometimes offer point predictions (e.g., “This job should require 60 hours of labor”) and sometimes offer intervals of varying widths (e.g., “This job should require 50–70 hours of labor”). Although past research has explored the association between precision and the credibility of forecasts, conclusions from these studies have been mixed: research on overconfidence has suggested that more precise predictions might be viewed as more credible, whereas research on the communication of uncertainty has generally suggested the opposite (see Table 1 for an overview). Here, we resolve this apparent conflict by proposing that the association between precision and credibility depends critically on recipients’ perceptions of the nature of relevant uncertainty.

Specifically, recent research has found that people intuitively distinguish “epistemic” uncertainty (“knowable” uncertainty, attributable to insufficient information/skill) from “aleatory” uncertainty (“random” uncertainty, attributable to chance/stochastic processes; Fox & Ülkümen, 2011). For example, the answer to a trivia question is knowable in principle (epistemic), and judges may express uncertainty about it due to insufficient information or skill; in contrast, the outcome of a chance gamble is inherently random (aleatory), and judges may express uncertainty about it given their beliefs about the underlying distribution of outcomes.¹ Perceptions of the nature of uncertainty (epistemicness and/or aleatoriness) are separate from

¹ Of course, most naturalistic events entail a mix of both epistemic and aleatory uncertainty. Further, while epistemic and aleatory uncertainty are not mutually exclusive, we treat them as complementary here and examine these dimensions separately — on a *post hoc* basis — in select analyses reported in the online supplement.

perceptions of the level of uncertainty, and they have wide-ranging implications. For example, perceptions of epistemicness and aleatoriness have been shown to influence the language used to describe uncertainty (Ülkümen et al., 2016), the extremity of probability judgments (Tannenbaum et al., 2017), decisions under ambiguity (Fox et al., 2021), investing decisions (Walters et al., 2023), evaluations of others' performance (Fox et al., 2026), and evaluations of sequential outcomes (Brimhall et al., 2026). However, extant work has not yet explored how the perceived nature of uncertainty influences the evaluation of more versus less precise predictions.

We note that prior studies examining how precision affects credibility have used predictions involving varying kinds of uncertainty. For example, studies arguing that greater precision confers greater credibility have tended to focus on quantities that may be more epistemic (e.g., estimating weights from photos; Radzevick & Moore, 2011; Van Zant, 2022), whereas studies arguing that less precision confers greater credibility have tended to focus on quantities that may be more aleatory (e.g., future sports scores; Gaertig & Simmons, 2018, 2023).

We propose that perceptions of the nature of uncertainty may help explain these discrepant findings because they affect how recipients interpret experts' expressed levels of precision. Specifically, we propose that when uncertainty is perceived to be primarily epistemic, recipients interpret precision as a signal of an expert's degree of confidence about a knowable outcome; in contrast, we propose that when uncertainty is perceived to be somewhat aleatory, precision is interpreted as a signal of an expert's beliefs about how widely a random outcome could vary.

We recruited 300 CloudResearch Connect participants to test this initial proposal. All participants considered a host of scenarios that naturally varied in perceived epistemicness and

aleatoriness (e.g., trivia, weather, gambling, etc.) where one expert provided a narrower interval and another provided a wider interval — conveyed visually using unmarked lines of different lengths. Next, participants chose whether they interpreted the different interval widths as indicating that: (a) the experts differed in how confident they were (“confidence” interpretation) or (b) the experts had different beliefs about how widely the outcome could vary (“variability” interpretation). All participants also rated their perceptions of the nature of uncertainty for each scenario on separate screens. As predicted, respondents were more likely to endorse the “confidence” interpretation when they saw uncertainty as more epistemic and the “variability” interpretation when they saw uncertainty as more aleatory ($b = 0.154$, 95% CI = [0.102, 0.206], $z = 5.860$, $p < .001$; Odds Ratio = 1.17, 95% CI = [1.11, 1.23]). Thus, this initial study (fully detailed in Supplement 1) supports our proposal that perceptions of the nature of uncertainty are associated with different interpretations of experts’ expressed levels of precision.

If precision signals confidence under epistemic uncertainty, then experts who provide wider intervals may appear less confident and therefore less credible. Conversely, if precision signals beliefs about how widely random outcomes could vary under aleatory uncertainty, then experts who provide wider intervals (at least up to a reasonable point) may in fact appear *more* credible.

We provide support for these predictions across three preregistered studies ($N = 1,185$), plus five additional preregistered studies reported in the online supplement ($N = 1,758$). Our results both reconcile contradictory findings from prior research and contribute to the growing literature on the perceived nature of uncertainty. All preregistrations, materials, data, and code are available via ResearchBox

(https://researchbox.org/3688?PEER_REVIEW_passcode=NBBUPL), and we note any deviations from our preregistrations where necessary.

Study 1

Study 1 provides correlational evidence for a link between the perceived nature of uncertainty and the relative credibility of point versus range predictions.

Participants

We recruited 300 CloudResearch-approved participants on Amazon Mechanical Turk (MTurk). We excluded participants who failed an attention check, yielding $N = 298$ participants ($M_{Age} = 40.39$, $SD_{Age} = 10.58$; 43.29% female, 55.70% male, 1.01% other).

Procedure

All participants considered three scenarios in a fixed order: (i) contractors estimating the number of hours required for a renovation (contractors scenario), (ii) financial advisors predicting returns on an investment (investments scenario), and (iii) meteorologists forecasting rainfall amounts (climate scenario). In each scenario, participants encountered one expert who provided a point prediction, and another who provided a range (see Table 2). Participants then evaluated the relative credibility of each expert using a five-item scale (adapted from Gaertig & Simmons, 2018) and indicated their preferences between them (see Table 3). On separate screens, participants indicated their perceptions of the nature of uncertainty in each scenario using the six-item Epistemic-Aleatory Rating Scale (EARS; Fox et al., 2026; see Table 4).

Results and Discussion

For each scenario, we constructed: (i) a “credibility” score by averaging responses from our five-item scale ($\alpha = 0.93$; higher values indicate that range predictions are relatively more credible); and (ii) an aleatory-epistemic difference score by averaging the aleatoriness and

epistemicness items from the EARS ($\alpha_{Aleatory} = 0.84$; $\alpha_{Epistemic} = 0.85$) and taking the difference between these averages (higher values indicate the perception that uncertainty in a given scenario is more aleatory and/or less epistemic). Next, we conducted a series of ordinary least squares (OLS) regressions to assess the relationship between the perceived nature of uncertainty and the relative credibility of point versus range predictions (as well as preferences between the corresponding experts) in each scenario.²

Figures 1A and 1B present scatterplots of each respondent's credibility and preference ratings, respectively, against aleatory-epistemic difference scores, for each scenario. As predicted, respondents deemed the range prediction more credible (and the point prediction less credible) to the extent they saw relevant uncertainty as more aleatory (and/or less epistemic). Specifically, we observed significant positive associations between the perceived nature of uncertainty (aleatory-epistemic difference scores) and the relative credibility of experts offering range versus point predictions for the contractors ($b = 0.288$, 95% CI = [0.202, 0.374], $t(296) = 6.577$, $p < .001$; $r = 0.357$, 95% CI = [0.254, 0.452]), investments ($b = 0.298$, 95% CI = [0.227, 0.369], $t(296) = 8.254$, $p < .001$; $r = 0.433$, 95% CI = [0.335, 0.521]), and climate ($b = 0.257$, 95% CI = [0.176, 0.338], $t(296) = 6.240$, $p < .001$; $r = 0.341$, 95% CI = [0.236, 0.438]) scenarios.

Respondents' preferences over experts echoed this pattern. Specifically, they favored the expert providing the range prediction over the expert providing the point prediction to the extent that they saw relevant uncertainty as more aleatory (and/or less epistemic) for the contractors ($b = 0.338$, 95% CI = [0.239, 0.437], $t(296) = 6.707$, $p < .001$; $r = 0.363$, 95% CI = [0.260, 0.458]), investments ($b = 0.361$, 95% CI = [0.284, 0.439], $t(296) = 9.141$, $p < .001$; $r = 0.469$, 95% CI =

² Our results hold when analyzing each of the items constituting our credibility measure separately, per our preregistration (see Supplement 2 for details). Our results also hold when treating aleatoriness and epistemicness as separate predictors (see Supplement 3 for details). We thank an anonymous reviewer for suggesting the latter *post hoc* analyses.

[0.376, 0.553]), and climate ($b = 0.269$, 95% CI = [0.177, 0.362], $t(296) = 5.718$, $p < .001$; $r = 0.315$, 95% CI = [0.209, 0.414]) scenarios.

In Supplement 4 we report *post hoc* robustness checks of these results,³ and in Supplements 5 and 6 we report two supplemental studies that extend these results to additional domains. Collectively, these observations provide strong correlational evidence for our hypotheses.

Study 2

Study 2 uses an experimental manipulation to provide evidence of a causal relationship between the perceived nature of uncertainty and the relative credibility of more versus less precision in the context of predictions about stock prices.

Participants

We recruited 1,000 CloudResearch-approved participants on MTurk. We excluded participants who failed comprehension and/or attention checks, yielding $N = 688$ participants ($M_{Age} = 42.49$, $SD_{Age} = 12.58$; 40.41% female, 58.14% male, 1.45% other).⁴

Procedure

Participants were randomly assigned to condition in a 2 (Chart: Absolute vs. Relative) $\times$ 2 (Format: Dollars vs. Percentages) between-subjects design. We varied “chart” as a manipulation

³ Specifically, we report the results of placebo tests to probe whether the general tendency to see events as more aleatory (and/or less epistemic) could influence our results independent of perceptions of the nature of uncertainty in a particular scenario. Critically, we observed that all matched correlations (i.e., correlations between credibility/preference ratings and aleatory/epistemic ratings in a particular scenario) were stronger than the mismatched correlations (i.e., correlations between credibility/preference ratings in a particular scenario and aleatory/epistemic ratings in another; see Supplement 4, Figures S2 and S3). Further, these differences were significant ($p < .05$) or marginally significant ($p < .10$) for 11 of the 12 possible comparisons (see Supplement 4, Table S9). We thank an anonymous reviewer for suggesting these *post hoc* analyses.

⁴ This large drop in sample size was largely due to a pre-planned comprehension check requiring participants to report details from the charts used in our manipulation. All remaining exclusions are attributable to a pre-planned attention check toward the end of the study (e.g., “... type ‘clock’ in the box next to ‘Other’”). Our results without these exclusions are qualitatively unchanged, and we report them in Supplement 7.

of perceived nature of uncertainty and “format” for robustness. Specifically, adapting a manipulation from Walters et al. (2023), we presented participants with a specific stock’s historical performance using either an “absolute” change chart (a time series of the stock’s price over time; see Figure 2A) or a “relative” change chart (a translation of the same time series into period-to-period changes in the stock’s price; see Figure 2B). The “absolute” chart’s discernible trend was designed to promote perceived epistemicness (knowability), and the “relative” chart’s scattered pattern was designed to promote perceived aleatoriness (randomness).⁵ Next, participants were asked to imagine that they were hiring a trader for an investment bank, and they encountered five candidates who made predictions concerning the future performance of the same stock in either dollar or percent terms. Their predictions spanned a point prediction (e.g., “Eli Lilly and Company[‘s] stock in one month’s time will [be \$575/change by +5%]”) and four successively wider range predictions centered on the same point (e.g., “Eli Lilly and Company[‘s] stock in one month’s time will [be between \$495 and \$655/change by –10% to +19%]”; see Table 5 for stimuli). All percent changes were equivalent to the corresponding dollar amounts (subject to rounding error); prediction format was merely manipulated to avoid possible confounds with chart type and we predicted no significant chart-format interactions. Next, participants chose — from ordered lists of the five candidates whose predictions they had just encountered — who was: (i) most credible, (ii) most competent, (iii) most thoughtful, and (iv) most trustworthy (together, these constituted our key dependent measures of “credibility,” and they appeared in a random order); on a separate screen, participants also chose which candidate they would prefer to hire (and we randomized the order of the credibility versus preference

⁵ Some readers may be concerned that this manipulation could influence perceived level of uncertainty alongside perceived nature, but we note that Walters et al. (2023) — the authors who originally developed the manipulation — successfully demonstrated that the charts only affected perceived nature of uncertainty without perturbing perceived level of uncertainty (see Study 2 of their paper as well as Study S1 of their Web Appendix).

measures). Finally, as a manipulation check, participants responded to the six-item EARS (see Table 4).

Results and Discussion

We first constructed: (i) a “credibility” score by treating responses to the four credibility items as ordinal measures (i.e., 1 = Point Prediction; 5 = Widest Range Prediction) and averaging them ($\alpha = 0.95$); (ii) an aleatory-epistemic difference score by averaging the aleatoriness and epistemicness items from the EARS ($\alpha_{Aleatory} = 0.80$; $\alpha_{Epistemic} = 0.86$) and taking the difference between these averages; and (iii) contrast-coded variables for chart ($-1 = \text{Relative}$; $+1 = \text{Absolute}$) and format ($-1 = \text{Percentages}$; $+1 = \text{Dollars}$).

We first probed our manipulation check: aleatory-epistemic difference scores across conditions. As predicted, respondents who reviewed the “absolute” chart (designed to promote perceived epistemicness) viewed the uncertainty in this domain as less aleatory (and/or more epistemic; $M = 0.200$, 95% CI = $[-0.023, 0.424]$) than those who reviewed the “relative” chart (designed to promote perceived aleatoriness; $M = 1.381$, 95% CI = $[1.176, 1.586]$). An OLS regression of aleatory-epistemic difference scores on chart confirmed that the effect of chart was significant ($b = -0.591$, 95% CI = $[-0.742, -0.439]$, $t(686) = -7.669$, $p < .001$; $d = 0.585$, 95% CI = $[0.432, 0.738]$). As expected, results did not differ by format (i.e., there was no chart-format interaction; $b = 0.067$, 95% CI = $[-0.084, 0.219]$, $t(684) = 0.875$, $p = .382$).

We next probed credibility. As predicted, respondents who reviewed the absolute chart (and thus viewed uncertainty as less aleatory and/or more epistemic) viewed candidates providing greater precision ($M = 2.957$, 95% CI = $[2.841, 3.073]$) as more credible than those who reviewed the relative chart ($M = 3.296$, 95% CI = $[3.175, 3.418]$). An OLS regression of credibility on chart confirmed that this difference was significant ($b = -0.170$, 95% CI = $[-0.254,$

$-0.085]$, $t(686) = -3.958$, $p < .001$; $d = 0.302$, 95% CI = $[0.151, 0.453]$), and again, results did not differ by format (i.e., there was no chart-format interaction; $b = -0.029$, 95% CI = $[-0.113, 0.056]$, $t(684) = -0.670$, $p = .503$).

Preferences echoed this pattern: respondents who reviewed the absolute chart preferred candidates providing greater precision ($M = 2.916$, 95% CI = $[2.794, 3.037]$) more than respondents who reviewed the relative chart ($M = 3.225$, 95% CI = $[3.096, 3.354]$). Again, an OLS regression confirmed that this difference was significant ($b = -0.155$, 95% CI = $[-0.243, -0.066]$, $t(686) = -3.426$, $p < .001$; $d = 0.261$, 95% CI = $[0.111, 0.412]$), and once more, results did not differ by format (i.e., there was no chart-format interaction; $b = -0.030$, 95% CI = $[-0.119, 0.059]$, $t(684) = -0.665$, $p = .507$).

In Supplement 8 we report two additional studies that extend these results to additional domains (and use an alternate manipulation). Overall, these observations provide evidence of a causal relationship between perceptions of the nature of uncertainty and the credibility of more versus less precise predictions.

Study 3

Studies 1 and 2 (along with studies reported in the online supplement) suggest that the perceived nature of uncertainty indeed influences the perceived credibility of more versus less precise predictions. Next, we probed the limits of how credibility judgments vary with precision (interval width) under epistemic versus aleatory uncertainty. Specifically, we predicted that when uncertainty is principally epistemic, wider intervals would convey less confidence/knowledge, so that credibility judgments would decrease monotonically with interval width. In contrast, we predicted that when uncertainty is somewhat aleatory, a reasonably (but not excessively) wide interval would convey appropriate recognition of inherent randomness, so that credibility

judgments would follow an “inverted-U” (i.e., negative quadratic) pattern — first increasing, then decreasing as interval width increased.

Participants

We recruited 199 CloudResearch-approved participants on MTurk ($M_{Age} = 40.74$, $SD_{Age} = 11.12$; 43.22% female, 56.78% male).

Procedure

All participants evaluated viewers’ judgments about the total number of points scored in two professional basketball games: one that had already been played (where uncertainty might be perceived as principally epistemic), and one that had yet to be played (where uncertainty might be perceived as somewhat aleatory; cf. Ülkümen et al., 2016, Study 1; we counterbalanced whether judgments about the past versus upcoming game were considered first). Specifically, participants were recruited after Game 4 (“epistemic” game) and before Game 5 (“aleatory” game) of the 2023 NBA Finals, and they considered the same series of judgments (i.e., intervals of varying widths) about the point totals for both games (see Table 6).⁶ After encountering each judgment, for each game, participants rated the relevant judge’s credibility, competence, and knowledge (see Table 7; together these constituted our measures of “credibility”). Next, participants explicitly chose one judge from each “pool” of judges (i.e., one judge of the past game’s score and one judge of the upcoming game’s score) to help them place bets on games in the subsequent NBA season (our secondary measure of “preference”).

Results and Discussion

As mentioned, we predicted that credibility ratings would monotonically decrease as interval width increased under epistemic uncertainty but would follow an inverted-U pattern

⁶ Importantly, prior beliefs concerning the scores should have been quite similar across games since the matchups were practically identical.

under aleatory uncertainty. We tested these predictions by analyzing respondents' mean ratings of each judge on our three-item credibility scale ($\alpha = 0.94$) and comparing the functional relationships between precision (interval width) and credibility across conditions (i.e., games). Visually, these data provide clear support for our hypotheses (see Figure 3) — i.e., there is a much stronger “inverted-U” pattern for the upcoming (somewhat aleatory) game than for the past (principally epistemic) game. We probed this difference statistically by regressing respondents' mean ratings of each judge upon a dummy-coded variable for game (0 = Epistemic; 1 = Aleatory), interval width, the square of interval width, and the interactions between condition and interval width along with the square of interval width (including respondent fixed effects and clustering standard errors by respondent). Our main preregistered outcome of interest was the interaction between condition (i.e., game) and the square of interval width. As predicted, this interaction was significant ($b = -0.0002$, 95% CI = $[-0.0003, -0.0001]$, $t(1,786) = -6.756$, $p < .001$), suggesting that the inverted-U pattern was indeed stronger for the upcoming game (where uncertainty might be perceived to be somewhat more aleatory) than for the past game (where uncertainty might be perceived to be principally epistemic). We next examined the simple linear effect of interval width for the epistemic game and observed that this effect was directionally consistent with our prediction ($b = -0.007$, 95% CI = $[-0.016, 0.002]$, $t(1,786) = -1.596$, $p = .112$). Further, in an alternate specification where we reverse-coded the condition dummy variable, we examined the simple quadratic effect of interval width for the aleatory game, and observed that this effect was significant and consistent with our prediction ($b = -0.0002$, 95% CI = $[-0.0003, -0.0002]$, $t(1,786) = -8.283$, $p < .001$). These latter two simple effects provide further support for our predictions concerning the functional relationship between credibility and precision under epistemic versus aleatory uncertainty.

For our secondary measure, we predicted that respondents would choose judges offering greater precision under epistemic uncertainty than under aleatory uncertainty since credibility should peak at wider intervals under the latter form of uncertainty. We tested this prediction by treating judge selections for each game as ordinal responses (i.e., 1 = Point Prediction; 5 = Widest Range Prediction) and regressing them on our dummy-coded variable for condition (including respondent fixed effects and clustering standard errors by participant). As predicted, respondents favored judges who provided more precise point totals for the past (epistemic) game ($M = 2.532$, 95% CI = [2.352, 2.714]) than the future (aleatory) game ($M = 3.065$, 95% CI = [2.917, 3.213]; $b = 0.533$, 95% CI = [0.369, 0.696], $t(198) = 6.432$, $p < .001$).

Altogether, these results support our assertion that there are distinct functional relationships between credibility and precision under epistemic versus aleatory uncertainty.

General Discussion

The present research helps explain prior discrepant findings concerning the credibility of more versus less precise predictions by identifying a key moderator: the nature of uncertainty. Specifically, we show that when forecasts concern epistemic (knowable) quantities, more precise predictions are more credible (presumably because they convey greater confidence/knowledge), but that when forecasts concern somewhat aleatory (random) quantities, less precise predictions are more credible (presumably because they acknowledge that outcomes can vary) up to a reasonable point.

This research has clear implications for experts seeking to project credibility when making predictions, but it may likewise provide insights concerning how people construct confidence intervals. Particularly, our research builds upon the observation that the level of precision communicated by experts signals their level of confidence under predominantly

epistemic uncertainty whereas it signals beliefs about outcome variability under somewhat aleatory uncertainty. Thus, interval judgments may be construed as “confidence” intervals under epistemic uncertainty and “variability” intervals under aleatory uncertainty. Whereas the present research focused on the impact of the perceived nature of uncertainty on recipients’ impressions of intervals of varying widths, future research should explore the impact of the perceived nature of uncertainty on experts’ choices of interval widths.

Naturally, prior beliefs should influence both the evaluation and construction of interval judgments. That is, domain-specific beliefs about the level of uncertainty should influence expectations about the “most credible” interval width in an absolute sense (i.e., the location of the maximum for the “inverted-U” documented in Study 3). While this is beyond the scope of the present investigation (which is chiefly concerned with the relative credibility of different interval widths), we likewise view this as a promising avenue for future research.

Finally, we acknowledge that the numeric judgments we explore here are not the only ways experts may choose to express uncertainty. For example, experts may provide: (i) more explicitly bounded confidence intervals (e.g., 90% confidence intervals); (ii) “best guesses” accompanied by a corresponding confidence band (e.g., $X \pm Y$); (iii) verbal hedges (e.g., “but that’s just a wild guess”); and so on. These and other ways of expressing uncertainty warrant further investigation as well.

Altogether, the present research highlights an important new domain in which the perceived nature of uncertainty influences judgment and choice. We hope that continuing research on this topic will help us better understand judgment, decision making, and communication under different variants of uncertainty.

Constraints on Generality

Our participants were all drawn from WEIRD samples (Henrich et al., 2010), recruited online. Results may differ across samples with different levels of education and/or numeracy, as well as for samples where linguistic or cultural norms give rise to different interpretations of uncertainty.

References

- Brimhall, C. I., Fox, C. R., Tannenbaum, D., & Imas, A. (2026). *Streaks as signals of skill* [Unpublished Working Paper].
- Du, N., Budescu, D. V., Shelly, M. K., & Omer, T. C. (2011). The appeal of vague financial forecasts. *Organizational Behavior and Human Decision Processes*, 114(2), 179–189. <https://doi.org/10.1016/j.obhdp.2010.10.005>
- Fox, C. R., Goedde-Menke, M., & Tannenbaum, D. (2021). Ambiguity Aversion and Epistemic Uncertainty. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.3922716>
- Fox, C. R., Tannenbaum, D., Ülkümen, G., Walters, D. J., & Erner, C. (2026). *Credit, blame, luck, and the perceived nature of uncertainty* [Unpublished Working Paper].
- Fox, C. R., & Ülkümen, G. (2011). Distinguishing Two Dimensions of Uncertainty. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.3695311>
- Gaertig, C., & Simmons, J. P. (2018). Do People Inherently Dislike Uncertain Advice? *Psychological Science*, 29(4), 504–520. <https://doi.org/10.1177/0956797617739369>
- Gaertig, C., & Simmons, J. P. (2023). Are people more or less likely to follow advice that is accompanied by a confidence interval? *Journal of Experimental Psychology: General*, 152(7), 2008–2025. <https://doi.org/10.1037/xge0001388>
- Henrich, J., Heine, S. J., & Norenzayan, A. (2010). The weirdest people in the world? *Behavioral and Brain Sciences*, 33(2–3), 61–83. <https://doi.org/10.1017/S0140525X0999152X>
- Howe, L. C., MacInnis, B., Krosnick, J. A., Markowitz, E. M., & Socolow, R. (2019). Acknowledging uncertainty impacts public acceptance of climate scientists' predictions. *Nature Climate Change*, 9(11), 863–867. <https://doi.org/10.1038/s41558-019-0587-5>
- Hu, B., Gaertig, C., & Dietvorst, B. J. (2024). How Should Time Estimates Be Structured to

- Increase Customer Satisfaction? *Management Science*, Articles in Advance, 1–19.
<https://doi.org/10.1287/mnsc.2023.00137>
- Jenkins, S. C., Harris, A. J. L., & Lark, R. M. (2019). When unlikely outcomes occur: The role of communication format in maintaining communicator credibility. *Journal of Risk Research*, 22(5), 537–554. <https://doi.org/10.1080/13669877.2018.1440415>
- Joslyn, S. L., & LeClerc, J. E. (2016). Climate Projections and Uncertainty Communication. *Topics in Cognitive Science*, 8(1), 222–241. <https://doi.org/10.1111/tops.12177>
- Radzevick, J. R., & Moore, D. A. (2011). Competing to Be Certain (But Wrong): Market Dynamics and Excessive Confidence in Judgment. *Management Science*, 57(1), 93–106.
<https://doi.org/10.1287/mnsc.1100.1255>
- Tannenbaum, D., Fox, C. R., & Ülkümen, G. (2017). Judgment Extremity and Accuracy Under Epistemic vs. Aleatory Uncertainty. *Management Science*, 63(2), 497–518.
<https://doi.org/10.1287/mnsc.2015.2344>
- Ülkümen, G., Fox, C. R., & Malle, B. F. (2016). Two dimensions of subjective uncertainty: Clues from natural language. *Journal of Experimental Psychology: General*, 145(10), 1280–1297. <https://doi.org/10.1037/xge0000202>
- van der Bles, A. M., van der Linden, S., Freeman, A. L. J., & Spiegelhalter, D. J. (2020). The effects of communicating uncertainty on public trust in facts and numbers. *Proceedings of the National Academy of Sciences*, 117(14), 7672–7683.
<https://doi.org/10.1073/pnas.1913678117>
- Van Zant, A. B. (2022). Strategically overconfident (to a fault): How self-promotion motivates advisor confidence. *Journal of Applied Psychology*, 107(1), 109–129.
<https://doi.org/10.1037/apl0000879>

Walters, D. J., Ülkümen, G., Tannenbaum, D., Erner, C., & Fox, C. R. (2023). Investor Behavior Under Epistemic vs. Aleatory Uncertainty. *Management Science*, 69(5), 2761–2777.

<https://doi.org/10.1287/mnsc.2022.4489>

Table 1

Summary of Prior Research. Alphabetical listing of prior studies speaking to recipients' preferences (or anticipated preferences) for more versus less precise predictions. For each highlighted article, we list whether the authors find: (a) a consistent preference for more precision (e.g., "point" predictions); (b) a consistent preference for less precision (e.g., "range" predictions); or (c) "mixed" results that reflect either a combination of both findings or inconclusive results. Finally, we list relevant domains considered in each article.

Paper	Preference (More Precision, Less Precision, Mixed)	Key Domain(s)
Du et al. (2011)	Mixed	Financial Forecasts
Gaertig & Simmons (2018)	Mixed	Outcomes of Sporting Events, Financial Forecasts
Gaertig & Simmons (2023)	Less Precision	Outcomes of Sporting Events, Effects of COVID-19
Howe et al. (2019)	Less Precision	Effects of Climate Change
Hu et al. (2024)	Less Precision	Delivery Times, ETAs
Jenkins et al. (2019)	Mixed	Effects of Climate Change
Joslyn & LeClerc (2016)	Less Precision	Effects of Climate Change
Radzevick & Moore (2011)	More Precision	Weight Estimation
van der Bles et al. (2020)	Mixed	Unemployment Statistics, Effects of Climate Change, Immigration Statistics
Van Zant (2022)	More Precision	Financial Forecasts, Weight Estimation, Professional Advice

Table 2

Text of Scenarios Used in Study 1. We counterbalanced whether expert “A” or “B” provided a point prediction or a range prediction across participants and scenarios.

Scenario	Prompt
Contractors	You are in the process of hiring a contractor to renovate your bathroom. As a part of that process, you solicited bids from two different local contractors who were highly-rated in online reviews. Although both charge the same rate of \$40/hour, they have provided two different estimates of how long the job will take (and therefore, how much the labor costs will be for the job). In your initial meetings, each contractor told you the following: • Contractor A told you that “I think this job will require 80 hours of labor, costing \$3,200” • Contractor B told you that “I think this job will require 70 to 90 hours of labor, costing \$2,800 to \$3,600”
Investments	You are looking to hire a financial advisor to help you invest in some individual stocks so that you can grow your money over time. You talk to two advisors, and ask both of them how much of a return you can expect to realize in 10 years’ time if you invest \$2,500 in NexusTech (NYSE: NXT) – a tech company that you have had your eye on. Each advisor replies as follows: • Advisor A tells you that “If you invest \$2,500 in NexusTech now, you should expect the value of your holdings to be \$6,000 in 10 years’ time” • Advisor B tells you that “If you invest \$2,500 in NexusTech now, you should expect the value of your holdings to be \$5,000 to \$7,000 in 10 years’ time”
Climate	You work for a government agency which manages water resources in the American south. In the coming year, the region is expected to experience “El Niño,” a weather phenomenon that is typically associated with more rain. In order to help better forecast and manage the increased rainfall over the course of the year, you are helping your employer hire a meteorologist as a consultant. You have narrowed your search down to two potential candidates, and in the course of interviewing them they each tell you the following: • Candidate A tells you that “Over the next year, you should plan on having 20 more inches of rainfall than you do on average” • Candidate B tells you that “Over the next year, you should plan on having 15 to 25 more inches of rainfall than you do on average”

Table 3

Dependent Measures Used in Study 1. Participants responded on a 1–6 scale where 1 = Definitely Advisor A and 6 = Definitely Advisor B, and the term “expert” was replaced with the relevant professional title. We analyzed the “Credibility” items and the “Preference” item separately, and we counterbalanced whether Advisor A or Advisor B had provided more precision (i.e., the point prediction) or less precision (i.e., the range prediction) in the text of the stimuli.

No.	Type	Item
1	Credibility	Which [expert] provided a more credible estimate?
2	Credibility	Which [expert] provided a more accurate estimate?
3	Credibility	Which [expert] provided an estimate reflecting more careful thought?
4	Credibility	Which [expert] provided a more honest estimate?
5	Credibility	Which [expert] provided a more realistic estimate?
6	Preference	Which [expert] would you rather hire?

Table 4

Six-Item Epistemic-Aleatory Rating Scale (EARS; Fox et al., 2026). Participants responded on a 1–7 scale where 1 = Not at all and 7 = Very much. Item order was randomized and each time the scale was presented it was preceded by a stem asking participants to contemplate the nature of uncertainty associated with the event/prediction in question (see ResearchBox for details and materials).

Subscale	No.	Item
Epistemic	1	...is knowable in advance, given enough information.
	2	...is something that becomes predictable with additional knowledge or skills.
	3	...is something that well-informed people would agree on.
Aleatory	1	...is determined by chance factors.
	2	...could play out in different ways on similar occasions.
	3	...is something that has an element of randomness.

Table 5

Text of Predictions Presented in Study 2. Participants were randomly assigned to view either the “Dollars” format or the “Percentages” format. For all participants, we counterbalanced whether intervals were ordered from narrowest to widest or widest to narrowest.

Candidate	Format	
	Dollars	Percentages
A	“I expect that the price of Eli Lilly and Company stock in one month’s time will be \$575”	“I expect that the price of Eli Lilly and Company stock in one month’s time will change by +5%”
B	“I expect that the price of Eli Lilly and Company stock in one month’s time will be between \$565 and \$585”	“I expect that the price of Eli Lilly and Company stock in one month’s time will change by +3% to +6%”
C	“I expect that the price of Eli Lilly and Company stock in one month’s time will be between \$555 and \$595”	“I expect that the price of Eli Lilly and Company stock in one month’s time will change by +1% to +8%”
D	“I expect that the price of Eli Lilly and Company stock in one month’s time will be between \$535 and \$615”	“I expect that the price of Eli Lilly and Company stock in one month’s time will change by –3% to +12%”
E	“I expect that the price of Eli Lilly and Company stock in one month’s time will be between \$495 and \$655”	“I expect that the price of Eli Lilly and Company stock in one month’s time will change by –10% to +19%”

Table 6

Judgments Presented to Participants in Study 3. All judgments concerned Games 4 and 5 of the 2023 NBA Finals, made after the conclusion of Game 4 and before the start of Game 5.

Bracketed text indicates wording that differed for the past versus upcoming game. Note that the point prediction corresponds to the actual outcome observed in Game 4, given that the game had already concluded by the time of data collection. We counterbalanced whether these intervals were ordered from narrowest to widest — or vice versa — at the participant level.

Judge	Judgment
A	“I think that the Denver Nuggets and the Miami Heat, combined, [scored/will score] <u>exactly</u> 203 points [this past Friday/tonight]”
B	“I think that the Denver Nuggets and the Miami Heat, combined, [scored/will score] somewhere between 200-205 points [this past Friday/tonight]”
C	“I think that the Denver Nuggets and the Miami Heat, combined, [scored/will score] somewhere between 190-210 points [this past Friday/tonight]”
D	“I think that the Denver Nuggets and the Miami Heat, combined, [scored/will score] somewhere between 170-230 points [this past Friday/tonight]”
E	“I think that the Denver Nuggets and the Miami Heat, combined, [scored/will score] somewhere between 140-260 points [this past Friday/tonight]”

Table 7

Credibility Measures Used in Study 3. After encountering each judge, participants used these measures to evaluate them on a 1–7 scale where 1 = Not at all and 7 = Extremely.

Item	Text
Credibility	How credible is this person?
Competence	How competent is this person?
Knowledge	How knowledgeable is this person?

Figure 1A

Credibility Results from Study 1. Each panel presents a scatterplot of respondents' perceptions of the nature of uncertainty (aleatoriness – epistemicness) against their relative credibility ratings of judges providing less precise predictions (i.e., range predictions) versus more precise predictions (i.e., point predictions; higher scores indicate that judges providing wider range predictions appear more credible). Error bands surrounding the regression lines represent 95% confidence intervals.

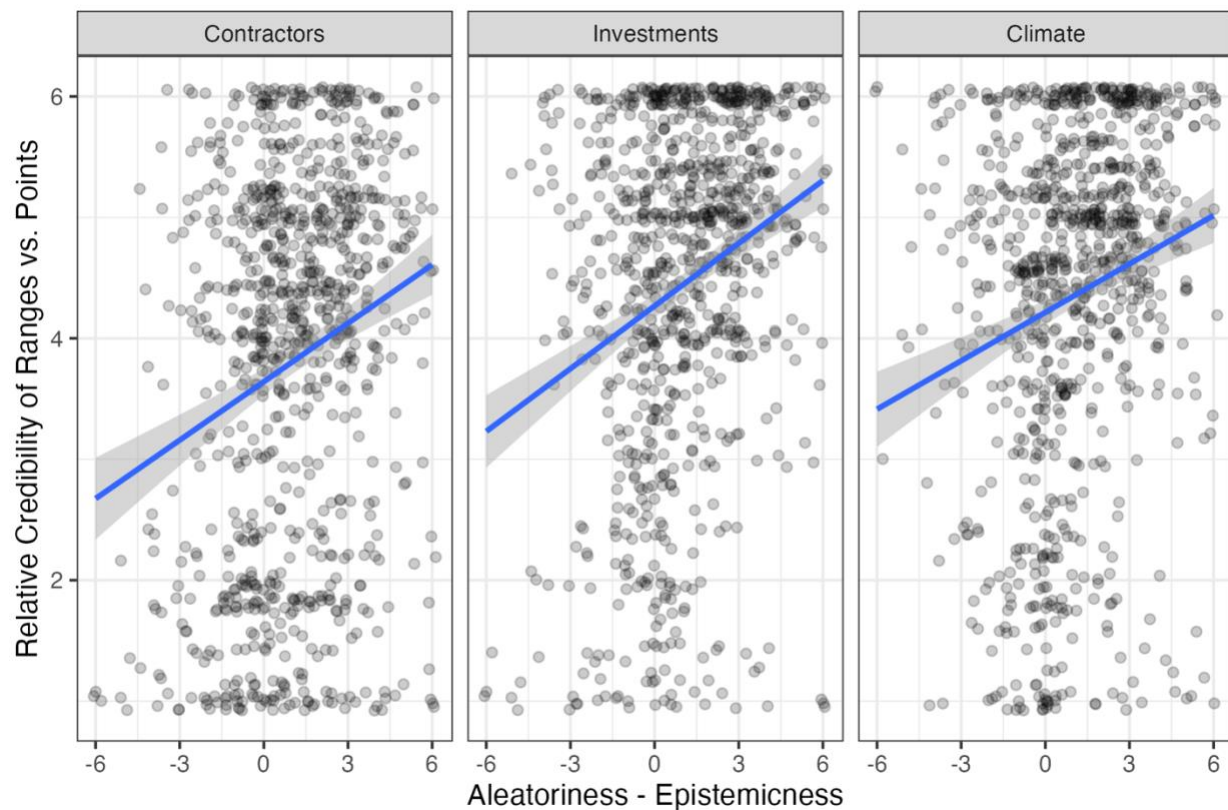


Figure 1B

Preference Results from Study 1. Each panel presents a scatterplot of respondents' perceptions of the nature of uncertainty (aleatoriness – epistemicness) against their relative preference ratings for judges providing less precise predictions (i.e., range predictions) versus more precise predictions (i.e., point predictions; higher scores indicate preferences for judges providing wider range predictions). Error bands surrounding the regression lines represent 95% confidence intervals.

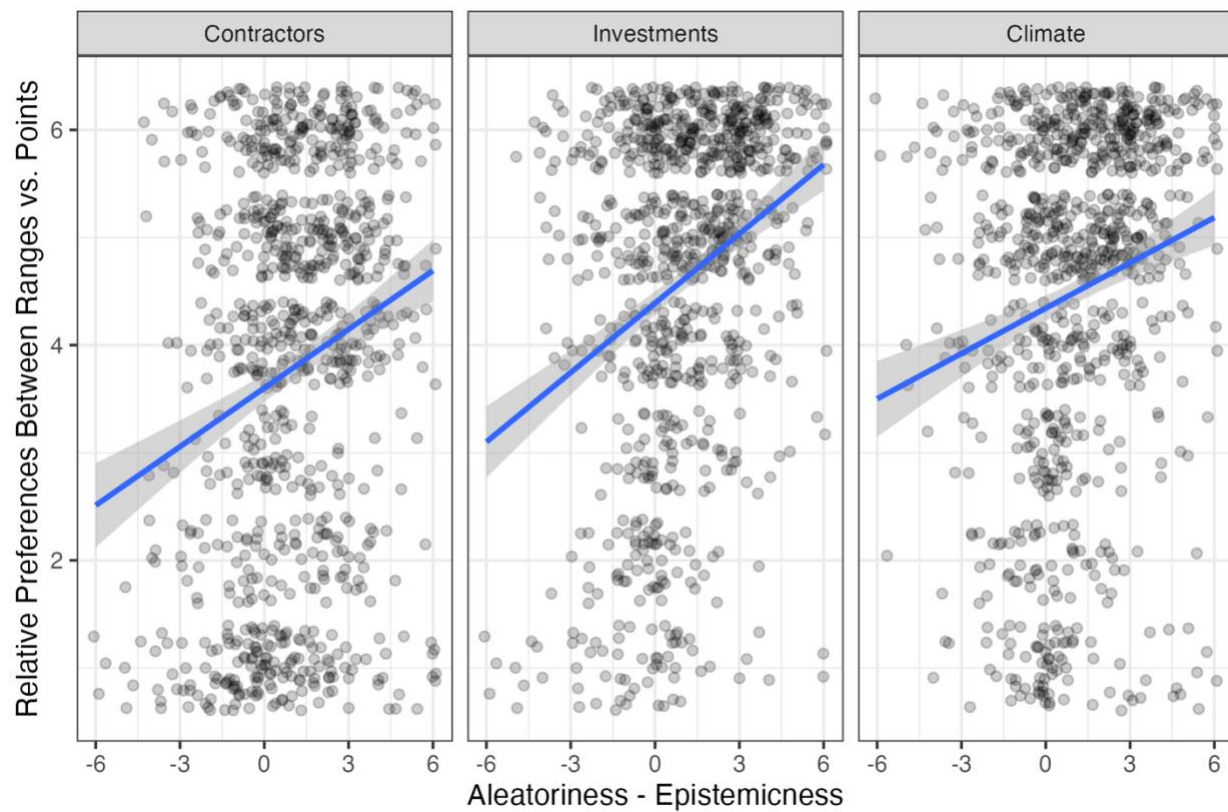


Figure 2A

“Absolute” Chart from Study 2.

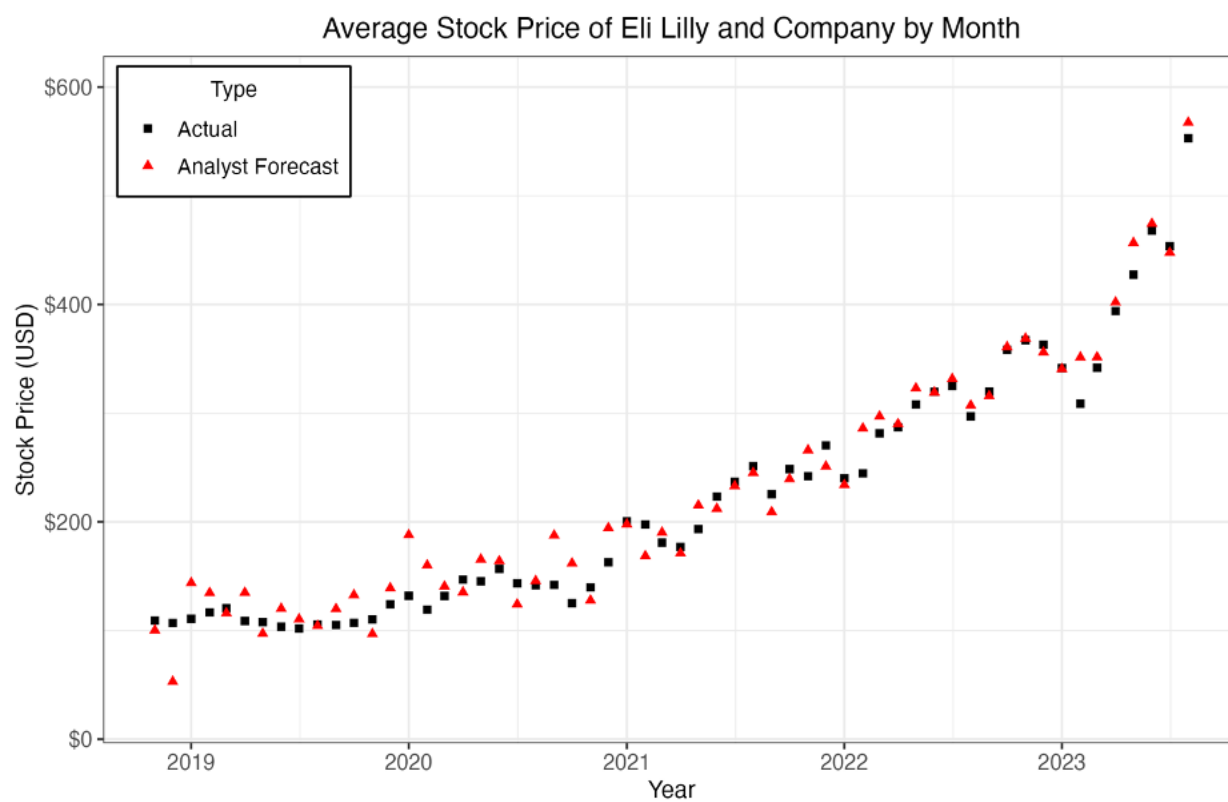


Figure 2B

“Relative” Chart from Study 2.

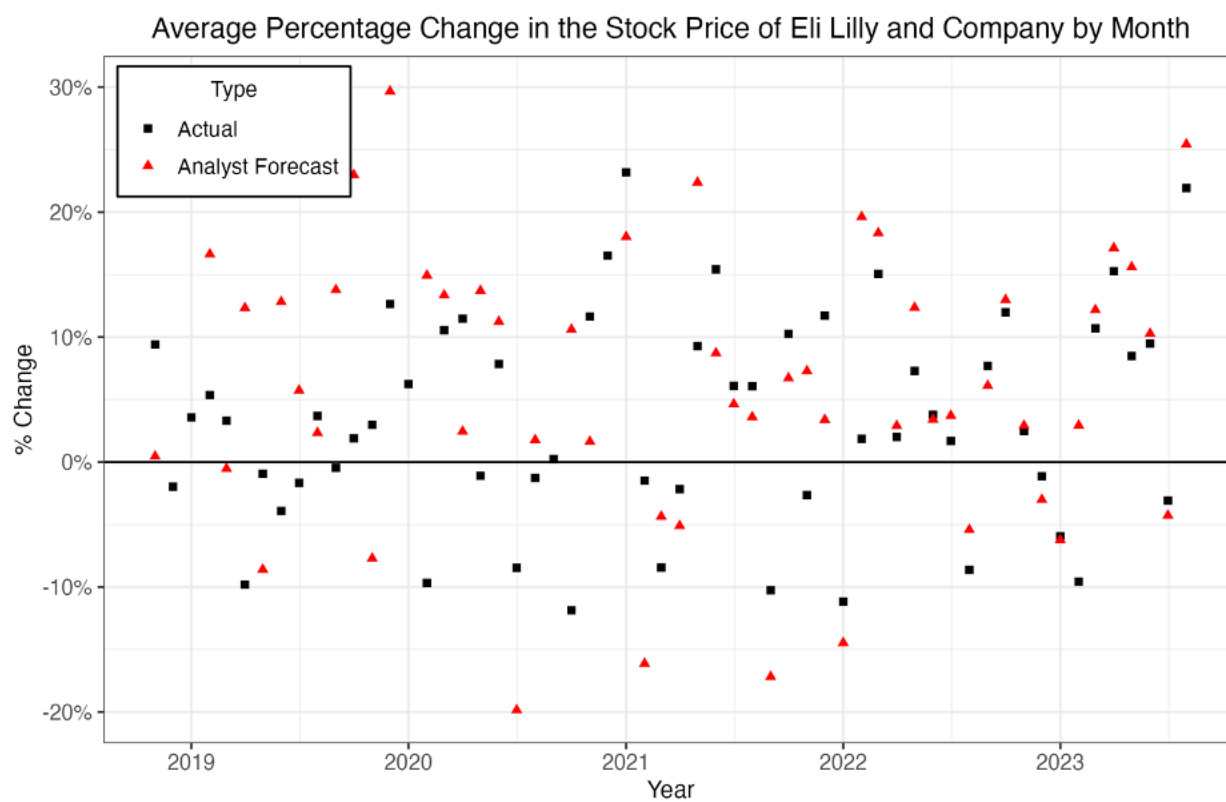
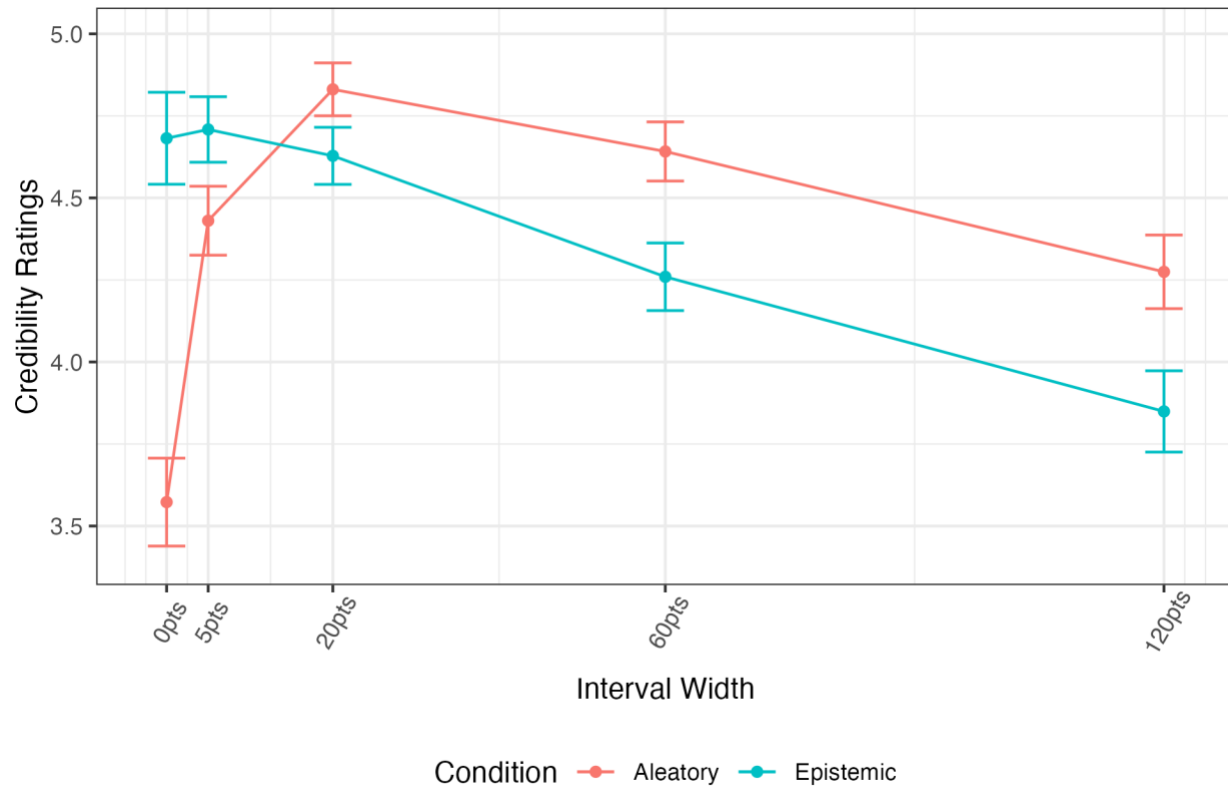


Figure 3

Results from Study 3. We plot mean credibility ratings for each hypothetical judge, by condition (i.e., game), as a function of interval width. Error bars represent standard errors.



Online Supplement:
Credibility of More vs. Less Precise Predictions Depends
on the Perceived Nature of Uncertainty

Table of Contents:

Supplement	Pages
1: Supplemental Study 1	35–41
2: Item-Level Analyses of Credibility Measures for Study 1	42–44
3: <i>Post Hoc</i> Analyses for Study 1 Treating Epistemicness and Aleatoriness as Separate Predictors	45–46
4: <i>Post Hoc</i> Placebo Tests for Study 1	47–49
5: Supplemental Study 2	50–52
6: Supplemental Study 3	53–54
7: Analyses Without Exclusions for Study 2	55–57
8: Supplemental Studies 4A/4B	58–65

Supplement 1: Supplemental Study 1

In Supplemental Study 1 we present initial evidence (reported in brief in the Introduction of the main text) that recipients interpret more versus less precise predictions in different ways as a function of their perceptions of the nature of uncertainty.

Participants

We recruited 300 participants from CloudResearch Connect ($M_{Age} = 41.95$, $SD_{Age} = 12.83$; 45.67% female, 52.33% male, 1.67% other, 0.33% prefer not to say).

Procedure

All participants considered six scenarios in a random order. In each scenario (reproduced fully in Table S1), participants were told that two experts about a particular topic had been asked to provide judgments about a relevant quantity (e.g., trivia buffs were asked to offer judgments about the height of Mount Everest, meteorologists were asked to predict rainfall amounts, gamblers were asked to offer judgments about winnings from playing a slot machine, and so on). Next, they were told that these experts had offered judgments using intervals of different widths (e.g., “They both offer ranges...centered on the same ‘best guess,’ but that differ in how wide they are...”), and we conveyed this visually using unmarked lines of different lengths (see Figure S1). On the same screen, participants were also asked: “If forced to choose, which do you think is the more plausible reason that these experts answered with ranges of different widths?” Their response options were either: (i) “The experts differ in how confident they are in their answers. Expert 2 has less confidence in their answer than Expert 1” (“confidence” interpretation) or (ii) “The experts differ in how widely they think the answer could vary. Expert 2 thinks the answer could vary more widely than Expert 1” (“variability” interpretation). Finally, on a separate screen, participants responded to the four-item Epistemic-Aleatory Rating Scale (EARS; Fox et

al., 2026; see Table S2) to indicate their perceptions of the nature of uncertainty in the relevant scenario. We randomized the order of the experts' predictions and all response options at the participant level.

Results and Discussion

For each scenario, we first binarized our key dependent measure by coding endorsements of the “confidence” interpretation with a value of one and by coding endorsements of the “variability” interpretation with a value of zero. Next, we constructed epistemicness and aleatoriness scores for each scenario by averaging the corresponding scale items from the EARS ($\alpha_{Epistemic} = 0.91$; $\alpha_{Aleatory} = 0.94$). We then subtracted the latter from the former to construct epistemic-aleatory difference scores for each scenario, with higher values indicating the perception that the uncertainty in a particular scenario was more epistemic and/or less aleatory. Finally, we conducted our primary preregistered analysis: a mixed effects logistic regression of responses to our key dependent variable upon epistemic-aleatory difference scores (with random intercepts for participant and scenario) to assess whether perceptions of the nature of uncertainty were associated with endorsing different interpretations of different levels of precision (interval widths).

As predicted, the extent to which respondents endorsed the “confidence” versus “variability” interpretations depended on the extent to which they saw the underlying uncertainty as more epistemic (and/or less aleatory). Specifically, we observed a significant positive association between the perceived nature of uncertainty (i.e., epistemic-aleatory difference scores) and the likelihood of endorsing the “confidence” versus “variability” interpretations in this omnibus analysis ($b = 0.154$, 95% CI = [0.102, 0.206], $z = 5.860$, $p < .001$; Odds Ratio = 1.17, 95% CI = [1.11, 1.23]).

Although our primary preregistered analysis entailed the foregoing omnibus analysis, we also analyzed our data at the scenario level. Specifically, we conducted logistic regressions of respondents' endorsements of the "confidence" versus "variability" interpretations upon epistemic-aleatory difference scores, scenario-by-scenario (see Table S3). These scenario-level results are more heterogeneous — which is not unexpected given that they rely solely on within-scenario variation in perceived nature of uncertainty. Nonetheless, they suggest that — at least to some extent — perceptions of the nature of uncertainty and different modes of interpreting precision (interval widths) may manifest across participants within-scenario, as well.

Altogether, these results suggest that perceptions of the nature of uncertainty may indeed be associated with interpreting expressed level of precision in different ways.

Table S1

Text of Scenarios Used in Supplemental Study 1. Each prompt was followed by the statement (edited for clarity by scenario) “They both offer ranges...centered on the same ‘best guess,’ but that differ in how wide they are. Specifically, here is what their answers look like...” and the visual depiction reproduced in Figure S1. We randomized the order of scenarios at the participant level.

Scenario	Prompt
Trivia Questions	You ask two trivia buffs the following question: “What is the height of Mount Everest?”
Stock Prices	You ask two skilled investors the following question about a particular stock: “What will the share price be when the market closes tomorrow?”
Basketball Scores	You ask two NBA basketball fans the following question about a particular player: “How many points will he score in tonight’s game?”
Rainfall Amounts	You ask two meteorologists the following question about a big storm that’s expected next week: “How many inches of rain will there be?”
Renovation Completion Times	You ask two contractors the following question about a particular kitchen renovation: “How many hours of labor will this job take to complete?”
Slot Machine Winnings	You ask two gamblers at a casino the following question: “How much money will I walk away with if I play this slot machine 100 times, at \$1 each?”

Table S2

Four-Item Epistemic-Aleatory Rating Scale (EARS; Fox et al., 2026). Participants responded on a 1–7 scale where 1 = Not at all and 7 = Very much. Item order was randomized at the participant level (i.e., order was constant across scenarios for each participant) and each time the scale was presented it was preceded by a stem asking participants to contemplate the nature of uncertainty associated with the event/prediction in question (see ResearchBox for details and materials).

Subscale	No.	Item
Epistemic	1	...is knowable in advance, given enough information.
	2	...is something that well-informed people would agree on.
Aleatory	1	...is something that has an element of randomness.
	2	...is determined by chance factors.

Table S3

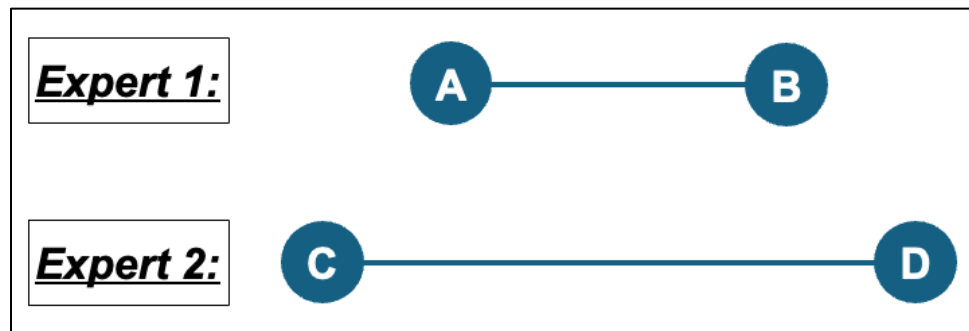
Secondary Analyses of Scenario-Level Results from Supplemental Study 1. Alongside our primary preregistered analysis, we also analyzed the results scenario-by-scenario. Specifically, we conducted logistic regressions of respondents' endorsements of the "confidence" versus "variability" interpretations of different interval widths upon epistemic-aleatory difference scores for each scenario, the results of which we report in the table below.

	Trivia Questions	Stock Prices	Basketball Scores	Rainfall Amounts	Renovation Completion Times	Slot Machine Winnings
Ep. - Al.	0.311***	0.076	-0.038	0.172**	0.227***	0.073
	[0.168, 0.460]	[-0.037, 0.191]	[-0.146, 0.069]	[0.050, 0.299]	[0.112, 0.347]	[-0.041, 0.187]
	$z = 4.208$	$z = 1.304$	$z = -0.687$	$z = 2.713$	$z = 3.783$	$z = 1.268$
	$p < .001$	$p = .192$	$p = .492$	$p = .007$	$p < .001$	$p = .205$
Log. Likelihood	-155.142	-193.987	-198.485	-178.545	-193.852	-182.464
N	300	300	300	300	300	300

+ $p < .1$, * $p < .05$, ** $p < .01$, *** $p < .001$

Figure S1

Visual Depiction of Interval Widths from Supplemental Study 1. We presented this image alongside each scenario in Supplemental Study 1 to illustrate the widths of experts' prediction intervals. We randomized whether Expert 1 (versus Expert 2) provided the narrower or wider prediction at the participant level using an alternate version of this image where the intervals depicted in each "row" were flipped.



Supplement 2: Item-Level Analyses of Credibility Measures for Study 1

Table S4

Item-Level Analyses for Study 1 (Contractors Scenario). Per our preregistration, we report ordinary least squares (OLS) regressions of each individual item from our credibility scale upon aleatory-epistemic difference scores in the table below, specifically for the contractors scenario. Note that only “credibility” measures are included since the choice measure is reported in the main text.

	Credible	Accurate	Careful Thought	Honest	Realistic
Al. - Ep.	0.285***	0.227***	0.335***	0.298***	0.294***
	[0.189, 0.382]	[0.126, 0.329]	[0.234, 0.437]	[0.200, 0.395]	[0.195, 0.392]
	$t = 5.808$	$t = 4.400$	$t = 6.510$	$t = 6.000$	$t = 5.865$
	$p < .001$	$p < .001$	$p < .001$	$p < .001$	$p < .001$
R^2	0.102	0.061	0.125	0.108	0.104
N	298	298	298	298	298

+ $p < .1$, * $p < .05$, ** $p < .01$, *** $p < .001$

Table S5

Item-Level Analyses for Study 1 (Investments Scenario). Per our preregistration, we report ordinary least squares (OLS) regressions of each individual item from our credibility scale upon aleatory-epistemic difference scores in the table below, specifically for the investments scenario. Note that only “credibility” measures are included since the choice measure is reported in the main text.

	Credible	Accurate	Careful Thought	Honest	Realistic
Al. - Ep.	0.342***	0.206***	0.284***	0.322***	0.335***
	[0.266, 0.418]	[0.116, 0.296]	[0.196, 0.372]	[0.244, 0.400]	[0.257, 0.413]
	$t = 8.833$	$t = 4.507$	$t = 6.383$	$t = 8.128$	$t = 8.463$
	$p < .001$	$p < .001$	$p < .001$	$p < .001$	$p < .001$
R^2	0.209	0.064	0.121	0.182	0.195
N	298	298	298	298	298

+ $p < .1$, * $p < .05$, ** $p < .01$, *** $p < .001$

Table S6

Item-Level Analyses for Study 1 (Climate Scenario). Per our preregistration, we report ordinary least squares (OLS) regressions of each individual item from our credibility scale upon aleatory-epistemic difference scores in the table below, specifically for the climate scenario. Note that only “credibility” measures are included since the choice measure is reported in the main text.

	Credible	Accurate	Careful Thought	Honest	Realistic
Al. - Ep.	0.288***	0.180***	0.287***	0.276***	0.253***
	[0.198, 0.378]	[0.083, 0.278]	[0.196, 0.378]	[0.188, 0.364]	[0.162, 0.344]
	$t = 6.319$	$t = 3.635$	$t = 6.195$	$t = 6.206$	$t = 5.480$
	$p < .001$	$p < .001$	$p < .001$	$p < .001$	$p < .001$
R^2	0.119	0.043	0.115	0.115	0.092
N	298	298	298	298	298

+ $p < .1$, * $p < .05$, ** $p < .01$, *** $p < .001$

**Supplement 3: *Post Hoc* Analyses for Study 1 Treating Epistemicness and
Aleatoriness as Separate Predictors**

Table S7

Results from Study 1 Treating Epistemicness and Aleatoriness as Separate Predictors (Credibility Measure). Below we report the results from ordinary least squares (OLS) regressions of our credibility measure upon epistemicness and aleatoriness for each scenario — these analyses differ from those reported in the main text (where, per our preregistration, we combine epistemicness and aleatoriness into a single aleatory-epistemic difference score). We thank an anonymous reviewer for suggesting these post hoc analyses.

	Contractors	Investments	Climate
Ep.	−0.441***	−0.273***	−0.213***
	[−0.593, −0.289]	[−0.385, −0.161]	[−0.338, −0.088]
	$t = -5.705$	$t = -4.809$	$t = -3.354$
	$p < .001$	$p < .001$	$p < .001$
Al.	0.168*	0.328***	0.312***
	[0.038, 0.299]	[0.199, 0.457]	[0.167, 0.456]
	$t = 2.544$	$t = 5.017$	$t = 4.244$
	$p = .011$	$p < .001$	$p < .001$
R^2	0.144	0.188	0.119
N	298	298	298

+ $p < .1$, * $p < .05$, ** $p < .01$, *** $p < .001$

Table S8

Results from Study 1 Treating Epistemicness and Aleatoriness as Separate Predictors (Preference Measure). Below we report the results from ordinary least squares (OLS) regressions of our preference measure upon epistemicness and aleatoriness for each scenario — these analyses differ from those reported in the main text (where, per our preregistration, we combine epistemicness and aleatoriness into a single aleatory-epistemic difference score). We thank an anonymous reviewer for suggesting these post hoc analyses.

	Contractors	Investments	Climate
Ep.	−0.500***	−0.329***	−0.166*
	[−0.676, −0.325]	[−0.451, −0.206]	[−0.309, −0.024]
	$t = -5.620$	$t = -5.281$	$t = -2.294$
	$p < .001$	$p < .001$	$p = .022$
Al.	0.211**	0.402***	0.399***
	[0.061, 0.361]	[0.261, 0.543]	[0.234, 0.564]
	$t = 2.765$	$t = 5.607$	$t = 4.769$
	$p = .006$	$p < .001$	$p < .001$
R^2	0.146	0.221	0.110
N	298	298	298

+ $p < .1$, * $p < .05$, ** $p < .01$, *** $p < .001$

Supplement 4: *Post Hoc* Placebo Tests for Study 1

Table S9

Statistical Tests for Domain Specificity of Relationships Between Aleatory-Epistemic Difference Scores and Credibility/Preference Ratings in Study 1. Below we report the results from a series of Steiger (1980) tests to probe whether the scenario-level correlations between aleatory-epistemic difference scores and credibility/preference ratings were significantly different from “mismatched” pairings of aleatory-epistemic difference scores and credibility/preference ratings across scenarios.

DV	A–E Domain	Mismatched DV Domain	<i>r</i> (matched)	<i>r</i> (mismatched)	Δr	<i>z</i>	<i>p</i>
Credibility	Contractors	Investments	0.357	0.223	0.134	2.195	.028
	Contractors	Climate	0.357	0.064	0.293	4.237	< .001
	Investments	Contractors	0.433	0.189	0.243	4.047	< .001
	Investments	Climate	0.433	0.278	0.154	2.735	.006
	Climate	Contractors	0.341	0.208	0.133	1.953	.051
	Climate	Investments	0.341	0.235	0.106	1.814	.070
Pref.	Contractors	Investments	0.363	0.238	0.125	2.045	.041
	Contractors	Climate	0.363	0.056	0.307	4.508	< .001
	Investments	Contractors	0.469	0.185	0.284	4.731	< .001
	Investments	Climate	0.469	0.262	0.207	3.582	< .001
	Climate	Contractors	0.315	0.188	0.127	1.881	.060
	Climate	Investments	0.315	0.284	0.032	0.524	.601

Figure S2

Matched and Mismatched Correlations for Credibility Ratings in Study 1. Below we visually depict the full set of “matched” and “mismatched” correlations between our credibility ratings and aleatory-epistemic ratings across the domains/scenarios considered in Study 1. Specifically, each cell displays the correlation coefficient between aleatory-epistemic difference scores and credibility ratings for a particular domain pairing, with darker blue shading indicating stronger positive correlations.

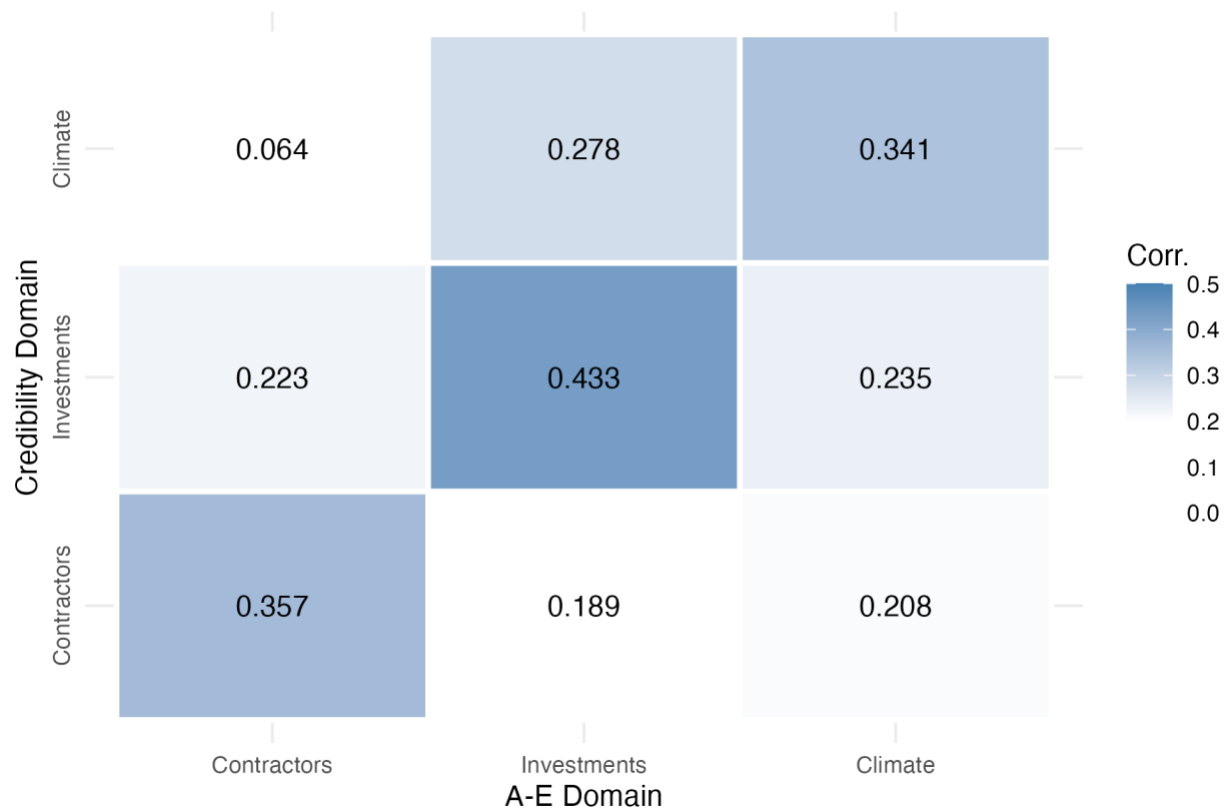
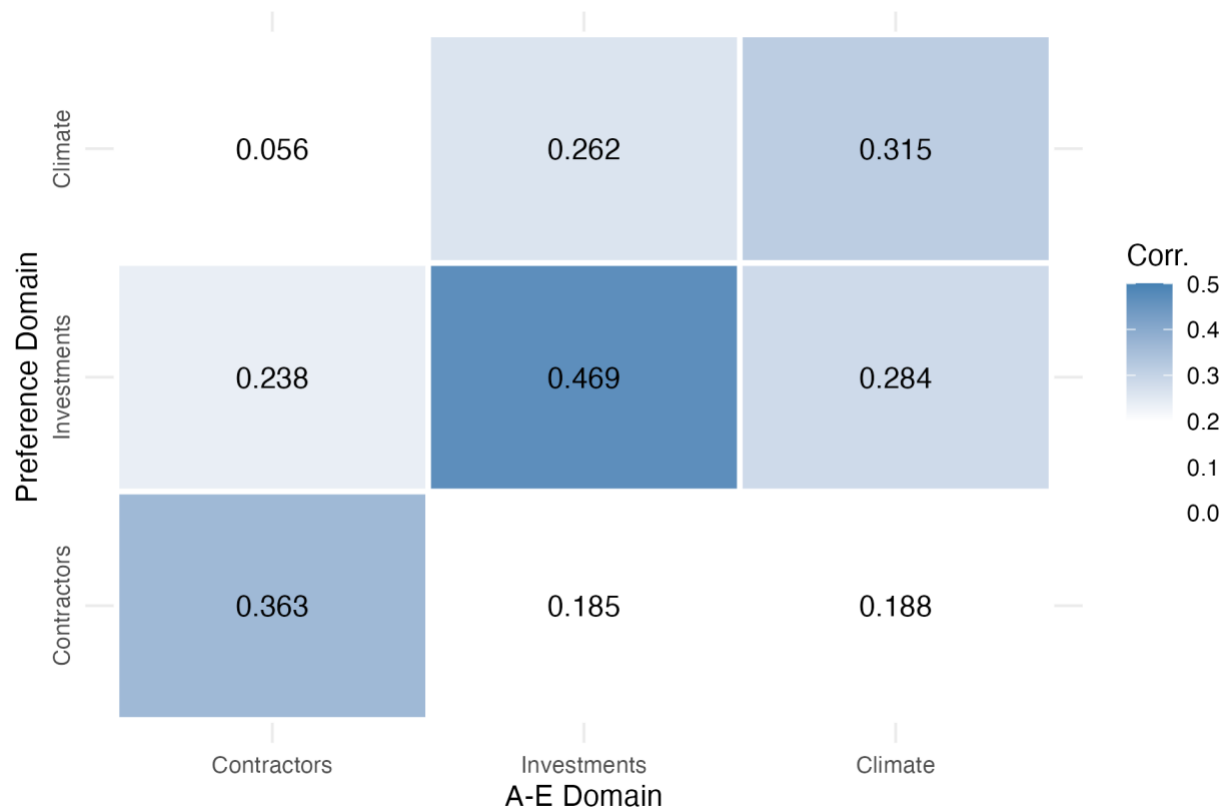


Figure S3

Matched and Mismatched Correlations for Preferences in Study 1. Below we visually depict the full set of “matched” and “mismatched” correlations between our preference ratings and aleatory-epistemic ratings across the domains/scenarios considered in Study 1. Specifically, each cell displays the correlation coefficient between aleatory-epistemic difference scores and preference ratings for a particular domain pairing, with darker blue shading indicating stronger positive correlations.



Supplement 5: Supplemental Study 2

In Supplemental Study 2 we present additional correlational evidence from an exploratory study probing the link between the perceived nature of uncertainty and the relative credibility of point versus range predictions. This study was similar to Study 1 (reported in the main text) but used slightly different measures and considered a different domain.

Participants

We recruited 99 participants from Prolific Academic ($M_{Age} = 36.54$, $SD_{Age} = 13.50$; 44.44% female, 53.54% male, 1.01% other, 1.01% prefer not to say).

Procedure

All participants considered a single scenario wherein one expert provided a point prediction about the outcome of an upcoming NFL game (“...Bengals will win by 4 points”), and another provided a range prediction (“...Bengals will win by 5-10 points”). Participants then rated, on 1–11 scales (1 = Expert A, who offered a point prediction; 11 = Expert B, who offered a range prediction), which expert was: (i) making a more credible claim; (ii) making a more believable claim; (iii) making a more reasonable claim; (iv) more knowledgeable; (v) more competent; and (vi) making a more accurate claim. Next, participants indicated their perceptions of the nature of uncertainty in this scenario using items adapted from the Epistemic-Aleatory Rating Scale (EARS; Fox et al., 2026; see Table S10 for adapted items).

Results and Discussion

First, we constructed a “credibility” score by averaging responses to our six dependent measures ($\alpha = 0.93$), with higher values indicating the perception that range predictions were relatively more credible. Next, we constructed aleatoriness and epistemicness scores by averaging the relevant subscale items ($\alpha_{Aleatory} = 0.79$; $\alpha_{Epistemic} = 0.83$). Finally, we constructed

an aleatory-epistemic difference score by subtracting epistemicness scores from aleatoriness scores. Thus, higher values indicated participants perceived uncertainty in this scenario to be more aleatory and/or less epistemic. Finally, we conducted an ordinary least squares (OLS) regression to assess the relationship between the perceived nature of uncertainty and relative credibility judgments.

As predicted, the extent to which participants judged the expert offering the range prediction to be more credible than the expert providing the point prediction depended on the extent to which respondents saw the underlying uncertainty as more aleatory (and/or less epistemic). Specifically, we observed a positive association between the perceived nature of uncertainty (i.e., aleatory-epistemic difference scores) and the credibility of experts offering range versus point predictions that was directionally consistent with our hypothesis ($b = 0.257$, 95% CI = $[-0.007, 0.521]$, $t(97) = 1.930$, $p = .056$).

Table S10

Epistemicness and Aleatoriness Measures Used in Supplemental Study 2 (adapted from Fox et al., 2026). Participants responded on a 1–7 scale where 1 = Not at all and 7 = Very much. Item order was randomized at the subscale level, and each subscale was preceded by a stem asking participants to contemplate the nature of uncertainty associated with NFL game scores (see ResearchBox for details and materials).

Subscale	No.	Item
Epistemic	1	...is in principle knowable in advance.
	2	...is knowable in advance, given enough information.
	3	...is something that well-informed people would agree on.
	4	...is something that could be better predicted by consulting an expert.
	5	...is something that becomes more predictable with additional knowledge or skills.
Aleatory	1	...is something that has an element of randomness.
	2	...feels unpredictable.
	3	...feels like it is determined by chance factors.
	4	...feels like it could play out in different ways on similar occasions.

Supplement 6: Supplemental Study 3

Supplemental Study 3 provided additional correlational evidence for a link between perceptions of the nature of uncertainty and preferences for point versus range predictions in the presence of incentives.

Participants

We recruited 400 participants from Prolific Academic. Due to a programming oversight, we did not obtain demographic information from participants.

Procedure

Upon entering the study, participants completed a filler task wherein they had to complete various word fragments (e.g., “S T R E _ _”). After completing this filler task, participants worked on the main task in this study (for our purposes). In this task, participants were awarded a \$1.00 bonus that they could either keep or donate (in part or in total) to a charity devoted to providing access to clean water. Importantly, between-subjects, we manipulated whether this charity’s potential impact was offered as a point prediction (“...every dollar donated...will pay for 1,100 liters of clean water”) or as a range prediction (“...every dollar donated...will pay for between 900 and 1,300 liters of clean water”). After encountering the relevant prediction, participants indicated how much of their bonus they might want to donate to charity on a sliding scale. Finally, on a separate screen, participants rated their perceptions of the nature of uncertainty associated with predicting this charity’s impact using the same measures as in Supplemental Study 2 (see Table S10).

Results and Discussion

First, we constructed aleatoriness and epistemicness scores by averaging the relevant subscale items ($\alpha_{Aleatory} = 0.92$; $\alpha_{Epistemic} = 0.87$). Next, we constructed an aleatory-epistemic

difference score by subtracting epistemicness scores from aleatoriness scores. Thus, higher values indicated participants perceived uncertainty in this scenario to be more aleatory and/or less epistemic. Additionally, we constructed a condition dummy variable (0 = Range Prediction; 1 = Point Prediction) to indicate how the charity's impact was framed to participants. Finally, we conducted an OLS regression of donations upon condition, aleatory-epistemic difference scores, and their interaction to probe whether different levels of precision (i.e., point versus range predictions) yielded different donation amounts depending on perceptions of the nature of uncertainty.

As predicted, we observed a significant interaction between condition (i.e., which prediction format respondents encountered) and aleatory-epistemic difference scores ($b = -0.064$, 95% CI = $[-0.103, -0.025]$, $t(396) = -3.244$, $p = .001$). This suggests that the effectiveness of more versus less precise predictions in spurring donations depends on perceptions of the nature of uncertainty. A floodlight analysis sheds further light on this pattern for participants who viewed uncertainty as more aleatory (and/or less epistemic), and vice versa. Specifically, for participants who viewed uncertainty as more aleatory and/or less epistemic (i.e., aleatory-epistemic difference scores = 0.786, +1 SD), point predictions yielded significantly smaller donations than range predictions ($b = -0.113$, 95% CI = $[-0.213, -0.014]$, $t(396) = -2.240$, $p = .026$). In contrast, for participants who viewed uncertainty as more epistemic and/or less aleatory (i.e., aleatory-epistemic difference scores = -2.844 , -1 SD), point predictions yielded significantly larger donations than range predictions ($b = 0.119$, 95% CI = $[0.020, 0.218]$, $t(396) = 2.358$, $p = .019$).

Supplement 7: Analyses Without Exclusions for Study 2

Table S11

Results from Study 2 Manipulation Check (No Exclusions). Below we report the results of our manipulation check from Study 2 without excluding any participants who failed our comprehension and/or attention checks.

	Aleatoriness–Epistemicness	
Condition	–0.530*** [–0.654, –0.406] $t = -8.366$ $p < .001$	–0.530*** [–0.655, –0.406] $t = -8.363$ $p < .001$
Format		0.059 [–0.066, 0.183] $t = 0.925$ $p = .355$
Condition × Format		0.009 [–0.116, 0.133] $t = 0.135$ $p = .892$
R^2	0.065	0.066
N	1,002	1,002

+ $p < .1$, * $p < .05$, ** $p < .01$, *** $p < .001$

Table S12

Results from Study 2 Credibility Measures (No Exclusions). Below we report the results concerning responses to our credibility measures from Study 2 without excluding any participants who failed our comprehension and/or attention checks.

	Credibility	
Condition	−0.094**	−0.094**
	[−0.164, −0.024]	[−0.164, −0.024]
	$t = -2.638$	$t = -2.638$
	$p = .008$	$p = .008$
Format		0.022
		[−0.048, 0.092]
		$t = 0.613$
		$p = .540$
Condition × Format		−0.004
		[−0.074, 0.066]
		$t = -0.124$
		$p = .901$
R^2	0.007	0.007
N	1,002	1,002

+ $p < .1$, * $p < .05$, ** $p < .01$, *** $p < .001$

Table S13

Results from Study 2 Preferences Measure (No Exclusions). Below we report the results concerning responses to our preferences measure (i.e., hiring choices) from Study 2 without excluding any participants who failed our comprehension and/or attention checks.

	Hiring Choices	
Condition	–0.074+	–0.074+
	[–0.149, 0.000]	[–0.149, 0.000]
	$t = -1.954$	$t = -1.953$
	$p = .051$	$p = .051$
Format		0.030
		[–0.044, 0.105]
		$t = 0.798$
		$p = .425$
Condition × Format		0.013
		[–0.061, 0.088]
		$t = 0.354$
		$p = .723$
R^2	0.004	0.005
N	1,002	1,002

+ $p < .1$, * $p < .05$, ** $p < .01$, *** $p < .001$

Supplement 8: Supplemental Studies 4A/4B

In Supplemental Studies 4A and 4B, we sought additional causal evidence for our hypothesis by using an alternate manipulation from the one used in the main text. Specifically, in each supplemental study, we used an “article manipulation” wherein we asked participants to read an article describing uncertainty in the relevant contexts in terms highlighting epistemicness or aleatoriness, after which they encountered (and rated) a corresponding set of interval predictions.

Supplemental Study 4A

Participants

We recruited 400 CloudResearch-approved participants from Amazon Mechanical Turk (MTurk). We excluded participants who failed attention checks, yielding $N = 367$ participants ($M_{Age} = 42.18$, $SD_{Age} = 12.42$; 49.86% female, 49.86% male, 0.27% other).

Procedure

In Study 4A we asked participants to consider predictions regarding the size of legal settlements from traffic accident cases. We randomly assigned participants to the “epistemic” or “aleatory” condition. Across conditions, participants read a short article describing the uncertainty associated with settlements from a traffic accident case in terms highlighting either epistemic uncertainty (e.g., “...the size of the settlement is largely determined by past precedent...”) or aleatory uncertainty (e.g., “...the size of the settlement is largely determined by how the legal process unfolds...”; see ResearchBox for materials). After reading the relevant article, we asked participants to briefly summarize what they had learned about the process of pursuing a legal settlement in a traffic accident case in no more than a few sentences. Next, on a separate screen, participants encountered one law firm who provided a point prediction regarding

the potential size of a settlement in a traffic accident case (“...you should expect to receive a settlement of \$5,000”) and another who provided a range prediction (“...you should expect to receive a settlement of \$4,000 to \$6,000”). On the same screen, participants then rated the relative credibility of each firm on a five-item scale (adapted from Gaertig & Simmons, 2018; these are the same measures used in Study 1 in the main text) and indicated their preferences between them (see Table S14 for details). Finally, on yet another screen, participants responded to the six-item EARS (Fox et al., 2026; see Table S15 for details) as a manipulation check.

Results and Discussion

We first made several transformations by constructing: (i) a credibility score by averaging responses to the five credibility items ($\alpha = 0.90$); (ii) an aleatory-epistemic difference score by averaging the aleatoriness and epistemicness items on the EARS ($\alpha_{Aleatory} = 0.75$; $\alpha_{Epistemic} = 0.70$) and taking the difference between these scores; and (iii) a contrast-coded variable for condition ($-1 = \text{Aleatory}$; $+1 = \text{Epistemic}$).

We first probed our manipulation check. As predicted, participants who reviewed the article describing the uncertainty associated with legal settlements for traffic accident cases in terms highlighting epistemicness indeed viewed the uncertainty associated with predictions in this domain to be more epistemic (and/or less aleatory; $M = -0.625$, 95% CI = $[-0.897, -0.353]$) than those who reviewed the article highlighting aleatoriness ($M = 0.639$, 95% CI = $[0.365, 0.913]$). An OLS regression of aleatory-epistemic difference scores on condition revealed this difference was significant ($b = -0.632$, 95% CI = $[-0.825, -0.440]$, $t(365) = -6.460$, $p < .001$).

We next probed credibility. As predicted, participants who reviewed the article highlighting epistemic uncertainty judged firms providing point predictions ($M = 4.126$, 95% CI = $[3.939, 4.313]$) to be relatively more credible than those who reviewed the article highlighting

aleatory uncertainty ($M = 4.368$, 95% CI = [4.178, 4.558]). An OLS regression of credibility on condition revealed this difference was directionally consistent with our hypothesis ($b = -0.121$, 95% CI = [-0.254, 0.012], $t(365) = -1.792$, $p = .074$). Given that visual inspection of these data suggested marked skewness in respondents' credibility ratings, we also conducted a nonparametric test. This alternative approach yielded similar results ($W = 18784$, $p = .055$).

Preference ratings followed the same pattern: respondents who read the article highlighting epistemic uncertainty favored firms providing point predictions ($M = 3.978$, 95% CI = [3.755, 4.201]) as compared to those who read the article highlighting aleatory uncertainty ($M = 4.284$, 95% CI = [4.062, 4.506]). An OLS regression of preference ratings on condition revealed this difference was also directionally consistent with our hypothesis ($b = -0.153$, 95% CI = [-0.310, 0.004], $t(365) = -1.919$, $p = .056$). Once more, given that visual inspection of these data suggested marked skewness in respondents' preference ratings, we also conducted a nonparametric test. This alternative approach suggested that the effect of condition on preference ratings might be significant ($W = 19024$, $p = .027$).

Supplemental Study 4B

Participants

We recruited 600 participants from Prolific Academic. We excluded participants who failed an attention check, yielding $N = 592$ participants ($M_{Age} = 41.38$, $SD_{Age} = 14.18$; 44.59% female, 54.05% male, 1.35% other).

Procedure

Study 4B was nearly identical to Study 4A, but in Study 4B we asked participants to consider predictions regarding the amount of weight one might lose by adopting a new personal health regimen with the help of a trainer. Once again, we randomly assigned participants to the

“epistemic” or “aleatory” condition, and across conditions, participants read a short article describing the uncertainty associated with weight loss in terms highlighting either epistemic uncertainty (e.g., “...how much weight [one] can expect to lose...is primarily determined by how well people adhere to the trainer’s specified routine”) or aleatory uncertainty (e.g., “...how much weight [one] can expect to lose...is primarily determined by differences in the body’s response to changes in food intake and exercise, which vary widely from person-to-person”; see ResearchBox for materials). After reading the relevant article, we asked participants to briefly summarize what they had learned about the process of losing weight with the help of a trainer in no more than a few sentences. Next, on a separate screen, participants encountered one trainer who provided a point prediction regarding weight loss (i.e., the point prediction “...you should expect to lose 20 lbs”) and another who provided a range prediction (i.e., the range prediction “...you should expect to lose 16 lbs to 24 lbs”). On the same screen, participants then rated the relative credibility of each trainer and indicated their preferences between them using the same measures from Study 4A (see Table S14). Finally, on a separate screen, participants responded to the six-item EARS (Fox et al., 2026; see Table S15 for details) as a manipulation check.

Results and Discussion

As in Study 4A, we first made several transformations by constructing: (i) a credibility score by averaging responses to the five credibility items ($\alpha = 0.92$); (ii) an aleatory-epistemic difference score by averaging the aleatoriness and epistemicness items on the EARS ($\alpha_{Aleatory} = 0.73$; $\alpha_{Epistemic} = 0.82$) and taking the difference between these scores; and (iii) a contrast-coded variable for condition ($-1 = \text{Aleatory}$; $+1 = \text{Epistemic}$).

We first probed our manipulation check. As predicted, participants who reviewed the article describing the uncertainty associated with weight loss in terms highlighting epistemicness

indeed viewed the uncertainty associated with predictions in this domain to be more epistemic (and/or less aleatory; $M = -0.760$, 95% CI = $[-0.975, -0.544]$) than those who reviewed the article highlighting aleatoriness ($M = 0.743$, 95% CI = $[0.527, 0.959]$). An OLS regression of aleatory-epistemic difference scores on condition revealed this difference was significant ($b = -0.751$, 95% CI = $[-0.903, -0.599]$, $t(590) = -9.687$, $p < .001$).

We next probed credibility. As predicted, participants who reviewed the article highlighting epistemic uncertainty judged trainers providing point predictions ($M = 4.979$, 95% CI = $[4.853, 5.105]$) to be relatively more credible than those who reviewed the article highlighting aleatory uncertainty ($M = 5.132$, 95% CI = $[5.009, 5.254]$). An OLS regression of credibility on condition revealed this difference was directionally consistent with our hypothesis ($b = -0.076$, 95% CI = $[-0.164, 0.011]$, $t(590) = -1.711$, $p = .088$). Given that visual inspection of these data suggested marked skewness in respondents' credibility ratings, we also conducted a nonparametric test. This alternative approach suggested that the effect of condition on credibility ratings might be significant ($W = 48406$, $p = .024$).

Preference ratings followed the same pattern: respondents who read the article highlighting epistemic uncertainty favored trainers providing point predictions ($M = 4.735$, 95% CI = $[4.580, 4.889]$) as compared to those who read the article highlighting aleatory uncertainty ($M = 5.010$, 95% CI = $[4.872, 5.148]$). An OLS regression of preference ratings on condition revealed this difference was significant and consistent with our hypothesis ($b = -0.138$, 95% CI = $[-0.241, -0.034]$, $t(590) = -2.618$, $p = .009$). Once more, given that visual inspection of these data suggested marked skewness in respondents' preference ratings, and for consistency with our approach to analyzing other dependent measures in Studies 4A and 4B, we also conducted a

nonparametric test. This alternative approach similarly suggested that the effect of condition on preference ratings was significant ($W = 49422, p = .004$).

Table S14

Dependent Measures Used in Supplemental Studies 4A and 4B. Participants responded on a 1–6 scale where 1 = Definitely Advisor A and 6 = Definitely Advisor B, and the term “expert” was replaced with the relevant professional title. We counterbalanced whether expert “A” or “B” provided a point prediction or a range prediction, and when analyzing the data we ensured that higher ratings on these measures always corresponded to range predictions.

No.	Type	Item
1	Credibility	Which [expert] provided a more credible estimate?
2	Credibility	Which [expert] provided a more accurate estimate?
3	Credibility	Which [expert] provided an estimate reflecting more careful thought?
4	Credibility	Which [expert] provided a more honest estimate?
5	Credibility	Which [expert] provided a more realistic estimate?
6	Preference	Which [expert] would you rather hire?

Table S15

Six-Item Epistemic-Aleatory Rating Scale (EARS; Fox et al., 2026). Participants responded on a 1–7 scale where 1 = Not at all and 7 = Very much. Item order was randomized and each time the scale was presented it was preceded by a stem asking participants to contemplate the nature of uncertainty associated with the event/prediction in question (see ResearchBox for details and materials).

Subscale	No.	Item
Epistemic	1	...is knowable in advance, given enough information.
	2	...is something that becomes predictable with additional knowledge or skills.
	3	...is something that well-informed people would agree on.
Aleatory	1	...is determined by chance factors.
	2	...could play out in different ways on similar occasions.
	3	...is something that has an element of randomness.

References

- Fox, C. R., Tannenbaum, D., Ülkümen, G., Walters, D. J., & Erner, C. (2026). *Credit, blame, luck, and the perceived nature of uncertainty* [Unpublished Working Paper].
- Gaertig, C., & Simmons, J. P. (2018). Do People Inherently Dislike Uncertain Advice? *Psychological Science*, 29(4), 504–520. <https://doi.org/10.1177/0956797617739369>
- Steiger, J. H. (1980). Tests for comparing elements of a correlation matrix. *Psychological Bulletin*, 87(2), 245–251. <https://doi.org/10.1037/0033-2909.87.2.245>